



Técnicas de análisis multivariante aplicadas a las ciencias sociales

Volumen I

Demian Ábrego Almazán
José Melchor Medina Quintero
Yesenia Sánchez Tovar
Manuel H. de la Garza Cárdenas



UAT

VERDAD, BELLEZA, PROBIDAD



**Técnicas de análisis multivariante
aplicadas a las ciencias sociales
Volumen I**

Ábrego Almazán, Demian

Técnicas de análisis multivariante aplicadas a las ciencias sociales Volumen I / Demian Ábrego Almazán, José Melchor Medina Quintero, Yesenia Sánchez Tovar, Manuel H. de la Garza Cárdenas. — Cd. Victoria, Tamaulipas : Universidad Autónoma de Tamaulipas ; Ciudad de México : Colofón , 2021.
146 páginas : 17 x 23 cm

1. Ciencias sociales – Métodos estadísticos. 2. Modelo de ecuaciones estructurales. 3. Mínimos Cuadrados. I. Medina Quintero, José Melchor, autor. II. Sánchez Tovar, Yesenia, autor. III. Garza Cárdenas, Manuel H. de la, autor. IV. Título.

LC: **HA29.5**

Dewey: **001.422 072 7**

Consejo de Publicaciones UAT

Tel. (52) 834 3181-800 • extensión: 2948 • www.uat.edu.mx

Centro Universitario Victoria

Centro de Gestión del Conocimiento. Tercer Piso

Cd. Victoria, Tamaulipas, México. C.P. 87149

consejopublicacionesuat@outlook.com



Fomento Editorial Una edición del Departamento de Fomento Editorial de la Universidad Autónoma de Tamaulipas

D. R. © 2021 Universidad Autónoma de Tamaulipas

Matamoros SN, Zona Centro Ciudad Victoria, Tamaulipas C.P. 87000

Edificio Administrativo, planta baja, CU Victoria

Ciudad Victoria, Tamaulipas, México

Libro aprobado por el Consejo de Publicaciones UAT

ISBN UAT: 978-607-8750-60-3

Colofón S.A. de C.V.

Franz Hals núm. 130, Alfonso XIII

Delegación Álvaro Obregón C.P. 01460, Ciudad de México

www.colofonlibros.com • colofonedicionesacademicas@gmail.com

ISBN Colofón: 978-607-635-236-6

Se prohíbe la reproducción total o parcial de esta obra incluido el diseño tipográfico y de portada, sea cual fuere el medio, electrónico o mecánico, sin el consentimiento por escrito del Consejo de Publicaciones UAT.

Impreso en México • *Printed in Mexico*

El tiraje consta de 400 ejemplares

Este libro fue dictaminado y aprobado por la Dra. Guadalupe del Carmen Briano Turrent, profesor investigador en la Universidad Autónoma de San Luis Potos, integrante del Sistema Nacional de Investigadores nivel I y a la Dra. Osiris María Echeverría Ríos, profesor investigador en la Universidad Politécnica Metropolitana de Hidalgo, candidata al Sistema Nacional de Investigadores. Asimismo, fue recibido por el Comité Interno de Selección de Obras de Colofón Ediciones Académicas para su valoración en la sesión del segundo semestre 2020 donde se sometió al sistema de dictaminación a “doble ciego” por especialistas en la materia, el resultado de los dictámenes fue positivo.

"PARA CREAR COSAS BUENAS
PRIMERO HAY QUE CREER
EN ELLAS"



UNIVERSIDAD
AUTÓNOMA DE
TAMAULIPAS
—1950-2020—

Técnicas de análisis multivariante aplicadas a las ciencias sociales Volumen I

Demian Ábrego Almazán
José Melchor Medina Quintero
Yesenia Sánchez Tovar
Manuel H. de la Garza Cárdenas





Ing. José Andrés Suárez Fernández
PRESIDENTE

Dr. Julio Martínez Burnes
VICEPRESIDENTE

Dr. Héctor Manuel Cappello Y García
SECRETARIO TÉCNICO

C.P. Guillermo Mendoza Cavazos
VOCAL

Dra. Rosa Issel Acosta González
VOCAL

Lic. Víctor Hugo Guerra García
VOCAL

Consejo Editorial del Consejo de Publicaciones de la Universidad Autónoma de Tamaulipas

Dra. Lourdes Arizpe Slogher • Universidad Nacional Autónoma de México | **Dr. Amalio Blanco** • Universidad Autónoma de Madrid, España | **Dra. Rosalba Casas Guerrero** • Universidad Nacional Autónoma de México | **Dr. Francisco Díaz Bretones** • Universidad de Granada, España | **Dr. Rolando Díaz Lowing** • Universidad Nacional Autónoma de México | **Dr. Manuel Fernández Ríos** • Universidad Autónoma de Madrid, España | **Dr. Manuel Fernández Navarro** • Universidad Autónoma Metropolitana, México | **Dra. Juana Juárez Romero** • Universidad Autónoma Metropolitana, México | **Dr. Manuel Marín Sánchez** • Universidad de Sevilla, España | **Dr. Cervando Martínez** • University of Texas at San Antonio, E.U.A. | **Dr. Darío Páez** • Universidad del País Vasco, España | **Dra. María Cristina Puga Espinosa** • Universidad Nacional Autónoma de México | **Dr. Luis Arturo Rivas Tovar** • Instituto Politécnico Nacional, México | **Dr. Aroldo Rodrigues** • University of California at Fresno, E.U.A. | **Dr. José Manuel Valenzuela Arce** • Colegio de la Frontera Norte, México | **Dra. Margarita Velázquez Gutiérrez** • Universidad Nacional Autónoma de México | **Dr. José Manuel Sabucedo Comeselle** • Universidad de Santiago de Compostela, España | **Dr. Alessandro Soares da Silva** • Universidad de São Paulo, Brasil | **Dr. Akexandre Dorna** • Universidad de CAEN, Francia | **Dr. Ismael Vidales Delgado** • Universidad Regiomontana, México | **Dr. José Francisco Zúñiga García** • Universidad de Granada, España | **Dr. Bernardo Jiménez** • Universidad de Guadalajara, México | **Dr. Juan Enrique Marciano Medina** • Universidad de Puerto Rico-Humacao | **Dra. Ursula Oswald** • Universidad Nacional Autónoma de México | **Arq. Carlos Mario Yori** • Universidad Nacional de Colombia | **Arq. Walter Debenedetti** • Universidad de Patrimonio, Colonia, Uruguay | **Dr. Andrés Piqueras** • Universitat Jaume I, Valencia, España | **Dr. Yolanda Troyano Rodríguez** • Universidad de Sevilla, España | **Dra. María Lucero Guzmán Jiménez** • Universidad Nacional Autónoma de México | **Dra. Patricia González Aldea** • Universidad Carlos III de Madrid, España | **Dr. Marcelo Urrea** • Revista Latinoamericana de Psicología Social | **Dr. Rubén Ardila** • Universidad Nacional de Colombia | **Dr. Jorge Gissi** • Pontificia Universidad Católica de Chile | **Dr. Julio F. Villegas** • Universidad Diego Portales, Chile | **Ángel Bonifaz Ezeta** • Universidad Nacional Autónoma de México

Índice

Prólogo.....	9
Capítulo 1. Introducción a las técnicas de análisis multivariante.....	11
1.1. Análisis multivariante de datos.....	11
1.2. Datos.....	12
1.3. Escalas de medidas.....	14
1.3.1. Escalas no métricas.....	14
1.3.2. Escalas métricas.....	15
1.4. Clasificación de las técnicas de análisis multivariante.....	15
1.4.1. Análisis de interdependencia.....	16
1.4.2. Análisis de dependencia.....	17
1.5. Conceptos complementarios.....	20
Capítulo 2. Análisis multivariante mediante Modelado de Ecuaciones Estructurales.....	23
2.1. Introducción.....	23
2.2. ¿Qué es la modelización de ecuaciones estructurales?.....	23
2.3. Tipo de modelización.....	26
2.4. Tipos de variables.....	27
2.4.1 Construcción de Modelos.....	29
2.4.2. Jerarquía de los componentes.....	30
2.5. Mediciones.....	32
2.6. Escala de medida.....	33
2.7. Tipos de codificación de datos y construcción de la escala.....	33
2.8. Distribución de datos.....	36
Capítulo 3. Supuestos básicos para la ejecución de la técnica de ecuaciones estructurales.....	39
3.1. Introducción.....	39
3.2. Linealidad.....	39
3.3. Normalidad.....	41
3.3.1. Normalidad univariante.....	41
3.3.2. Normalidad multivariante.....	43
3.4. Multicolinealidad.....	44
3.5. Valores perdidos.....	46
3.6. Valores atípicos.....	47

Capítulo 4. Modelización de ecuaciones estructurales basada en la covarianza..	49
4.1 Introducción.....	49
4.2. Elección del método SEM por covarianzas.....	50
4.3. Tipos de relaciones en los modelos estructurales.....	52
4.3.1. Causalidad.....	52
4.3.2. Covarianza o correlación.....	53
4.3.3. Relación espuria.....	54
4.3.4. Relación causal directa e indirecta y efectos totales.....	55
4.4. Métodos de estimación en los modelos por covarianza.....	57
4.4.1. Máxima verosimilitud (ML).....	57
4.4.2. Mínimos cuadrados generalizados (GLS).....	59
4.4.3. Mínimos cuadrados no ponderados (ULS).....	60
4.4.4. Distribución Asintóticamente libre (ADF).....	60
4.5. Etapas para la construcción de SEM por covarianzas.....	60
4.5.1. Paso 1. Especificación teórica del modelo de medida y del modelo estructural.....	61
4.5.2. Paso 2. Valoración del modelo de medida.....	62
4.5.2.1. Fiabilidad interna.....	63
4.5.2.2. Validez.....	66
4.5.3. Paso 3. Valoración del modelo estructural.....	68
4.5.3.3 Bondad de ajuste.....	69
Capítulo 5. Aplicación práctica de ecuaciones estructurales con base en la covarianza.....	73
5.1.Introducción.....	73
5.3. Configuración de parámetros para valoración de un modelo en el software AMOS.....	83
5.4. Evaluación del modelo de medida.....	86
5.5. Evaluación del modelo estructural.....	90
5.5.1. Valoración de criterios de calidad.....	90
5.5.2. Interpretación de resultados.....	92
Capítulo 6. Modelado de ecuaciones estructurales basado en la varianza.....	97
6.1. Introducción.....	97
6.2. Fundamentos de Mínimos Cuadrados Parciales (PLS).....	100
6.2.1. Modelación PLS en el Modelo de Medida.....	102
6.2.2. Modelación PLS en el Modelo Estructural.....	105
6.2.3. Estimaciones PLS, PLS _c y PLS _{predictivo}	106

6.2.4. Tamaño de la Muestra.....	106
6.2.5. Software PLS.....	108
6.3. Evaluación de la Modelación con PLS.....	108
6.3.1. Modelos de Medida Estimados en Modo A (Reflectivos).....	109
6.3.2. Modelos de Medida Estimados en Modo B (Medidas Formativas).....	111
6.3.3. Evaluación del Modelo Estructural.....	112
6.3.4. Otros índices para valoración en PLS.....	113
Capítulo 7. Aplicación práctica de ecuaciones estructurales con base en la varianza.....	115
7.1. Introducción.....	115
7.2. Creación de un nomograma en el software SmartPLS.....	115
7.3. Evaluación de los modelos de medida.....	119
7.3.1. Estimación de constructos de tipo reflectivo.....	119
7.3.2. Estimación de constructos de tipo formativo.....	123
7.4. Evaluación del modelo estructural.....	126
7.4.1. Evaluación del modelo global.....	126
7.4.2. Valoración de la colinealidad.....	129
7.4.3. Valoración de las rutas del modelo estructural.....	129
7.4.4. Estimación del coeficiente de determinación (R^2).....	132
Bibliografía.....	135

Prólogo

Las técnicas estadísticas han evolucionado en la investigación científica; en el ayer quedaron las herramientas como análisis de regresión, ANOVA, chi-cuadrada y otras. Hoy las necesidades son distintas, se requiere de análisis multivariante que permita crear una red entramada de variables de las que son codependientes. Surge con ello la segunda generación y la aparición del Modelado de Ecuaciones Estructurales (SEM, por sus siglas en inglés de Structural Equation Modeling), que es una herramienta que identifica más de una relación entre variables y su análisis inferencial se realiza en forma simultánea. AMOS y PLS son dos técnicas progresistas de SEM que han sobresalido en la investigación de alto impacto, ya que sus algoritmos, desde una sola corrida de datos y aunque no exista una conexión directa, obtienen los resultados de las conexiones establecidas por el investigador.

En primera instancia, el SEM por covarianzas ejecutado a través de diversos softwares como AMOS para SPSS es de uso extendido en las ciencias sociales cuando se busca confirmar teorías o modelos ampliamente soportados por la literatura, por lo que tiene un objetivo orientado a su causalidad y contrastación, dando al investigador la oportunidad de corroborar el ajuste del modelo a través de diversos indicadores de bondad de ajuste. Una de las ventajas de los modelos SEM por covarianzas es que permiten el desarrollo de modelos no recursivos o bidireccionales; sin embargo, por su orientación causal, requiere muestras grandes para evitar el sesgo en sus resultados.

En cuanto a los Mínimos Cuadrados Ordinarios (PLS, por sus siglas en inglés de Partial Least Squares), en comunión con las tecnologías de información facilitan los análisis multivariantes; de modo que ni el número de variables ni el número de casos son un obstáculo para aquel investigador que busca evidencia y no solo relacionar variables para encontrar la causalidad. Hoy se busca la predicción de los problemas, eventos o fenómenos que ayude a investigadores e instituciones a contar con más y mejores opciones para su toma de decisiones con base en datos tratados estadísticamente, con una alta validación de sus procedimientos y resultados.

Esta obra se enfoca esencialmente a temas de las ciencias sociales, del comportamiento y administrativas, y tomando en cuenta que los trabajos empíricos llevados a cabo con AMOS y PLS destacan por su alto valor académico, sus aportaciones al conocimiento y un sustento robusto que combina la parte teórica con la inferencia estadística; precisamente por ello su estudio y análisis.

Estas herramientas generan índices básicos, si bien, de igual forma, destacan otros que han evolucionado en su fortaleza para evaluar relaciones entre variables de trabajo, a la vez que se han ido desarrollando nuevos índices que exigen más a sus algoritmos, pero que, sin duda, es para mejorar los resultados. Dentro de

estos índices sobresalen: alfa de Cronbach, AVE, HTMT, f^2 , q^2 , R^2 , coeficiente path estandarizado, t-statistic, CFI, RMSEA, rho_A, entre otros, que son esenciales para la validación de los constructos, la fiabilidad de ítems, la validez convergente, la validez discriminante, la consistencia interna y el poder de decidir, con base en sus resultados, en la aceptación o rechazo de una hipótesis de trabajo.

Sin duda, este documento es una guía para investigadores y practicantes noveles y expertos que les proporcionará la oportunidad de evaluar sus trabajos de investigación e incursionar en las publicaciones de alta calidad, ya que a través de su lectura se conoce la teoría, para aplicarlo en un ejemplo ilustrativo que lleva al lector desde el diseño de un modelo gráfico con AMOS y PLS hasta el desarrollo robusto de un análisis estadístico.

Esperamos que esta obra sea de utilidad y un medio para crecer en el ámbito personal y profesional dentro de la investigación científica, que estamos convencidos que será un punto de partida para muchas cosas mejores por venir.

Los autores

Capítulo 1. Introducción a las técnicas de análisis multivariante

El Modelado de Ecuaciones Estructurales (SEM) es una técnica específica dentro del análisis multivariante de datos con un valor importante para el sector académico, el científico y algunos sectores productivos. La técnica requiere de un conocimiento previo sobre estadística inferencial, por lo que es conveniente repasar sus elementos básicos antes de, más adelante, entrar de lleno con el SEM.

En este primer capítulo se abordan aspectos generales sobre el análisis multivariante, su utilidad y aplicación. También, los datos, estructura, tipos y escalas de medida, ya que su manejo es fundamental para la ejecución de cualquier técnica estadística. Asimismo, se explican brevemente las diferentes técnicas de análisis multivariante, para que el lector tenga una perspectiva amplia de las diferentes opciones de las que puede disponer, así como de su utilidad. Finalmente, se incluye un apartado que expone algunos conceptos complementarios que se deben conocer para facilitar al lector el tránsito por esta obra.

1.1. Análisis multivariante de datos

Se le denomina análisis multivariante de datos al conjunto de técnicas estadísticas empleadas para el estudio de diversos objetos o sujetos con características o atributos identificados o magnificados mediante diversas escalas, y que son analizadas de forma simultánea (Hair et al., 2017). Estos tipos de análisis se utilizan en diversas áreas profesionales como la industria manufacturera, el ámbito empresarial, las ciencias de la salud, las ciencias naturales, las ciencias sociales, por mencionar algunas. La amplia aplicación de estas técnicas estadísticas se debe a que sus resultados permiten la comprensión de fenómenos complejos, plantear soluciones a problemáticas, representar el mundo en una versión simplificada, entre otras (Hair et al., 1999; Veliz, 2016).

Las técnicas multivariantes se han popularizado en las últimas décadas debido a que el análisis simultáneo de variables permite resultados más complejos, con información más profunda, además de que los avances tecnológicos facilitan el tratamiento de grandes cantidades de información mediante diversos *softwares* especializados. Sin embargo, debe quedar claro que la calidad de los resultados y, por lo tanto, su nivel de utilidad está directamente asociada con la correcta aplicación de la técnica estadística (Hair et al., 2017).

En cuanto a la aplicación académica, las técnicas multivariantes de datos son empleadas para desarrollar nuevo conocimiento, así como para poner a prueba los argumentos teóricos establecidos en realidades particulares. Tal es el caso de las ciencias administrativas, donde se emplean las distintas técnicas para diversos fines

(Everitt y Dunn, 1991). Por ello, el estudio, aprendizaje y comprensión por parte del interesado en aplicarlas es fundamental para: primero, seleccionar la técnica multivariante adecuada para cumplir el objetivo trazado; segundo, ejecutar el análisis estadístico pertinente de forma correcta y, tercero, interpretar los resultados obtenidos para los fines deseados.

En razón de lo anterior, el presente capítulo tiene por objetivo explorar algunos de los conceptos básicos para la comprensión de las técnicas de análisis multivariante, así como su clasificación en función de su utilidad. Como punto de partida se abordará la descripción general de los datos, los cuales son la materia prima que emplean las diversas técnicas estadísticas para obtener los resultados.

1.2. Datos

Los datos son la forma más simple de información y representan la medición de una cualidad o magnitud, la cual recibe el nombre de variable, que en esta obra también se puede identificar como indicador o ítem. La denominación se debe a que los datos que la representan lo hacen mediante un valor que difiere de un sujeto a otro. Dicho de otra forma, un dato es la medición que recibe una variable que pertenece a un sujeto de estudio (Everitt y Dunn, 1991; Gujarati y Porter, 2010).

Ahora, a medida que se agregan sujetos de estudio, evaluando la misma variable, se obtiene un conjunto de datos, sobre los que a través de las técnicas estadísticas se puede obtener una mayor información que va más allá de los datos recabados. Cuando se tiene un conjunto de éstos, es conveniente ordenarlos en una Matriz de datos, que ayudará a aplicar las técnicas estadísticas pertinentes.

Tabla 1.1 Ejemplo de matriz de datos univariante

Día de la semana ⁴	Temperatura ²
Lunes ³	35° ¹
Martes	33°
Miércoles	27°
Jueves	32°
Viernes	35°
Sábado	37°
Domingo	35°

Nota: ¹Dato, ²Variable, ³Sujeto de estudio, ⁴Unidad de análisis.

Fuente: Elaboración propia.

La Tabla anterior (1.1) ejemplifica una Matriz de datos compuesta por distintos valores de una variable correspondientes a sujetos de estudio que conforman una unidad de análisis. Además, una Matriz de datos facilita la aplicación de las técnicas de análisis, inclusive de forma manual. Como puede observarse, cada sujeto de estudio (día de la semana) tiene un valor determinado para la variable de interés (temperatura); consecuentemente, cada conjunto de sujetos de estudio ofrecerá la posibilidad de contar con una serie de datos que posibilitará comprender un fenómeno de estudio (Miller, 2001).

La estadística tiene técnicas que permiten conocer el comportamiento de un conjunto de datos, como las medidas de tendencia central (media, mediana y moda) y las medidas de dispersión (varianza y desviación estándar), además de otras técnicas de análisis univariante, y si bien es cierto que su aplicación permite obtener información útil, a medida que se agregan más variables se pueden implementar técnicas con resultados aún más valiosos, como se comentó anteriormente.

A medida que se aumenta el número de datos, el uso de la Matriz de datos adquiere mayor relevancia, debido a que el ordenamiento facilita la identificación y evaluación de series de datos y, de manera general, las técnicas estadísticas representan su comportamiento mediante diversos indicadores. Otro punto a resaltar es que el *software* especializado demanda la conformación de estas matrices para poder desarrollar las técnicas. En la Tabla 1.2 se presenta una forma genérica de conformar una matriz de datos con diversas variables.

Tabla 1.2 Ejemplo de matriz de datos multivariante

	V ₁	V ₂	V ₃	...	V _n
S ₁	D _{1,1}	D _{1,2}	D _{1,3}		D _{1,n}
S ₂	D _{2,1}	D _{2,2}	D _{2,3}		D _{2,n}
S ₃	D _{3,1}	D _{3,2}	D _{3,3}		D _{3,n}
S ₄	D _{4,1}	D _{4,2}	D _{4,3}		D _{4,n}
...					
S _n	D _{n,1}	D _{n,2}	D _{n,3}		D _{n,n}

Nota: D =Dato; V= Variable, S= Sujeto de estudio.

Fuente: Elaboración propia.

En este punto aparece una característica importante que debe ser considerada para la selección adecuada de la técnica de análisis multivariante a utilizar, ese elemento es el tipo de escala de medida de cada variable. Sabemos que un dato es un valor

resultante de la medición de un atributo o magnitud, ahora profundizaremos en lo que representa ese valor, mismo que está en función de la escala de medida (Hair et al., 1999; Hair et al., 2017).

1.3. Escalas de medidas

La escala de medida es el *catálogo de valores* que representa una variable determinada y da el significado al valor o dato asignado. Aquí es importante denotar que la forma de medir es inherente a una variable, por lo que debe existir una coherencia entre ambas; en caso de no presentar coherencia, el valor no representará la información de la característica o magnitud deseada. Por lo tanto, la técnica estadística elegida no podrá ser aplicada o, en algunos casos, resultará en información que llevará a generar conclusiones falsas o a toma de decisiones inadecuada (Hair et al., 2017). Para ello existen diferentes tipos de escala de medida y dependiendo de esa escala será el significado del valor asignado a un sujeto. El tipo de escala estará en función del tipo de variable y lo que pretende representar, ya sea la cuantificación de algún elemento o representar una cualidad (Veliz, 2016).

Comprender el tipo de variable y su escala de medida es importante, puesto que la selección del análisis multivariante a aplicar está en función del objetivo que se busca y el tipo de variable o escala de medida de la cual se dispone. A continuación se describen los tipos de escala de medida y el tipo de variables en las cuales se aplican.

1.3.1. Escalas no métricas

Las escalas **no métricas** son medidas que buscan representar cualidades y se aplican a las denominadas *variables cualitativas*. En estas escalas los números representan, etiquetan o identifican categorías dentro de la unidad de análisis (Hair et al., 1999; Veliz, 2016). La característica que debe cumplir este tipo de escalas es que para una variable se deben de incluir todas las posibles categorías en las que pueda incurrir y que cada una de ellas sea excluyente de las demás opciones; de esta manera cada valor asignado representará una única categoría (Brace, 2018; Hair et al., 2017).

En las escalas no métricas se pueden distinguir dos tipos:

- Escalas nominales. Identifican o etiqueta categorías, es decir, que el valor asignado solo indica la presencia/ausencia de un elemento o pertenencia a una categoría determinada (Hair et al., 2017). Con este tipo de variables se puede contar el número de ocurrencias de cada categoría que compone la escala (Brace, 2018). A manera de ejemplo, se puede mencionar el sexo, tipo de sector donde opera una empresa o identificador (ID) nacional de una persona (credencial de elector en México, DNI en España, por ejemplo).

- Escalas ordinales. Poseen una complejidad mayor, ya que además de identificar categorías, las ordenan o clasifican en función de la calidad o cantidad del atributo representado (Hair et al., 2017; Véliz, 2016). De esta manera, tanto el incremento o decremento de la variable tendrá un significado (Hair et al., 2017). La escala representa una comparación entre las categorías que la integran puesto que es el ordenamiento de una escala nominal bajo un criterio determinado (Brace, 2018), como podrían ser el orden de preferencia de consumo de un tipo de producto o el nivel de satisfacción.

1.3.2. Escalas métricas

Las escalas de medida que cuantifican o magnifican algún elemento son denominadas **métricas**, empleadas en las *variables cuantitativas* (Hair et al., 1999; Veliz, 2016). Esta clase de escalas tiene un tipo de unidad definida para dar valor a la medición en función de un *ranking* preestablecido, lo que permite interpretar las magnitudes con mayor facilidad (Hair et al., 2017).

Algunas escalas métricas son:

- Escala de Intervalo. Se emplea para cuantificar magnitudes en función de una escala determinada. Cada posible valor se encuentra ordenado en la escala y cada uno de ellos es equidistante de otro, por lo que cada medición realizada puede ser clasificada, lo que posibilita comprender la diferencia entre distintos elementos medidos (Hair et al., 2017). Lo característico de esta escala es que el cero no significa la ausencia de la variable medida, sino que tiene un significado dentro de la escala (Véliz, 2016). Por ejemplo, en la medición de la temperatura el cero representa un valor, y aún existen valores por debajo de él.
- Escala de Razón. Es similar a la de intervalo, pero aquí el cero sí representa la ausencia de la variable medida. Además, los valores que la componen no necesariamente son equidistantes entre ellos (Brace, 2018).

1.4. Clasificación de las técnicas de análisis multivariante

La clasificación de las diversas técnicas está en función del objetivo o finalidad que persigue el análisis multivariante. Un primer grupo tiene la función de distinguir, entre el conjunto de variables a analizar, si son dependientes o independientes. En otras palabras, si con variables contempladas en los estudios se puede establecer una relación de dependencia (Hair et al., 1999), la cual se revela cuando el resultado de una o varias variables se da en función de los valores que presentan una o algunas de ellas (Gujarati y Porter, 2010); en este caso, se realizan los denominados análisis de dependencia.

El segundo grupo de análisis multivariante se conoce como análisis de interdependencia, el cual se realiza cuando todas las variables involucradas en el

análisis son tratadas simétricamente, es decir, que no se distinguen categorías, por lo que se deben analizar conjuntamente. Pérez (2004, pp. 2) las describe como *técnicas multivariantes descriptivas*, por los resultados obtenidos del manejo de los datos.

Enseguida se describirán brevemente las técnicas multivariantes más comúnmente empleadas para los análisis de dependencia y para los de interdependencia, haciendo mención en su utilidad, así como en las escalas, en los casos necesarios, para desarrollarlas. Así mismo, es puntual indicar que el objetivo de esta obra no es profundizar en la totalidad de las Técnicas de Análisis Multivariante, sin embargo, si es deseo del lector, se invita a remitirse a obras especializadas, como Análisis Multivariante (Hair et al., 1999), Microeconometrics using STATA (Cameron y Trivedi, 2009) y Técnicas de análisis multivariante de datos Aplicaciones con SPSS (Pérez, 2004).

1.4.1. Análisis de interdependencia

Buscan localizar las relaciones que se presentan entre las distintas variables que conforman la matriz de datos. Son técnicas que muestran la estructura que subyace al grupo de variables en función de las interrelaciones que presentan (Clavijo y Gómez, 1979). En esta categoría existen algunas que son útiles para crear conjuntos en función de un criterio estadístico objetivo, y otras que ayudan a eliminar información redundante (Pérez, 2004). Estas aplicaciones por sí solas no determinan relaciones de dependencia, sin embargo, los resultados obtenidos por ellas pueden ser insumo para técnicas que sí lo hacen, lo que aumenta su valor y utilidad.

La Tabla 1.3 muestra las distintas técnicas que pertenecen a la categoría de análisis de interdependencia. La selección de la apropiada, en primera instancia se da en función al elemento asociativo o disociativo que se empleará como base para realizar la estructura de interrelaciones. En segundo término se considera el tipo de escalas que miden a los elementos del estudio.

Tabla 1.3 Clasificación de técnicas de interdependencia

Elementos Asociativos	Tipo de escala empleado	Técnica Multivariante
Variable	Métrica	Análisis Factorial/Componentes
	No Métrica	Principales
Casos/Encuestados	Métricas	Análisis Clúster
	No Métricas	
Objetos	Métrica	Análisis Multidimensional
	No Métrica	Análisis Multidimensional
		Análisis de Correspondencia

Fuente: Elaboración propia a partir de Hair et al. (1999).

Análisis Factorial/Componentes Principales. El análisis factorial y el análisis de componentes principales son similares, pero cuentan con diferencias que no se tratarán en esta obra, por lo que es importante aclarar al lector que los elementos que las distinguen deben ser analizadas por el interesado, para optar por la técnica adecuada.

Estas técnicas se emplean cuando se cuenta con una gran cantidad de variables (ítems/indicadores). Su propósito es simplificar la información contenida en un conjunto de variables mediante la extracción de la información más relevante, para crear nuevos factores o componentes que faciliten la interpretación del estudio. Los resultados pueden fungir como insumos para aplicar otros análisis multivariantes (Hair et al., 1999, Véliz, 2019). El objetivo es reducir el número de datos mediante la eliminación de información redundante, así como la selección de elementos que proveen una mayor cantidad de información, generando *Factores* o *Componentes* con valores calculados mediante aquellos insumos de mayor relevancia (Pérez, 2004).

Análisis Clúster. También conocido como análisis por conglomerados. Su objetivo es la creación de grupos de sujetos con características similares pero lo suficientemente discordantes como para ser excluyentes de los demás grupos conformados. Se aplica cuando no existe una preconcepción de los grupos existentes, sino que el resultado es el que dictará la descripción de los colectivos en función de las características que generan la convergencia de los sujetos y los distinguen de los demás grupos (Hair et al., 1999).

Análisis Multidimensional. Se emplea para la clasificación de objetos en función de n dimensiones o variables. De acuerdo con los valores que obtengan las dimensiones o variables, los objetos serán asociados entre sí en función de las concordancias o discordancias de las múltiples dimensiones consideradas.

Análisis de Correspondencia. Similar al análisis factorial o de componentes principales, pero el análisis por correspondencia emplea variables cualitativas, por lo que su objetivo es reducir la cantidad de datos, manteniendo la información más relevante para encontrar una tendencia general. Para ello se emplea una tabla de contingencia, para representar y cuantificar la ocurrencia de las variables estudiadas en función de diversas características del sujeto de estudio (Hair et al., 1999; Pérez, 2004).

1.4.2. Análisis de dependencia

El objetivo básico de este tipo de técnicas estadísticas es conocer o predecir el valor de una variable (variable dependiente) en función de una o más variables conocidas (variable independiente). El elemento principal para emplear estas técnicas es la distinción de la *jerarquía* o clasificación de las variables en dependientes e independientes. Pérez (2004, pp. 2) las denomina técnicas multivariantes analíticas o inferenciales, puesto que el resultado obtenido permite realizar inferencias estadísticas.

La Tabla 1.4 muestra las distintas técnicas que pertenecen a la categoría de análisis de dependencia. La selección de alguna de ellas está primeramente en función del número de variables dependientes involucradas en el análisis; después, al número de relaciones establecidas y finalmente, de acuerdo con las escalas de medida empleadas para la variable dependiente.

Tabla 1.4 Clasificación de las técnicas de dependencia

Número de variables dependientes	Número de relaciones	Escala de medida de Variable dependiente	Técnica de Análisis Multivariante
Única	Única	Métrica	Regresión Lineal Simple
			Análisis de conjunto
		No métrica	Análisis discriminante
Varias	Única		Regresión Logística
		Métrica	Análisis de Correlación canónica ¹
			Análisis Multivariante de la Varianza ²
Varias	Múltiples	No Métrica	Análisis de correlación canónica con variables ficticias
		Métricas	Modelo de Ecuaciones Estructurales
		No métricas	

Nota: ¹Para ejecutar la técnica de Análisis de correlación canónica, además de las condiciones mostradas en la tabla se deberá contar con variables independientes métricas. ²Para ejecutar la técnica de análisis Multivariante de la varianza se deberá contar con variables independientes no métricas.

Fuente: Elaboración propia.

Regresión Lineal. La regresión lineal simple (una variable independiente) y la múltiple (dos o más variables independientes) son el análisis de dependencia más comúnmente usado. Establecen la relación entre una variable dependiente y la o las variables independientes. Buscan estimar o predecir el valor medio de la variable dependiente en función de los valores conocidos de la(s) variable(s) independiente(s) (Gujarati y Porter, 2010). Para ejecutar esta técnica, tanto la variable dependiente como la(s) variable(s) independiente(s) deben ser representadas por escalas métricas (Pérez, 2004).

Análisis Conjunto. Parte de la definición de características que podrán ser reales o hipotéticas, mismas que serán evaluadas por los sujetos de estudio y definirán la preferencia en función de la valoración subjetiva del conjunto de atributos. La técnica es útil para evaluar objetos con una serie de propiedades combinadas y definir las preferencias de las combinaciones presentadas o, en su defecto, la combinación de características más aceptada (Hair et al., 1999).

Análisis Discriminante. Se usa cuando la variable dependiente es no métrica y cuenta con más de dos nomenclaturas o categorías, mientras que las variables independientes son métricas. El objetivo del análisis es comparar los grupos que forman las diversas categorías en función de los valores generalizados de las variables de escala métrica. Se emplea para conocer los factores que llevan a los sujetos a pertenecer a determinada categoría o etiqueta (Hair et al., 1999), en otras palabras, mediante las variables independientes se busca predecir a qué categoría pertenece un caso determinado (Pérez, 2004).

Regresión Logística. Esta técnica genera un modelo de probabilidad de eventos y se emplea tradicionalmente cuando la variable dependiente es binaria, es decir, cuando solo cuenta con dos categorías o etiquetas. También son conocidos como modelos logísticos, *logit* (Hair et al., 1999), o modelos de probabilidad lineal (Pérez, 2004). La regresión *logit* busca medir la probabilidad y los cambios en ella de incurrir en una de las dos categorías; la probabilidad de ocurrencia estimada está en función de variables independientes incluidas en el modelo multivariante (Cameron y Trivedi, 2009).

Análisis de correlación/regresión canónica. Se emplean en los casos en que se pretenda analizar más de una variable dependiente respecto a una o más variables independientes. Otra aplicación es cuando tanto las variables dependientes como las independientes no son métricas. El objetivo de estas técnicas es similar al de la regresión lineal y de correlación lineal (Pérez, 2004).

Análisis multivariante de la varianza. Su objetivo es medir, comparar y establecer semejanzas o diferencias existentes entre dos o más variables métricas dependientes divididas por una variable categórica. El análisis multivariante de la varianza es conocido como MANCOVA y es un paso más avanzado que el ANOVA (análisis de la varianza), en donde solo se considera una variable para diferentes conjuntos.

Modelo de ecuaciones estructurales (SEM, por sus siglas en inglés de *Structural Equation Modeling*). Es una técnica multivariante que se diferencia de las mencionadas anteriormente en que permite establecer más de una relación de forma simultánea mediante inferencias estadísticas. Hay otras técnicas que también pueden hacerlo, pero de forma individual, no simultánea (Byrne, 2016).

El SEM se describe más ampliamente en el resto de los capítulos de la presente obra, por lo que el breve repaso sobre las técnicas de análisis multivariantes pretende poner en contexto al lector sobre algunos elementos básicos para, de a poco, introducirlo en la complejidad tanto teórica como práctica que implican los Modelos SEM.

1.5. Conceptos complementarios

El empleo de las técnicas multivariantes implica el uso de diversos términos que forman parte del lenguaje manejado entre los profesionales que las aplican. En la

mayoría de los textos especializados se asume el conocimiento de algunos conceptos básicos, ya que solo se trata de una explicación breve de ellos. A continuación se describen algunos de esos elementos con los que el lector debe familiarizarse para una mejor comprensión de las diversas técnicas de análisis de datos.

Análisis Exploratorio. Se refiere a la aplicación de diferentes técnicas estadísticas descriptivas para analizar la composición del conjunto de datos, para conocer su estructura.

Consistencia. Grado en el que los resultados obtenidos se acercan al valor real (Gujarati y Porter, 2010).

Residuos. Diferencia entre los valores observados y los valores estimados (Gujarati y Porter, 2010).

Error estadístico. También conocido como perturbación. Es una variable aleatoria que no es incluida en el modelo de forma consciente, sin embargo, está presente; incluye elementos que afectan el fenómeno de estudio (no contemplados) o modifican el resultado que se obtiene del análisis (Gujarati y Porter, 2010).

Inferencia estadística. Es el proceso de comprobar o refutar hipótesis con base en los resultados obtenidos por las técnicas de análisis multivariantes (Gujarati y Porter, 2010) y lleva a una conclusión entre lo que se plantea y la evidencia observada (Everitt y Dunn, 1991).

Normalidad. Se refiere al cumplimiento de las características de una distribución normal (también conocida como distribución *gaussiana*), una distribución de frecuencias que presenta una serie de datos (Pérez, 2004).

Significancia estadística. Resulta del procedimiento para probar la existencia de una aseveración o una propiedad (Triola, 2004).

Validez. Comprobar o medir el grado de fiabilidad de un modelo estadístico (Pérez, 2004).

Varianza. Medida de dispersión que representa el grado de separación de un grupo de datos entre sí (Triola, 2004). Gran parte de las técnicas multivariantes buscan comprender el comportamiento de una variable, para lo cual se toma a la varianza como la representación de ese comportamiento que se pretende explicar.

Después de este breve repaso de lo más importante de las Técnicas de Análisis Multivariantes, el lector se encuentra preparado para adentrarse en el estudio del SEM; para ello, se inicia con un preámbulo a la Modelización de Ecuaciones Estructurales, para presentar sus elementos fundamentales, y se continúa con el capítulo dedicado a exponer los supuestos básicos que se deben cumplir para realizar el análisis.

El siguiente par de capítulos aborda el SEM basado en la covarianza; el primero, desde un punto de vista teórico y el segundo, desde su aplicación práctica.

Finalmente, los últimos dos capítulos analizan el SEM basado en la varianza, desde la perspectiva teórica el VI, y el último, desde su aplicación práctica.

Capítulo 2. Análisis multivariante mediante Modelado de Ecuaciones Estructurales

2.1. Introducción

Las herramientas estadísticas han sido usadas durante mucho tiempo por profesionistas y por académicos para el desarrollo de diversos trabajos. En la década de 1980 las técnicas estadísticas que dominaron fueron llamadas de primera generación. Para principios de 1990 la evolución tecnológica permitió el desarrollo e implementación de técnicas de análisis multivariantes más avanzadas (Hair et al., 2016), a las que se les denominó de segunda generación, y su uso proliferó por la relativa facilidad de desarrollarlas en los diversos *softwares* especializados. Dentro de las técnicas de segunda generación se encuentra el Modelado de Ecuaciones Estructurales (SEM). Hair et al. (1999, pp. 14) la catalogan como una técnica emergente empleada para lograr una explicación eficiente de un determinado fenómeno.

Como se mencionó en el capítulo anterior, la principal ventaja que presenta SEM respecto a otras técnicas es la posibilidad de evaluar de manera simultánea diversas relaciones de dependencia, además de poder emplear tanto variables observables como variables teóricas (también conocidas como variables latentes), que se componen de diversas variables observables (Byrne, 2016; Escobedo et al., 2016; Hair et al., 2016). La técnica multivariante SEM engloba distintos tipos de modelos estadísticos cuya utilidad y características los hacen peculiares y que más adelante se explicarán. Hair et al. (1999, pp. 612) mencionan que cualquier técnica SEM debe cumplir con dos condiciones: *i*) estimar múltiples relaciones simultáneamente y *ii*) tener la capacidad de representar variables no observables.

SEM se usa mayormente para comprobar la existencia de relaciones entre variables y evaluarlas mediante inferencias estadísticas, y resulta sumamente útil para comprobar hipótesis, por lo que es ampliamente recurrida en investigaciones en ciencias sociales, aunque también tiene aplicaciones en otros campos (Byrne, 2016). El objetivo de esta obra es tratar los SEM desde la perspectiva académica de las ciencias administrativas.

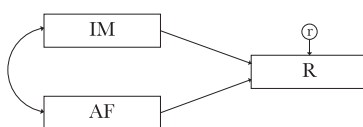
2.2. ¿Qué es la modelización de ecuaciones estructurales?

Las técnicas SEM, a pesar de presentar relaciones complejas, poseen la ventaja de que estas relaciones son expuestas de forma gráfica y no solo numérica, por lo que la comprensión del fenómeno puede llegar a ser más simple, ya que muestra las variables que intervienen en el estudio (y en muchos casos, la conformación de las variables), las relaciones existentes y los valores estadísticos resultantes. A esta representación gráfica y sistemática se le conoce como *path diagram* en inglés y nomograma en español (Byrne, 2016).

La modelización de las ecuaciones estructurales sigue una nomenclatura estandarizada y sistemática que permite la ejecución de las técnicas con los *softwares* estadísticos y la lectura de los resultados por los interesados. La modelización corresponde a la etapa de definición de las relaciones hipotéticas entre variables que se pretende probar (Byrne, 2016; Escobedo et al., 2016). Esta construcción hipotética de relaciones puede realizarse a partir de teorías establecidas, la experiencia del investigador o simplemente por las relaciones que se pretende evaluar o comprobar (Hair et al., 2018). Como ejemplo que ayuda a introducir al modelado dentro de SEM, en la Figura 2.1 se presentan dos relaciones sencillas entre distintas variables suponiendo que se busque la relación que existe entre la inversión en mercadotecnia (IM) y la rentabilidad (R), y entre el apalancamiento financiero (AF) y la rentabilidad (R).

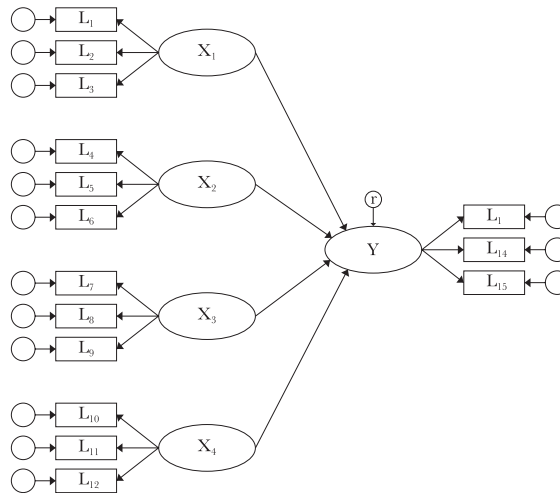
Como puede apreciarse, la modelización de variables y relaciones emplea una tipografía determinada que auxilia a la sistematización de la técnica, misma que resulta lógica, por lo que su comprensión es relativamente fácil y simple su lectura. En la Figura 2.1 se demuestra la condición de analizar múltiples relaciones de forma simultánea, lo que distingue a SEM de las otras técnicas de análisis multivariante; aunque esta es una representación básica que no contempla una de las ventajas del SEM, que es la posibilidad de incluir variables no observables (que se describirán más adelante). Esta inclusión sí se visualiza en el ejemplo que se muestra en la Figura 2.2.

Figura 2.1. Modelización de la relación entre variables observables.



Fuente: Elaboración propia.

Figura 2.2. Modelización con variables no observables.



Fuente: Elaboración propia.

En la Figura 2.2. se distinguen cuatro variables independientes (X_1 , X_2 , X_3 y X_4), que son variables latentes y se miden por tres diferentes variables observables, que reciben el nombre de indicadores o ítems y son representadas por un rectángulo. Cada una de ellas se relaciona con la variable dependiente Y , que a su vez es una variable no observable conformada por tres indicadores (Véliz, 2016). En la Tabla 2.1. se describen los elementos que componen el modelo anterior, para una mayor comprensión de los diagramas y un panorama más amplio sobre la estructura de los *path diagram*.

Tabla 2.1 Nomenclatura básica del modelo SEM

Elemento	Nombre	Significado
	Variable Observable	Representa una variable que es medible por un único indicador, conocido también como ítem.
	Flecha unidireccional	Representa las relaciones establecidas. Cada flecha unidireccional significa una relación, donde al origen la flecha se le denomina independiente y a donde llega se le llama variable dependiente.

Elemento	Nombre	Significado
	Flecha bidireccional	Indica que las variables a las que apunta se correlacionan o covarían.
	Variable no observable/variable latente/constructo	Muestra una variable no observable que es construida a partir de otras variables, observables o no (indicadores/ítems).
	Error de medida y residuo	Término error de medida de la variable observable (indicador/ítem). Término residual de la variable dependiente.

Fuente: Elaborado a partir de Byrne (2016) y Véliz (2016).

2.3. Tipo de modelización

Las modelizaciones SEM se dividen en dos tipos, en función de la construcción del modelo. El primer tipo consta de los modelos que nacen de las teorías definidas dentro de la literatura y buscan ser comprobadas o rechazadas, de acuerdo con los resultados del análisis. En otras palabras, son modelos cuya estructura es conocida y se pretende verificar que la realidad estudiada se estructura de la misma forma que el postulado teórico. Se les llama Modelo de Ecuaciones Estructurales Basados en la Covarianza (*Covariance-Based Structural Equation Modeling*, CB-SEM) y, de acuerdo con Hair et al. (2017, p. 4), “determinan qué tan bien una propuesta teórica explica la covarianza de un conjunto de datos”, y son usados ampliamente para confirmar teorías o hipótesis de investigación (Hair et al., 1999).

El CB-SEM, como su nombre lo indica, para medir la variable no observada considera la varianza común entre los indicadores, por lo que esa varianza compartida entre los indicadores generará el valor de la variable latente (Hair et al., 2018). En los casos en que se desarrolle la técnica CB-SEM, el investigador deberá estar pendiente de la concordancia entre la modelización teórica usada y los resultados obtenidos, verificar distintos indicadores estadísticos que tienen como finalidad medir la idoneidad de los resultados, y la adecuación de la información empleada, entre otras (Escobedo et al., 2016; Hair et al., 1999).

Evidentemente, los modelos de ecuaciones estructurales basados en covarianzas siguen un método minucioso tanto previo como posterior al análisis. Inclusive, no todos los *softwares* estadísticos especializados en análisis de datos pueden ejecutar la técnica, y los que lo hacen requieren de pasos específicos, lo cual denota la dificultad para desarrollarla. Los pormenores conceptuales y teóricos del modelo basado en covarianzas serán profundizados en el Capítulo 4; además, en el Capítulo 5 se presentará un ejemplo práctico que desarrolla esta técnica, que pueda servir como guía práctica para su ejecución.

El segundo tipo de las modelizaciones SEM es la Modelación de Ecuaciones Estructurales Basada en la Varianza, conocido como Mínimos Cuadrados Parciales (*Partial Least Squares*, PLS-SEM), que se emplea en investigaciones de carácter exploratorio, es decir, cuando no se tiene una teoría definida como base para desarrollar el modelo, sino que se emplean las propias variables de estudio para encontrar una estructura que represente la realidad estudiada. Hair et al. (2017, pp. 4) plantean que el PLS-SEM “se enfoca en explicar la varianza de la variable dependiente al examinar el modelo”. Dicho de otro modo, para medir la variable no observable el PLS-SEM emplea la varianza de cada variable manifiesta o indicador (ítem) que conforma la variable latente, dándole valor mediante la combinación lineal de las varianzas de cada ítem considerado (Hair et al., 2018). A manera de información adicional, en el último lustro los avances en la técnica indican que tiende hacia la predicción más que a la causalidad entre las variables.

La aplicación del PLS-SEM ha aumentado en los últimos años, en parte por el desarrollo de *software* que facilita su uso, pero sobre todo por su flexibilidad para analizar los datos (Hair et al., 2018). Esta flexibilidad radica en que, a diferencia el CB-SEM, el PLS-SEM no debe cumplir con los supuestos o asunciones sobre las cuales trabaja el modelo de covarianzas (Hair et al., 2018).

El PLS-SEM es ideal para desarrollar o reformular modelos, también es útil para comparar representaciones de realidades diferentes; no obstante, esto debe realizarse con respaldo en las evidencias teóricas existentes y no limitarse a la evidencia empírica resultante, lo que significa que los resultados deben tener un argumento lógico más allá de que presenten los mejores indicadores estadísticos (Hair et al., 1999). De igual manera, es una técnica con un incremento reciente de su aplicación en diversos ámbitos (Hair et al., 2017) y, al igual que el CB-SEM, se deben atender ciertas consideraciones tanto teóricas como prácticas para implementarla y ejecutarla de forma adecuada. En el Capítulo 6 se profundiza en el PLS-SEM y en el Capítulo 7 se presenta un ejemplo de su aplicación, a manera de guía para ejecutarla.

2.4. Tipos de variables

En apartados anteriores se ha mencionado que para los SEM existen dos tipos de variables, una de las cuales no se presenta en las demás técnicas de análisis multivariantes. SEM permite emplear variables que no son observables directamente del sujeto u objeto de estudio, sino que son constructos teóricos que son medidos por la conjunción de otras variables, que reciben el nombre de indicadores (ítems) (Véliz, 2017).

Se ha explicado la forma de representar este tipo de variables en el Modelo gráfico (*path diagram/nomograma*), sin embargo, aún se debe profundizar en algunos aspectos de este tipo de variables que son conocidas como variables latentes y son construidas teóricamente, de tal forma que pueden definirse por variables manifiestas cuyos valores pueden indicar el valor de la variable latente. Es decir, las variables latentes no se observan, pero se infiere su existencia mediante la presencia de otras variables que en conjunto las representan (Byrne, 2016).

En resumen: una variable latente es un elemento de gran complejidad que se desagrega o divide en elementos de mayor simplicidad cuya medición es menos problemática, de modo que el valor de la variable no observable puede obtenerse mediante la suma de las variables manifiestas que la conforman a (Byrne, 2016). Así, las mediciones de los fragmentos en los que se dividió la variable latente unirán su información para crear una medida de ésta. La Tabla 2.2 ejemplifica una matriz de datos con variables latentes. Como se puede apreciar, se muestran dos variables que son medidas mediante tres indicadores cada una, por tanto, el valor de cada variable estará compuesto por los valores codificados para cada indicador.

Tabla 2.2 Ejemplo de Matriz de datos con variables no observables

	V1			V2		
	I ₁	I ₂	I ₃	I ₄	I ₅	I ₆
S ₁	D _{1,1}	D _{1,2}	D _{1,3}	D _{1,4}	D _{1,5}	D _{1,6}
S ₂	D _{2,1}	D _{2,2}	D _{2,3}	D _{2,4}	D _{2,5}	D _{2,6}
S ₃	D _{3,1}	D _{3,2}	D _{3,3}	D _{3,4}	D _{3,5}	D _{3,6}
S ₄	D _{4,1}	D _{4,2}	D _{4,3}	D _{4,4}	D _{4,5}	D _{4,6}
S _n	D _{n,1}	D _{n,2}	D _{n,3}	D _{n,4}	D _{n,5}	D _{n,6}

Fuente: Elaboración propia.

Existen técnicas estadísticas específicas para determinar los valores de las variables latentes, como algunas técnicas de interdependencia mencionados en el Capítulo 1, no obstante, en capítulos posteriores, dedicados a profundizar en los modelos SEM, se retomará la medición de estas variables.

Por otro lado, las variables observables son más simples de entender debido a que son indicadores únicos cuya información completa se encuentra en ese valor. Son elementos que se manifiestan de forma evidente en el fenómeno de estudio (Véliz, 2017). Una manera de sintetizar este apartado es que las variables observables son

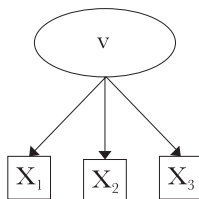
aquellas que pueden medirse directamente, mientras que las latentes se cuantifican de forma indirecta (Byrne, 2016). Es importante hacer esta distinción debido a que en las ciencias sociales, en las administrativas o en las del comportamiento los fenómenos estudiados requieren una medición indirecta, ya que se emplean construcciones teóricas complejas difíciles de medir, o se analizan fenómenos multidimensionales que implican el uso de diversos indicadores complementarios Byrne, (2016).

2.4.1 Construcción de Modelos

En SEM existen dos tipos de modelos, en función de la forma en la construcción de las variables no observables, definidos como modelos reflectivos y modelos formativos. De manera general, el modo en cómo los indicadores conforman la variable no observable distinguen a ambos modelos (Hair et al., 2018). La explicación teórica de cada uno de ellos, así como sus diferencias, requieren de una revisión profunda. En esta obra se examinarán brevemente, con el objetivo de conocer ambas vertientes y sus generalidades. Para una exploración más amplia se recomienda el libro “Advanced issues in partial least squares structural equation modeling” de Hair et al., (2018).

Modelos Reflectivos. Esta perspectiva es la más común en SEM. Las variables latentes están conformadas por indicadores que están correlacionados, por lo tanto, cuentan con una varianza común, misma que da origen a la variable no observable. Son desarrollados en los modelos CB-SEM y en PLS-SEM. La forma de identificar a las variables no observables reflectivas es en las flechas unidireccionales que van de la variable no observable (elipse) a los indicadores (cuadros o rectángulos), como se ilustra en la Figura 2.3. En este caso, la varianza compartida entre los indicadores (X_1, X_2, X_3) refleja el valor de la variable latente (V).

Figura 2.3. Variable latente de un modelo reflectivo.

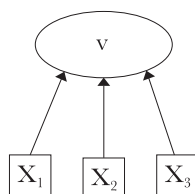


Fuente: Elaboración propia.

Modelos Formativos. También son desarrollados en los PLS-SEM. Se distinguen de los modelos reflectivos en que los indicadores que conforman la variable latente no

necesariamente están relacionados, pues pueden ser manifestaciones distintas que dan como resultado una variable.

Figura 2.4. Variable latente de un modelo formativo.



Fuente: Elaboración propia.

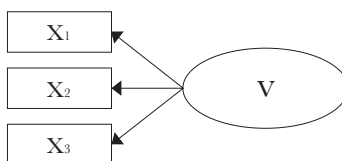
Como muestra la Figura 2.4, las variables latentes de los modelos formativos se caracterizan gráficamente por que las flechas unidireccionales van de los indicadores (cuadros o rectángulos) hacia la variable no observable (elipses), lo cual significa que los indicadores contribuyen a causar la variable latente (Hair et al., 2018).

Cabe recalcar que la definición o elección del tipo de modelo SEM a desarrollar, ya sea formativo o reflectivo, no sigue una regla o conjunto de reglas, sino que debe resultar de la estructura lógica y análisis del desarrollo del fenómeno estudiado (Hair et al., 2018).

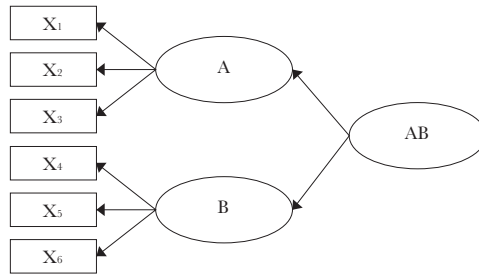
2.4.2. Jerarquía de los componentes

Las relaciones entre distintos conceptos teóricos pueden ser estudiadas empíricamente a través de un modelo de ecuaciones estructurales. Estos conceptos teóricos son medidos a través de variables latentes que se construyen a través de indicadores o ítems. No obstante, en algunas ocasiones la teoría o concepto analizado es complejo o multidimensional, por lo que su medición requiere la conjunción de distintos constructos. Cuando se tiene una variable no observable que se mide por una serie de indicadores, se le conoce como variable de primer orden; por otro lado, cuando un constructo se conforma por varias variables latentes, se le denomina variable de segundo orden (Hair et al., 2018), la Figura 2.5 lo representa gráficamente.

Figura 2.5. Ejemplo de la jerarquía de variables no observables



Ejemplo de Variable de primer orden



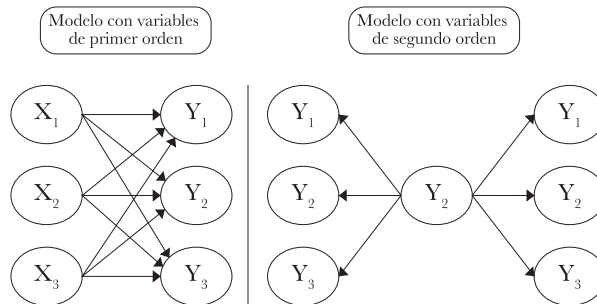
Ejemplo de Variable de segundo orden

Fuente: Elaboración propia.

Una vez que los conceptos teóricos son construidos a través de variables latentes de primero o segundo orden, se introducen en el modelo estructural, para trazar la relación que guardan con otros conceptos teóricos, que a su vez pueden ser de primero o de segundo orden.

La Figura 2.6 ilustra las relaciones analizadas por modelos de primer orden respecto a los modelos de segundo orden. El primer gráfico muestra las relaciones entre variables latentes independientes de primer orden y variables latentes dependientes de primer orden. El segundo, en cambio, muestra la relación de una variable latente independiente de segundo orden y tres variables dependientes latentes de primer orden.

Figura 2.6. Comparativo entre SEM con variables de primero y segundo orden



Fuente: Basado en Hair et al. (2018)

Como se puede observar, el primer gráfico analiza nueve relaciones directas entre las variables independientes y dependientes, mientras que el segundo solo estima tres relaciones directas, puesto que la variable latente independiente es de segundo orden.

Ambos tipos de modelos presentados en la figura anterior tienen sus características particulares y elementos que deben considerarse, e

independientemente de la complejidad del fenómeno, SEM es una herramienta que puede ser empleada en estudios con diferentes niveles de análisis. Y si bien este tema es amplio, el objetivo de la presente obra solo es dar una base al lector.

2.5. Mediciones

Medir es esencial para realizar los análisis estadísticos y, como se explicó en el capítulo anterior, la forma en la que se miden las variables es fundamental para desarrollar determinada técnica estadística. En algunos casos, medir es un procedimiento relativamente sencillo, sin embargo, para las ciencias administrativas, la medición de la mayoría de los fenómenos de estudio presenta importantes retos al investigador (Brace, 2018).

Medir consiste en asignar un valor a la variable, mismo que tiene un significado y representa información. La medición es fácil cuando se trata de dar valor a una variable simple y cuando dicho valor puede ser interpretado claramente; también lo es cuando existe una escala o unidad predeterminada para realizar mediciones. Por ejemplo, si se trata de medir las ventas de una empresa, se puede obtener por el número de unidades vendidas o en función de una unidad monetaria; en ambos casos la lectura del valor es clara y fácilmente interpretable. Sin embargo, existen otras variables cuya complejidad o abstracción impiden que su valoración sea por un único índice, debido a que su manifestación es multidimensional. Muchos de los objetos de estudio dentro de las ciencias administrativas se encuentran dentro de estos fenómenos complejos, ya que se refieren al comportamiento o actitudes de individuos (Brace, 2018).

El enfoque más empleado para realizar la medición es dividir el fenómeno complejo, es decir, tomar esa construcción teórica y seccionarla en elementos más simples. De esta manera se contará con una serie de indicadores, también llamados ítems, que tendrán una valoración individual y en conjunto representarán la variable no observable. Indudablemente, a medida que se incrementa el número de ítems para valorar una variable latente será más complejo conseguirlo, no obstante, múltiples ítems pueden volver la medición más exacta al evaluar distintos elementos del constructo complejo.

Uno de los objetivos es realizar la medición lo más exacta posible y reducir los errores que puedan conducir a obtener un valor que no represente la variable que se pretende medir (Hair et al., 2017). Si bien anteriormente se han mencionado elementos que puedan causar errores en la medida, como el uso inadecuado de una técnica o una mala elección en la escala para valorar una variable, es conveniente retomar algunos de ellos ahora que la discusión se centra en los modelos de ecuaciones estructurales.

2.6. Escala de medida

Como se describió en el capítulo anterior, las escalas de medida son los catálogos de valores asignados a una variable. La adecuada elección de la escala es fundamental para la ejecución de la técnica, y la mayoría de las técnicas de análisis multivariante, entre ellas SEM, precisan de escalas de intervalo para ser ejecutadas. El hecho de que SEM requiera de una escala determinada abre grandes retos para las ciencias administrativas, debido a que pocas de sus variables cuentan con escalas que las representen adecuadamente. Por el contrario, la mayoría de los fenómenos exigen implementar escalas especiales, y en ocasiones es necesario desarrollarlas.

Como se mencionó, la mayoría de las técnicas multivariantes emplean determinada escala, por lo que es comprensible que los instrumentos de medición deben estar en función de escalas que cumplan con las características y propiedades de las escalas de intervalo (Joshi et al., 2015). El procedimiento para lograrlo es la codificación de los datos, que se describirá más adelante.

Cuando se trabaja con variables latentes, actitudes o comportamiento, como en las ciencias administrativas, es común desarrollar escalas con ítems de ranqueo, que emplean la escala de intervalo para valorar una serie de indicadores asociados a variables latentes (Brace, 2018). Este mismo investigador señala que las escalas con ítems de ranqueo empleadas en cuestionarios o encuestas, mismas que fungan como instrumento de medida para la serie de variables de interés, ofrecen la posibilidad de estandarizar la medición de diversos conceptos y facilitan la colecta de la información.

Cabe indicar que el desarrollo de escalas de medida en las ciencias administrativas es un campo amplio y especializado, no obstante, en la presente obra se exponen los elementos básicos para que el lector conozca lo esencial para el desarrollo de SEM.

2.7. Tipos de codificación de datos y construcción de la escala

La codificación es el proceso de asignar números o valores a categorías a fin de facilitar la medida de una variable. Codificar es útil para poder hacer conversiones de escalas de medida, proporcionando un mayor peso a la información que representa. Por ejemplo, una escala ordinal puede ser recodificada asignando valores numéricos, transformándola a una escala métrica (Hair et al., 2017).

Estas acciones, además de convertir a escalas más útiles para los análisis multivariantes, pueden estandarizar la forma de medir distintas variables (Joshi et al., 2015), incluyendo las no observables, lo que permite a la vez un tratamiento estadístico uniforme, es decir, usar una escala idéntica para medir variables distintas. Este procedimiento es fundamental en las ciencias administrativas debido a que muchas variables no son observables y además son factores medidos a través de las

percepciones, actitudes o sentimiento de las personas sobre un fenómeno, lo que aumenta la complejidad de la medición (Albaum, 1997; Chomeya, 2010; Joshi et al., 2015), por ejemplo:

- Calidad de un producto
- Satisfacción laboral
- Nivel de motivación de los empleados
- Utilidad de servicios tecnológicos

Lo más seguro es que distintos sujetos de una unidad de análisis midan de forma diferente un mismo fenómeno. Esto se debe a distintas razones, como una conceptualización distinta (sobre todo con construcciones teóricas complejas), no contar con una escala definida para dar la valoración, ser diferente el significado de un valor entre los observadores, y otras más (Albaum, 1997). Estos inconvenientes disminuyen con el uso de escalas que permitan homogeneizar la valoración que realizan los sujetos de estudio, además de codificar las magnitudes de distintas variables en un mismo *idioma*. La escala más utilizada para estos fines se conoce como escala tipo Likert (Albaum, 1997; Joshi et al., 2015; Likert, 1932).

La escala tipo Likert consiste en una sucesión de planteamientos con una serie de opciones limitadas, para evaluar la percepción o sentimiento de un individuo de forma estandarizada (Likert, 1932); así mismo, facilita la conformación de cuestionarios, debido a que pueden incluirse varios ítems o indicadores (conocidos como batería de ítems) de fácil aplicación para medición (Brace, 2018). La escala es presentada en forma ordinal, pero se codifica para obtener una escala de intervalo, siempre que se cumplan los principios de equidistancia y simetría.

Para observar el principio de equidistancia se deben establecer categorías claramente distinguibles entre sí y ordenadas de una forma ascendente o descendente en una dirección (por ejemplo, grado de concordancia o discordancia, idoneidad o nivel de fuerza) (Albaum, 1997), de forma que el sujeto que realizará la valoración pueda ver la escala como una escalera recta, donde cada escalón sea fácilmente identificable de los demás. Brace (2018, pp. 75) considera a este ordenamiento como una de las características necesarias para desarrollar una escala Likert, sobre todo para facilitar tanto la lectura como la emisión de una valoración a un ítem determinado por parte del sujeto objeto del estudio.

En la Figura 2.7 se incluyen algunos de los ejemplos más comunes empleados como escala tipo Likert y, como puede apreciarse, cada una de las categorías se distingue de las demás, pero en conjunto forman una serie ordenada de categorías a las que se les puede recodificar con valores, lo que ayuda a conseguir la equidistancia entre las distintas categorías debido a que cambiar de una categoría a otra contigua

representará un cambio en la misma cuantía que si se replicara el mismo cambio en cualquier otro punto de la escala (Hair et al., 2017).

Figura 2.7. Ejemplos de escalas Likert.

Completamente en desacuerdo	Algo en desacuerdo	Ni en desacuerdo ni de acuerdo	Algo de acuerdo	Completamente de acuerdo
Completamente débil	Algo débil	Ni débil ni fuerte	Algo fuerte	Completamente fuerte
Inefectivo	Poco inefectivo	Ni efectivo ni inefectivo	Poco efectivo	Efectivo
Muy malo	Malo	Regular	Bueno	Muy Bueno
Codificación				
1	2	3	4	5

Fuente Elaboración propia.

Para que la escala cumpla el principio de simetría se recomienda identificar un punto medio en la escala Likert y verificar que esté equilibrada hacia los extremos. Dicho de otra forma, se puede partir de un punto medio y comprobar que en ambas direcciones de la escala se encuentre el mismo número de categorías (Joshi et al., 2015). Otra forma es equilibrar el número de opciones positivas respecto de las negativas; esto es útil en los casos donde no exista un valor central; no obstante, se recomienda que exista, para contar con una opción en caso de una valoración neutra (Brace, 2018).

No encontrar una escala simétrica podría sesgar hacia cierto punto la valoración de la variable medida, generando resultados poco confiables (Brace, 2018; Joshi et al., 2015). Esto se debe a que puede inducirse a una determinada respuesta, consciente o inconscientemente, en función del número de opciones positivas o negativas incluidas en la escala. Un punto controversial respecto a la escala tipo Likert es el número de categorías o escalones con los que debe contar (Chomeya, 2010). Se considera que escalas entre cinco a diez puntos son adecuadas (Brace, 2018). Frecuentemente se argumenta que el número de categorías se da en función del tipo de sociedad o de cultura donde se desarrolla la investigación (Lee et al., 2002), aunque lo recomendado y tradicionalmente implementado son escalas de cinco y siete puntos.

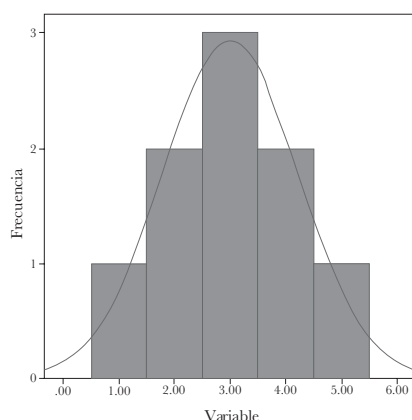
Lo anterior permite que la codificación realizada pueda generar variables cuya naturaleza se asemeje a las variables métricas tanto en valuación como en comportamiento; específicamente, logra que la distribución de los datos sea de tipo normal, lo que posibilita la implementación de las técnicas multivariantes.

2.8. Distribución de datos

Una vez que se definieron las variables de interés para el análisis, se determinó el tipo de escala de medida y se procedió a realizar las mediciones, se necesita analizar las diversas series de datos dentro de la matriz conformada. Este paso es importante, sobre todo si se usaron cuestionarios con escalas tipo Likert; primero, para comprobar que se hizo adecuadamente, y segundo, para asegurarse de que la técnica multivariante se puede ejecutar.

Básicamente, se debe evaluar la distribución que presenta cada serie de datos, debido a que el SEM, especialmente en el CB-SEM, debe usar variables que presentan una distribución normal. La distribución de datos se da en función de la tasa de respuesta que recibe cada una de las categorías incluidas en la escala; dicho de otro modo, se debe evaluar la distribución de preferencias de respuestas para cada categoría. Esta distribución puede visualizarse mediante un histograma que represente las frecuencias de respuestas que recibe un ítem determinado. La Figura 2.8 representa un histograma de frecuencia hipotético de un ítem formulado con una escala de medida tipo Likert.

Figura 2.8. Representación gráfica de la distribución de datos



Fuente: Imagen extraída del software SPSS versión 21.

La imagen presenta una distribución normal, la cual se logra cuando la mayoría de las frecuencias se dan en torno a una categoría media y las ocurrencias disminuyen gradualmente a medida que las demás categorías se alejan del centro.

Para determinar si existe o no una distribución de datos normal se cuenta con diversos métodos, tanto gráficos (diagrama de caja y bigotes, diagrama de hoja y tallos, gráfica Q-Q) como estadísticos (prueba Kolmogorov-Smirnov, prueba Shapiro-Wilk, prueba de asimetría, prueba de curtosis) que verifican la normalidad. En el capítulo 3 se profundiza en la asunción de normalidad para análisis multivariante y en algunas pruebas para comprobarla.

Es importante valorar los métodos gráficos, pero, sobre todo, los métodos estadísticos que verifican la normalidad, debido a que existen umbrales regularmente aceptados para afirmar la presencia o ausencia de una distribución normal. Esto es recomendable, ya que la evaluación mediante métodos gráficos podría ser subjetiva, por eso se aconseja el uso de ambos métodos.

Finalmente, la distribución de datos es útil para identificar patrones o tendencias en las respuestas de la unidad de análisis, que pueden ser efecto del fenómeno de estudio o que las respuestas hayan sido influenciadas; por ello la importancia de evaluar si la distribución corresponde a un adecuado conjunto de datos, es decir, que refleja fielmente la realidad (Barce, 2018).

Capítulo 3. Supuestos básicos para la ejecución de la técnica de ecuaciones estructurales

3.1. Introducción

La aplicación de la técnica de ecuaciones estructurales, al igual que las diversas técnicas del análisis multivariante, exige cumplir con una serie de supuestos y requisitos para lograr buenos resultados. Las ecuaciones estructurales descansan en dos técnicas estadísticas importantes: la regresión lineal múltiple y el análisis factorial o de componentes principales, según sea el caso de un SEM desarrollado por método de covarianzas o de varianzas.

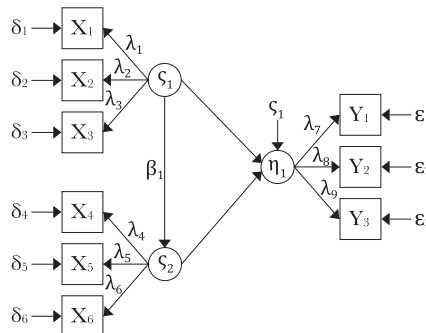
De los supuestos básicos de un modelo de regresión, los modelos SEM deberán observar la linealidad, normalidad, ausencia de multicolinealidad exacta y de autocorrelación de los errores (Hair et al., 1999; Kline, 2005). También es necesario, como en cualquier técnica, que se revisen los valores perdidos (*missing*) y los valores atípicos (*outliers*) que puedan estar presentes en la base de datos y que podrían distorsionar los resultados (Pérez et al., 2013).

La forma de revisión de estos supuestos previamente a la aplicación de la técnica SEM se aborda en este capítulo.

3.2. Linealidad

Un modelo estructural es un grupo de ecuaciones lineales que se ejecutan de manera conjunta y que en su representación matemática deberán estar expresadas de forma lineal. Hablar de relaciones lineales va más allá de una simple inspección matemática de la linealidad de los parámetros a estimar, ya que debe estar soportada por el modelo teórico que se pretende estimar, por lo que el investigador que intenta hacer uso de un modelo SEM tiene que haber comprobado que efectivamente las relaciones que plantea a través de su revisión de literatura obedecen a una relación lineal.

Figura 3.1. Modelo SEM.



Fuente: Elaboración propia.

Para empezar, es necesario comprender la terminología de un modelo estructural. La Figura 3.1 es una representación gráfica de un modelo de ecuaciones estructurales y la Tabla 3.1 enlista la nomenclatura básica utilizada en estos modelos.

Tabla 3.1 Nomenclatura modelo SEM

Variable	Nombre	Descripción
X	Equis	Indicadores (ítems) exógenos
Υ	Ye	Indicadores (ítems) endógenos
ξ	Xi	Variable latente exógena independiente
η	Eta	Variable latente endógena dependiente
β	Beta	Relación entre variables latentes exógenas
γ	Gamma	Relación variable latente exógena con variable latente endógena
ζ	Zeta	Error del modelo estructural
λ	Lambda	Cargas factoriales o pesos de regresión de los indicadores
δ	Delta	Errores de indicadores (ítems) exógenos
ϵ	Épsilon	Errores de indicadores (ítems) endógenos

Fuente: Leyva y Olague (2014)

En la figura anterior puede observarse que los modelos SEM desarrollan múltiples relaciones y que traen consigo la estimación de diferentes parámetros, determinados por: *i)* todos los coeficientes conectan a los indicadores o variables observadas con sus variables latentes, *ii)* los coeficientes se conectan a las variables latentes entre sí, *iii)* por lo que se puede asumir que se encuentran representados por las λ (cargas factoriales), ξ (coeficiente de regresión entre variables latentes exógenas), y *iv)* γ (coeficiente de regresión entre variables latentes exógenas y endógenas) (Bollen, 2009; Manzano, 2017; Hair et al., 1999).

Dentro de los modelos SEM por covarianzas suelen considerarse además dos tipos de parámetros: fijos o restringidos y libres o a estimar, que serán tratados en el siguiente capítulo. Lo interesante al momento de revisar la linealidad del modelo estructural es observar que efectivamente los diferentes coeficientes se encuentren con una potencia de 1 (uno), lo cual no implica que los indicadores (ítems) deban ser también lineales, puesto que las variables observables pueden verse afectadas en su potencia y eso no significa que el modelo deje de ser lineal (Gujarati, 2009).

No obstante, es conveniente aclarar que los modelos SEM han seguido evolucionando, y ello ha llevado a que en los últimos años se haya iniciado el desarrollo de modelos más complejos que responden a la realidad de las ciencias sociales y se acercan a patrones de relaciones no lineales, pero el *software* estadístico aún no incorpora dichas opciones de modelaje, pues requieren cambios a nivel de diseño (Wall, 2012).

3.3. Normalidad

Los modelos estructurales estimados por covarianzas requieren que los datos sigan una distribución normal tanto a nivel univariante como a nivel multivariante. Esto se ha ido suavizando a partir del uso de diferentes métodos para el cálculo de los parámetros del modelo; sin embargo, el método de máxima verosimilitud sigue necesitando el cumplimiento de este supuesto de distribución debido a que el error estándar y el Chi cuadrado calculados podrían estar sobrestimados, generando resultados significativos cuando en verdad no lo son (Gao et al., 2008), por lo que es importante revisar el cumplimiento de este supuesto.

3.3.1. Normalidad univariante

Este concepto se refiere a la distribución normal de cada uno de los indicadores observables que componen el modelo sujeto de estudio. Para la verificación de la normalidad univariante suele utilizarse la inspección visual de los datos a través de un histograma con curva de distribución, así como el *test* de Kolmogorov y la revisión de la asimetría y curtosis de los datos.

El *test* de Kolmogorov-Smirnov es una de las pruebas no paramétricas más recurrentes para constatar si una distribución muestral sigue una distribución normal, por lo que es una prueba de bondad de ajuste. Parte de una hipótesis nula que plantea que los datos siguen una distribución normal, es decir, que existe una homogeneidad en las varianzas de los datos. El valor del estadístico de Kolmogorov-Smirnov no tiene un rango específico para su contraste, sino más bien permite la obtención de un nivel de significatividad o valor *p*, el cual considera el 0.05 como punto de corte. Si el valor del estadístico resulta significativo, por debajo del 0.05, no podría asumirse homogeneidad en la varianza, por lo que se rechazaría la hipótesis nula, implicando que los datos no siguen una distribución normal.

A través de sus comandos, el *software* SPSS permite obtener los valores de esta prueba para cada uno de los indicadores o ítems, siguiendo la ruta Analizar → Estadísticos descriptivos → Explorar → Gráficos → Gráficos con prueba de normalidad. El visor de resultados arrojará el valor de la prueba, como se muestra en la Figura 3.2.

Figura 3.2. Resultado de prueba de normalidad

Pruebas de normalidad						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Estadístico	gl	Sig.	Estadístico	gl	Sig.
INN1	.266	127	.000	.865	127	.000
a. Corrección de la significación de Lilliefors						

Fuente: Elaboración propia a partir del software SPSS versión 21.

Como se observa en el ejemplo, se obtienen dos pruebas, la de Kolmogorov-Smirnov y la de Shapiro-Wilk; ambas buscan analizar la normalidad, pero se elige la interpretación en función al tamaño de la muestra; si es mayor a 50 observaciones se opta por el *test* de Kolmogrov, en tanto que para muestras menores a 50 se usa la prueba de Shapiro-Wilk. En el ejemplo, los grados de libertad (gl) 127 indican una muestra de 127 datos, por lo que se seleccionó el *test* de Kolmogrov para el análisis de la normalidad.

Siguiendo con el ejemplo, el estadístico arrojó un valor de 0.266, con una significancia asintótica de 0.000 y valor $p < 0.05$, por lo que se rechaza la hipótesis nula y no puede comprobarse la normalidad de los datos para la variable observable analizada.

Además del *test* de Kolmogorov-Smirnov para la revisión de la normalidad, se sugiere revisar la asimetría y curtosis de las variables observadas. La simetría se refiere a la distribución equidistante de los datos respecto a su media, mientras que la curtosis compara la distribución de los datos muestreados con la campana de Gauss, para verificar el apuntamiento de la curva de frecuencias, buscando su coincidencia con el modelo de la campana de Gauss (Gujarati, 2009).

Se espera que los coeficientes de asimetría y curtosis se acerquen al 0 (cero), valor que marcaría una distribución simétrica y mesocúrtica y por lo tanto normal. Sin embargo, se suelen considerar como valores aceptables para los estadísticos de asimetría y curtosis los niveles de ± 0.5 , de acuerdo con Bradley (1978). Otros, como George y Mallery (2003), establecen que valores de ± 1.0 son excelentes y como óptimos podrían aceptarse valores de hasta ± 1.6 .

Los estadísticos de asimetría y curtosis pueden obtenerse con el *software* SPSS a través de diferentes rutas, siendo la más habitual Analiza → Estadísticos descriptivos → Descriptivos → Opciones → Asimetría y Curtosis. El visor de resultados presentará el cuadro que se ilustra en la Figura 3.3.

Como se observa, la Figura 3.3 muestra los valores del estadístico de asimetría y de curtosis, así como el error típico calculado para cada uno. También puede

apreciarse que, de acuerdo con los valores dados por George y Mallery (2003), el indicador analizado cuenta con valores excelentes de asimetría y curtosis, siendo inferiores al ± 1.0 .

Figura 3.3. Resultados de asimetría y curtosis

	Estadísticos descriptivos					
	N	Media	Asimetría		Curtosis	
	Estadístico	Estadístico	Estadístico	Error típico	Estadístico	Error típico
INN1	127	3.72	-.771	.215	.225	.427
N válido	127					

Fuente: Elaboración propia a partir del software SPSS versión 21.

En algunos casos los investigadores optan por utilizar la asimetría y curtosis como indicadores de normalidad univariante, aun cuando el valor del estadístico de Kolmogrov-Smirnov no logre asegurar una distribución normal. No obstante, se sugiere la transformación de las variables como una forma de solucionar problemas de normalidad univariante (Hair et al., 1999).

3.3.2. Normalidad multivariante

La normalidad univariante o del indicador es una condición necesaria para la ejecución de la técnica SEM; sin embargo, como se mencionó en párrafos anteriores, en el caso de los métodos por covarianzas, y en particular el de máxima verosimilitud, se requiere además la revisión de la normalidad multivariante, que examina la similitud de distribuciones a lo largo de las distintas variables observadas que componen el modelo de investigación.

Existen diferentes métodos para calcular la normalidad multivariante, como la prueba de Mardia, la Henze-Zirkler, la de Shapiro Wilk Generalizada, entre otras (Porrás, 2015). Siendo de uso generalizado la de Mardia (1975), que parte de la siguiente fórmula:

$$k^2 = [z(G_1)]^2 + [z(G_2)]^2$$

Donde:

k^2 = coeficiente de normalidad multivariante

$z(G_1)$ = asimetría multivariante

$z(G_2)$ = curtosis multivariante

La prueba de Mardia considera que la asimetría y curtosis siguen una distribución asintótica normal a partir del valor de chi cuadrado, estableciendo como hipótesis nula la distribución normal multivariante, por lo que valores de significancia superiores al 0.05 permitirían no rechazar la hipótesis nula y, en consecuencia, confirmar la normalidad multivariada.

Asimismo, Bollen (1989) explica que cuando el coeficiente de normalidad multivariada (coeficiente de Mardia) sea inferior a $p*(p+2)$, en el que p identifica el número de variables observadas, entonces se cuenta con normalidad multivariante.

Actualmente, a través de *macros* puede obtenerse el valor de la asimetría y curtosis multivariada, así como el valor del estadístico de Mardia (Cain et al., 2016). Específicamente para SPSS, DeCarlo (1997) ha diseñado una *macro* para dichos cálculos, que puede encontrarse en el enlace: https://webpower.psychstat.org/wiki/manual/tools/univariate_and_multivariate_skewness_and_kurtosis_all.

Una vez instalada la *macro*, SPSS arroja los resultados de las pruebas con sus niveles de significatividad. Finalmente, Bollen (1989) señala que en el caso de la ejecución de estructurales por el método de máxima verosimilitud, puede garantizarse una adecuada estimación de parámetros y de bondad de ajuste si se logra una distribución mesocúrtica multivariada, aun cuando no se cumpla la hipótesis de normalidad multivariada.

3.4. Multicolinealidad

El supuesto de multicolinealidad, presente en SEM y derivado de los modelos de regresión lineal múltiple, pide que no exista una linealidad exacta entre las variables observables incluidas en el modelo. Es difícil que no exista esa relación, pero se sugiere que no sea una relación fuerte, puesto que la presencia de multicolinealidad trae consigo que los estimadores del modelo se vuelvan ineficientes (Gujarati, 1009).

Existen diferentes formas para detectar problemas de multicolinealidad en las variables, entre las más usuales se encuentra el análisis de correlaciones bivariadas, la revisión del índice de tolerancia y del índice de inflación de la varianza.

En primer lugar se deberán revisar las correlaciones bivariadas entre las variables observadas, considerando que correlaciones superiores al 0.85 podrían indicar problemas fuertes de colinealidad (Pérez et al., 2013). Esto se realiza a través del SPSS siguiendo la ruta Analizar → Correlaciones → Bivariadas.

Enseguida, deberán elegirse las variables observables y el tipo de coeficiente a calcular. Al respecto, la elección dependerá del tipo de medida utilizada para las variables: *i*) correlación de Pearson, cuando las variables sean cuantitativas continuas y con una distribución normal, *ii*) coeficiente de correlación de Spearman, cuando se

utilizan datos ordinales o de intervalo que no satisfacen la condición de normalidad, o *iii*) coeficiente de correlación de Kendall, cuando se trate de datos ordinales y que procedan de una muestra pequeña. Una vez seleccionado se ejecutarán los resultados, que se muestran como en la Figura 3.4.

Figura 3.4. Resultados de correlaciones bivariadas.

		Correlaciones			
		CH11	INN2	EX3	
Rho de Spearman	CH11	Coeficiente de correlación	1.000	.150	.292**
		Sig. (bilateral)	.	.093	.001
		N	127	127	127
	INN2	Coeficiente de correlación	.150	1.000	.439**
		Sig. (bilateral)	.093	.	.000
		N	127	127	127
	EX3	Coeficiente de correlación	.292**	.439**	1.000
		Sig. (bilateral)	.001	.000	.
		N	127	127	127
**. La correlación es significativa al nivel 0.01 (bilateral).					

Fuente Elaboración propia a partir del software SPSS versión 21.

En este caso, lo importante es revisar los valores de los coeficientes de correlación, donde se observa que en el ejemplo ninguna de las correlaciones entre variables es superior a 0.85, lo que hasta el momento permitiría asegurar que no existen problemas de multicolinealidad.

Posteriormente, se revisa el índice de tolerancia y el Factor de Inflación de la Varianza (FIV) a través de una serie de regresiones auxiliares en la que se genera una regresión para cada uno de los indicadores o variables observadas. El índice de tolerancia se refiere a la parte de la varianza de un indicador que no es explicada asociada al resto de las variables observadas; sus valores pueden oscilar entre 0 y 1, por lo que valores de tolerancia menores de 0.20 indican que la variable tiene poco poder explicativo y que comparte una gran cantidad de varianza con el resto de los indicadores, lo cual dejaría entrever problemas fuertes de colinealidad (Hair et al., 2011).

En tanto que el FIV permite conocer hasta qué punto la varianza de un coeficiente se incrementa por la colinealidad, su umbral no debe exceder los 10 puntos (Kutner, 2004). Sin embargo, Diamantopoulos y Siguaw (2006) consideran que este valor no debe ser superior de 3.3 puntos para garantizar la no multicolinealidad entre indicadores.

Estos estadísticos pueden obtenerse en el SPSS siguiendo la ruta Analiza → Regresión → Lineales; en el cuadro de diálogo se selecciona como variable dependiente alguno de los indicadores de la base de datos y como variables independientes se incorporan el resto de los indicadores a usar en el modelo. Enseguida, se selecciona Estadísticos → Diagnóstico de colinealidad. De esta forma se obtendrán los siguientes resultados (Figura 3.5).

Figura 3.5. Estadísticos de colinealidad.

Modelo	Coeficientes ^a						Estadísticos de colinealidad	
	Coeficientes estandarizados		Coeficientes tipificados	t	Sig.			
	B	Error típ.	Beta			Tolerancia	FIV	
1 (constante)	1.677	.493		3.401	.001			
INN2	.053	.108	.046	.489	.626	.820	1.220	
EX3	.353	.119	.281	2.979	.003	.820	1.220	

a. Variable dependiente CH1

Fuente: Elaboración propia a partir del software SPSS versión 21.

Los resultados obtenidos en el ejemplo no muestran problemas graves de multicolinealidad, al conseguir un FIV inferior a 3.3, como denotan Diamantopoulos y Sigauw (2006), y con un índice de tolerancia lejano a los valores inferiores de 0.20 (Hair et al., 2011).

Finalmente, en el caso de encontrarse con problemas de multicolinealidad grave, se sugiere suprimir la redundancia de ítems que explican el mismo concepto.

3.5. Valores perdidos

Los valores perdidos representan la ausencia de respuestas por parte de los sujetos de estudio, y pueden presentarse en dos tipos; el tipo MCAR, que representa valores perdidos que obedecen a un comportamiento aleatorio o al azar; y el tipo MNAR, en el que los valores perdidos siguen un patrón de conducta, por lo que se trata de datos perdidos no aleatorios (Rubin, 1976). Hasta hace poco se consideraba que el contar con valores perdidos era un impedimento para la ejecución de técnicas multivariantes como SEM; sin embargo, hoy en día puede optarse por algunas medidas para la sustitución y manejo de datos perdidos (Montenegro et al., 2015).

Los *softwares* actuales cuentan con opciones que permiten identificar si los valores perdidos obedecen a un comportamiento *aleatorio* o no, también pueden

recuperarlos a través de un proceso de imputación, siempre y cuando tengan un comportamiento aleatorio, sin que se pierda el poder estadístico (Graham, 2009).

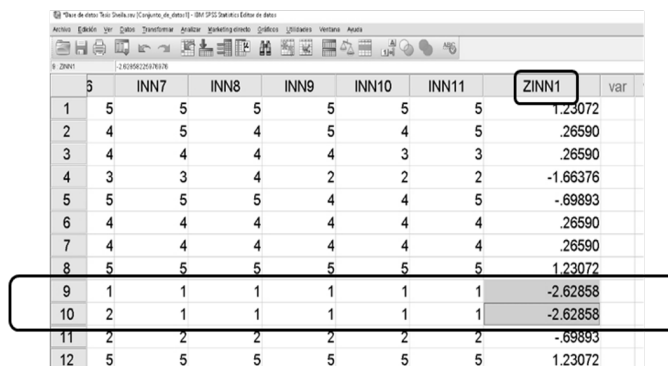
3.6. Valores atípicos

Representan observaciones que se encuentran fuera de los rangos equidistantes del resto de los datos. Son resultado de un posible error de codificación o de un acontecimiento extraordinario que ocasiona para ese sujeto de estudio una observación fuera de los percentiles del resto de la población. Los valores atípicos deben ser revisados de manera univariante y multivariante, ya que la presencia de *outliers* trae consigo estimadores con bajo poder predictivo y con problemas de sesgo. Además, su presencia afecta la distribución normal de los datos.

La revisión de los *outliers univariantes* se puede realizar de manera visual a través del gráfico *box plot* (caja y bigotes) o a través de las puntuaciones tipificadas o estandarizadas para cada una de las variables observables que componen el modelo. En el caso del cálculo de las puntuaciones tipificadas, cuando su puntuación es superior a 2.5 se considera que se encuentra ante la presencia de un valor atípico.

La ruta seguida en SPSS para la identificación de los valores atípicos univariados es: Analizar → Estadísticos descriptivos → Descriptivos; a continuación, se seleccionan cada uno de los indicadores y la opción Guardar valores tipificados como variables. Como resultado se generarán dentro de la base de datos las nuevas variables tipificadas, donde el investigador deberá revisar aquellos casos con puntuaciones superiores a 2.5, que le permitirán identificar el dato como un valor atípico. En la figura 3.6, se muestra el resultado de generar la variable tipificada Z, y en sus casos 9 y 10 tiene valores superiores al 2.5, lo cual es indicativo de la presencia de datos atípicos para esta variable.

Figura 3.6. Variables tipificadas.



	5	INN7	INN8	INN9	INN10	INN11	ZINN1	var
1	5	5	5	5	5	5	1.23072	
2	4	5	4	5	4	5	.26590	
3	4	4	4	4	3	3	.26590	
4	3	3	4	2	2	2	-1.66376	
5	5	5	5	4	4	5	-.69893	
6	4	4	4	4	4	4	.26590	
7	4	4	4	4	4	4	.26590	
8	5	5	5	5	5	5	1.23072	
9	1	1	1	1	1	1	-2.62858	
10	2	1	1	1	1	1	-2.62858	
11	2	2	2	2	2	2	-.69893	
12	5	5	5	5	5	5	1.23072	

Fuente: Elaboración propia a partir del software SPSS versión 21.

La revisión de los valores atípicos multivariantes se realiza para identificar *outliers* en una combinación de indicadores o variables observables. El criterio para examinar los valores atípicos es el de la distancia de Mahalanobis; de este modo se calculan las distancias entre los datos y el valor de Mahalanobis (cero), identificando aquellos puntos lejos del centro como valores atípicos multivariantes. La prueba de Mahalanobis opera con una distribución de chi cuadrado que permite establecer niveles de significatividad.

Por último, en el *software* SPSS es posible guardar los valores de significatividad de la prueba de Mahalanobis como una variable adicional, y otros *softwares* como, el AMOS, presentan en su visor de resultados los valores estadísticos de esta prueba. De esta forma es posible verificar la probabilidad o significancia de la prueba de Mahalanobis, en la que valores inferiores a 0.01 representan una combinación atípica multivariada, lo cual implica que esta observación debe ser eliminada de la muestra.

Capítulo 4. Modelización de ecuaciones estructurales basada en la covarianza

4.1 Introducción

En las últimas décadas en las ciencias sociales se ha intensificado el uso de ecuaciones estructurales, con el objetivo de desentrañar la complejidad de los fenómenos sociales y explicarlos. De las técnicas estadísticas utilizadas para estimar los resultados de un modelo estructural hay dos grandes grupos: el método basado en el análisis de covarianzas y el método basado en el análisis de la varianza.

En este capítulo se aborda específicamente el método basado en el análisis de covarianzas, cuya comprensión inicia recordando que las ecuaciones estructurales forman parte de las técnicas de dependencia del análisis multivariante y que podrían considerarse una combinación de la regresión lineal múltiple, el análisis factorial y el *path analysis* (Kahn, 2006).

La regresión lineal múltiple es una técnica utilizada frecuentemente en econometría y otras áreas afines para estimar los parámetros de un modelo causal o predecir el comportamiento de una variable endógena (Gujarati, 2009). De acuerdo con Wooldridge (2016), la regresión lineal múltiple permite el contraste de una hipótesis fundamentada en una teoría desarrollada previamente. Para estimar los parámetros de regresión existen diferentes métodos basados en la función de la recta; de ellos, los de mayor aceptación son i) los métodos de mínimos cuadrados ordinarios, los cuales se basan en la elección de los estimadores que permitan minimizar los residuos y aproximarlos a 0 (cero), de manera que los parámetros estimados sean lo más próximos a los valores reales, y ii) el método de máxima verosimilitud, que estima los parámetros que tienen mayor probabilidad de precisar los datos observados. En tanto que las ecuaciones estructurales por covarianza tratan de minimizar o ajustar las covarianzas observadas con las covarianzas pronosticadas (Ruíz et al., 2010). De aquí que las ecuaciones estructurales por covarianza sean asociadas de manera directa con el método de regresión.

Los métodos por covarianza son la base en los modelos de análisis factorial, recordando que es una técnica multivariante de interdependencia que busca la identificación de variables latentes en función de una gran cantidad de indicadores (ítems) o variables observables, agrupándolos a partir de sus correlaciones o covarianzas. El análisis factorial permite la identificación de esas variables teóricas difíciles de medir pero que se pueden construir a partir de indicadores operativos que las caracterizan (Pérez y Medrano, 2010). La forma en que estos indicadores se agrupan obedece a que cada uno cuenta con una varianza común, una varianza específica y un error, lo que significa que cada ítem comparte con el resto de los indicadores un rasgo común, pero también tiene un rasgo propio o único y que toda

variable está sujeta a un error de medición. El análisis factorial a través del análisis de la varianza compartida o varianza común identificará las dimensiones de una variable latente sin tomar en cuenta los elementos de varianza única y varianza por el error (Hair et al., 1999).

De esta forma, los modelos por covarianzas comparten como elemento clave la explicación de la varianza común y excluyen de la conformación de las variables latentes la parte de la varianza específica y el error aleatorio de cada variable observable. Es decir, ajustan la variable latente, lo que favorece que los parámetros calculados sean más eficientes y que el modelo resulte bondadoso (Joreskog y Wold, 1982; Anderson y Gerbin, 1988), lo cual se refleja en una mejor explicación de la variable endógena, aunque ello conlleve la pérdida de la fiabilidad de su predicción. El método por covarianzas descompone la covarianza existente entre las variables de un modelo estructural para a partir de ello estimar los parámetros, logrando definir a qué se debe dicha covarianza en función de los argumentos teóricos (Medrano y Muñoz-Navarro, 2017). En resumen, el análisis por covarianza es un método estadístico que busca analizar la relación entre variables, eliminando el error que no es controlado por el investigador y considerando las variaciones de cada una de las observaciones o casos de estudio analizados (Badii y Castillo, 2007).

Visto lo anterior, puede entenderse por qué para muchos investigadores el uso de ecuaciones estructurales está intrínsecamente ligado al análisis por el método de covarianzas que puede desarrollarse a través de los paquetes informáticos como AMOS, EQS, LISREL, entre otros (Hair et al., 2011).

El propósito de este capítulo es profundizar en el modelado de ecuaciones estructurales a través del método basado en el análisis de covarianzas, identificando su utilidad, fundamento, sus métodos de estimación y el proceso seguido para la comprobación de sus parámetros.

4.2. Elección del método SEM por covarianzas

Como se mencionó en la introducción, los modelos estructurales pueden desarrollarse por dos métodos, el de covarianzas y el modelo por varianzas. La elección del método dependerá en buena medida de las características de la investigación empírica a realizar. En este apartado se procura responder a la pregunta ¿cuándo es conveniente utilizar el modelo SEM por covarianzas?

En primer lugar, se debe tener claro el objetivo de la investigación, discerniendo si es de alcance causal y qué tanto bagaje teórico se tiene sobre las relaciones a analizar. Los estudios que cuentan con una fundamentación teórica suficiente pueden desarrollarse a través del método por covarianzas (Anderson y Gerbing, 1988). Si se pretende comprobar cuánto el modelo hipotetizado se asemeja

a una realidad a través de la revisión de los indicadores y constructos que la integran y contrastar las relaciones que de ellos surjan, se recurre a medidas de bondad, que permiten identificar el grado de ajuste del modelo (Castro et al., 2007). Cuando se trata de probar modelos con poca evidencia teórica o empírica se recomienda el uso del SEM por el método de varianzas (Jöreskog y Wold, 1982). De acuerdo con Beltrán y Espejel (2016), con un modelo por covarianzas se busca confirmar una teoría previamente estudiada y basada en la realidad.

Asimismo, los modelos estructurales a través de covarianza se caracterizan por su orientación hacia la causalidad de las relaciones estudiadas más que en su poder predictivo. Esto se debe a la forma de operar del SEM por covarianzas, basado en el análisis factorial, centrado en explicar la varianza común de los constructos (Medrano y Muñoz-Navarro, 2017), mientras que los modelos contrastados por Partial Least Squares (PLS) se basan en el valor predictivo al buscar maximizar la varianza explicada del modelo a través del coeficiente de determinación (R^2) (Castro et al., 2007).

Uno de los requisitos importantes al elegir trabajar un modelo por covarianzas tiene que ver con el tamaño de la muestra u observaciones mínimas de las que dispone el investigador, lo que en cierta medida se convierte en una limitación, puesto que el SEM por covarianzas requiere una muestra de por lo menos 200 observaciones para su correcta ejecución (Beltrán y Espejel, 2016). Sin embargo, otros como Hair et al. (1999) y Hu y Bentler (1999) explican que el tamaño de la muestra dependerá del número de indicadores y de parámetros a estimar por el modelo, y advierten que la muestra mínima para que no exista error en el ajuste del modelo debe ser por lo menos de 250 a 500 observaciones cuando el número de variables latentes es amplio, por lo que se recomienda que en modelos SEM por covarianzas se trabaje con un número reducido de variables latentes, para evitar estos sesgos por tamaño.

Una de las ventajas del modelo SEM por covarianzas es que permite el establecimiento de relaciones no recursivas o bidireccionales (Jöreskog y Wold, 1982), lo cual no es hasta el momento posible en modelos por varianza, por lo que si el modelo del investigador plantea una relación bidireccional, el método por covarianza será el adecuado. Finalmente, los modelos por covarianza solo pueden operar con variables latentes construidas a través de indicadores de tipo reflectivo, así que, si el modelo a probar por el investigador involucra variables latentes de tipo formativo, habría que optar por un modelo a través de PLS (Chin, 1998).

4.3. Tipos de relaciones en los modelos estructurales

4.3.1. Causalidad

El término causalidad resulta complejo en la época moderna y sobremanera en las ciencias sociales, por lo que resulta conveniente iniciar con la definición conceptual. Causalidad hace referencia al origen de algo, ese algo puede ser un elemento de la física, de la estadística, de las ciencias sociales, de la conducta, etcétera. Se habla de causalidad cuando se pretende explicar los determinantes de un fenómeno, es decir, comprender el funcionamiento de un sistema a través del estudio de sus componentes (García-Ferrer, 2008). Es un posicionamiento explicativo de los elementos que conducen a la ocurrencia de un fenómeno. El conocimiento de una relación de causalidad permite, mediante el pensamiento deductivo, predecir un efecto. (Simon, 1988).

La causalidad implica la presencia de una variable causal y de una variable resultado o efecto. En la mayoría de las ocasiones estas relaciones traen consigo un sentido unidireccional, considerando que es la causa la que determina el efecto, de manera que dicha relación da lugar a una causalidad recursiva, representada gráficamente a través de una flecha de un solo sentido. La flecha genera un coeficiente *path*, similar a los coeficientes beta de una regresión lineal múltiple, que representan el cambio que genera la variable causal en la variable efecto (Aaron y Aaron, 2001).

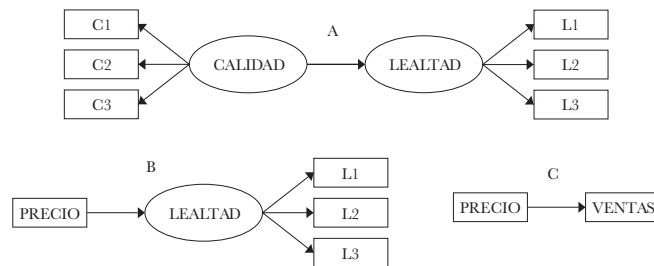
En los modelos estructurales las causas y los efectos pueden ser medidos de dos formas: a través de indicadores observables o a través de una variable latente, siempre y cuando una de las variables sea medida a través de un constructo latente, como se muestra en la Figura 4.1. Lo anterior se recalca para no confundir un modelo SEM con un modelo de *path analysis*, donde las causas y efectos son medidos a través de variables observables, en tanto que la riqueza de los modelos estructurales radica en la captura de conceptos complejos a través de indicadores que dan lugar a una variable latente o variable no observable (Pérez et al., 2013).

Es importante enfatizar que en un modelo SEM, aun cuando se encuentran presentes los elementos de una relación causal, lo que se prueba es la relación entre las variables observadas y la variable latente, pero SEM no aporta las pruebas suficientes para demostrar una causalidad; no obstante, los resultados arrojados por los estimadores de SEM permiten contrastar las hipótesis causales basadas en asociaciones directas (Bollen y Pearl, 2013). En este sentido, algunos autores advierten que la causalidad sólo es posible comprobarla a través de modelos experimentales y no a través de modelos con datos transversales; sin embargo, los modelos SEM por covarianzas son contruidos con fundamentación en una conceptualización teórica que permite suponer relaciones entre constructos, y las covarianzas obtenidas a

partir de los datos podrían servir para contrastar las teorías y derivar los efectos causales (Ruiz et al., 2010).

Conviene concluir que, aunque se opte por los modelos estructurales, por su similitud con las regresiones múltiples, un modelo estructural no probará causalidad, sino que permitirá la selección de las hipótesis que de acuerdo con la evidencia empírica tienen un soporte y desechar aquellas que no lo tengan (Aaron y Aaron, 2001). Considerando que cualquier teoría causal es factible de ser rechazada si los datos observados no la soportan estadísticamente (Medrano y Muñoz-Navarro, 2017).

Figura 4.1. Relaciones de causalidad



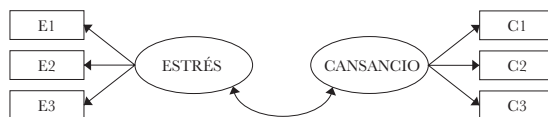
Fuente: Elaboración propia.

4.3.2. Covarianza o correlación

La correlación mide el grado de asociación que existe entre dos variables latentes. Aunque pudiera pensarse que la covariación y la correlación son sinónimos, no lo son. La covarianza se refiere a la asociación que guardan entre sí las variables, es decir, permite observar si la variable A y la variable B tienen una asociación positiva (directa) o negativa (indirecta), implicando que si una aumenta, la otra aumenta, o viceversa. Mientras que la correlación indica la intensidad de la asociación entre las variables. Aunque podría decirse que son medidas análogas, la diferencia principal es que la correlación es una medida estandarizada, por lo que sus valores oscilan entre el -1 y 1, y cuando se encuentran próximos a 0 se habla de una nula asociación, en tanto que los valores que se aproximan a la unidad reflejan una fuerte correlación. Mientras que la covarianza es un índice que mide la variación conjunta de dos variables expresada en su escala de medida original, por lo que no tiene un valor mínimo ni máximo, dificultando su comparación y, consecuentemente, su interpretación.

El cálculo de la correlación es el cociente de la covarianza de las variables x , y dividido entre el producto de las desviaciones estándar de las variables x , y (Gujarati, 2009). La covarianza y su correlación se representan en los modelos SEM con una flecha de doble punta, como se observa en la Figura 4.2.

Figura 4.2. Relaciones de correlación-covarianza



Fuente: Elaboración propia.

La correlación se calcula generalmente por tres estimadores, y la selección del método depende del tipo de variable analizada. En el caso de variables continuas suele utilizarse el coeficiente de correlación de Pearson; cuando se manejan variables cualitativas de tipo ordinal se recomienda el cálculo de la correlación a través del coeficiente de Tau de Kendall, en tanto que cuando se trate de variables en su mayoría ordinales o continuas pero que no presentan una distribución normal es preferible el uso del coeficiente de correlación de Spearman (Morales y Rodríguez, 2016). En el SEM por covarianzas las correlaciones son calculadas a partir de los coeficientes de Pearson.

Es importante asentar que si bien toda causalidad trae implícita una correlación, no toda correlación implica una causalidad (García-Ferrer, 2008). Un ejemplo: dos variables latentes, como el estrés y el cansancio, donde puede existir una asociación entre ambas, inclusive fuerte y directa, pero eso no conlleva a que una sea la causa y la otra el efecto; puede ser que el estrés cause el cansancio o que el cansancio sea el causante del estrés o que ambas sean causadas por otra variable no expresada en el modelo. Distinto es cuando se ha demostrado la causalidad entre las variables, ya que, de ser así, la causalidad siempre vendrá antecedida de una correlación.

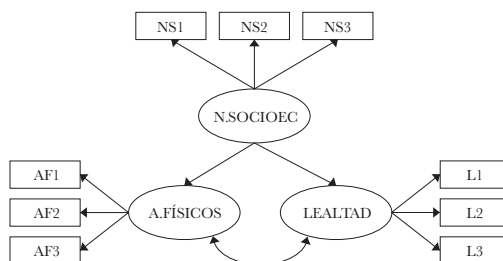
4.3.3. Relación espuria

El término espurio significa algo falso o no auténtico. En estadística, una relación espuria se refiere a la presencia de correlación significativa de dos variables cuando teóricamente no puede ser explicada. Esta falsa correlación es causada por una tercera variable que no se encuentra en el modelo (Simon, 1954). En su momento puede abonar un poco a la teoría de la falta de causalidad en presencia de correlaciones significativas, ¿por qué? En el apartado anterior se explicó que no toda correlación trae consigo causalidad, entonces, ¿cómo es posible que se encuentre asociación entre variables que no deberían estar relacionadas? Esto puede ser porque la relación es falsa y está siendo provocada por una tercera variable que afecta la relación entre ambas. Dicho de otra forma, una correlación significativa que posteriormente busca probar causalidad y no resulta con efectos significativos puede involucrar una relación espuria.

Un ejemplo clásico lo encontramos en el estudio realizado por Neyman en 1952 (citado por Lahura, 2003), quien encontró una fuerte relación entre el número de nacimientos de niños y la población de cigüeñas en diferentes regiones. Sería absurdo suponer que esta relación tenga un fundamento teórico o que sea cierta, el sentido común lo impide (Simon, 1954); no obstante, el coeficiente de correlación lo señala. ¿Cómo se podría demostrar que esta situación carece de certeza? La respuesta se encuentra en la inclusión de una tercera variable que puede ser la causa de esta falsa relación. Siguiendo con el ejemplo, se incorporó al estudio la variable número de habitantes, puesto que, entre más habitantes, más edificios, y con ello se explica por qué aumenta la población de cigüeñas, al anidar en los edificios.

La forma para determinar si la correlación es falsa es mediante los coeficientes de correlación parcial, que permiten aislar el efecto de una variable dentro de una relación, para de esta manera identificar la variabilidad entre el resto de las variables (Kronmal, 1993; Simon, 1954). Para valerse de la correlación parcial es necesario contar por lo menos con tres variables, sean estas X , Y , Z , donde X representa los atributos físicos de las personas, Y la lealtad hacia una marca y Z el nivel socioeconómico. La relación entre estas variables latentes resulta significativa, por lo que se procede al cálculo de las correlaciones parciales. Eliminando la variabilidad de Z , la correlación parcial mostrará si efectivamente X y Y (los atributos físicos de las personas y la lealtad hacia la marca) tienen una asociación significativa, y en caso de ser espuria, el coeficiente calculado será no significativo. La Figura 4.3 presenta un ejemplo de relación espuria.

Figura 4.3. Relación espuria



Fuente: elaboración propia.

4.3.4. Relación causal directa e indirecta y efectos totales

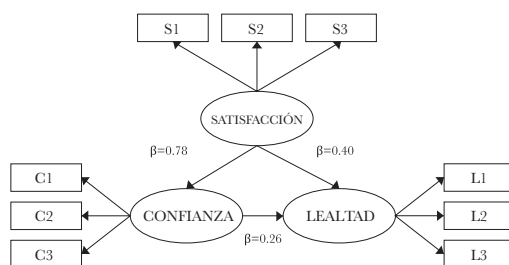
Una de las ventajas principales de la técnica SEM sobre la regresión múltiple es que además de examinar el efecto de una serie de variables latentes independientes en una variable latente dependiente, permite adicionar terceras variables y observar

el efecto de la variable dependiente inicial en una nueva variable dependiente, logrando un análisis simultáneo de varias relaciones de dependencia (Hair et al., 2011). Por lo tanto, en este tipo de modelos se puede hablar de dos momentos de afectación en un fenómeno estudiado; en un primer momento se reconocen los efectos de X en Y , y en un segundo momento se examinan los efectos de Y en W ; no obstante, en este flujo de relaciones X logra también una afectación de forma indirecta en W (Schumacker y Lomax, 2010). Lo anterior es posible visualizarlo con el *path analysis* o gráfico de senderos del que se apoya la técnica SEM.

Por lo que en un modelo estructural pueden observarse, si así se plantea, tres efectos causales: el efecto directo, el indirecto y el total. El efecto directo se refiere a la magnitud del impacto de una variable independiente en una variable dependiente, es por tanto una relación causal, mientras que las relaciones indirectas se encuentran cuando una tercera variable se interpone en la relación entre dos variables, modulando o potencializando el efecto entre ellas (Baron y Kenny, 1986). La suma de los efectos directos e indirectos es conocida como efecto total.

Por ejemplo, supóngase las variables Confianza, Satisfacción y Lealtad, como se aprecia en la Figura 4.4. En este modelo la variable latente Lealtad es una variable dependiente que recibe el efecto directo de las variables Satisfacción y Confianza, pero a su vez recibe un efecto indirecto de la variable Confianza, en tanto el efecto total será la suma del efecto indirecto de Confianza, el efecto directo de Satisfacción y el efecto directo de Confianza.

Figura 4.4. Efectos indirectos y efectos directos



Fuente: Elaboración propia.

Para conocer la magnitud del efecto de una variable sobre otra se toma el coeficiente estandarizado, que es similar al parámetro de la pendiente estimada de la regresión múltiple. Todo modelo estructural presentará sus coeficientes estandarizados y con ellos podrán determinarse los efectos directos, indirectos y totales que recibe una variable estudiada. Los efectos directos son equivalentes a los coeficientes

estandarizados de una relación directa, en tanto que para el cálculo de la magnitud de los efectos indirectos se recurre al producto de los coeficientes estandarizados de los efectos directos contenidos en la línea causal que lleva a la relación entre dos variables (Arbuckle, 2003).

En el modelo estructural del ejemplo y considerando que SEM parte de una regresión múltiple, los efectos directos para Lealtad podrían plantearse por la ecuación $\text{Lealtad} = b\text{Confianza} + b\text{Satisfacción} + u$. Lo anterior sería $\text{Lealtad} = 0.40 \text{ Satisfacción} + 0.26 \text{ Confianza}$.

En cuanto a los efectos indirectos, se observa que Confianza tiene un efecto indirecto siguiendo la ruta Confianza-Satisfacción-Lealtad, ese efecto sería de 0.312, que corresponde al producto de sus coeficientes estandarizados (0.78×0.40). Sin embargo, para un estudio más a detalle de los efectos indirectos debe revisarse el trabajo de Baron y Kenny (1986).

4.4. Métodos de estimación en los modelos por covarianza

Como ya se ha mencionado, SEM puede ser calculado por covarianzas o por mínimos cuadrados parciales. En el caso de los modelos por covarianzas, se busca estimar los parámetros a partir de la descomposición de la matriz de covarianzas, y tal descomposición puede lograrse por medio de distintos métodos de estimación. De acuerdo con Medrano y Muñoz-Navarro (2017), la descomposición de la varianza sigue dos reglas básicas:

- La suma de los efectos directos, indirectos y espurios representa la covarianza entre dos variables.
- La varianza de la variable dependiente es igual a la varianza explicada por las variables independientes del modelo más la varianza del error.

Ahora bien, los métodos de estimación que siguen los modelo por covarianzas son el de Máxima verosimilitud (MV), Mínimos cuadrados generalizados (GLS), Mínimos cuadrados no ponderados (ULS) y Distribución Asintótica libre (ADF). Estos métodos persiguen el mismo propósito de las covarianzas pronosticadas y las observadas, solo que la diferencia radica en la forma utilizada para ajustar las diferencias de las covarianzas muestrales (Medrano y Muñoz-Navarro, 2017). A continuación, se explica el proceso de selección del tipo de método.

4.4.1. Máxima verosimilitud (ML)

Es el método más utilizado en las estimaciones por covarianzas, sin embargo, no siempre es el más recomendable, depende de las características de los datos y su distribución. Este método fue presentado por Fisher entre 1912 y 1922 y busca estimar los parámetros que tengan mayor probabilidad de lograr los valores de la

muestra observada (Gujarati, 2009). Estos estimadores tendrán las características asintóticas de ser consistentes, insesgados y eficientes; en otras palabras, resultan sólidos, sus valores medios coinciden con los valores auténticos y la varianza del estimador es la más pequeña posible.

La ejecución del modelo estructural por este método descansa en el cumplimiento de los supuestos de normalidad univariada y multivariante y de multicolinealidad imperfecta, así como en la ausencia de valores atípicos y valores perdidos. Del mismo modo, se solicita el uso de una medida cuantitativa continua o por lo menos intervalar (Schermelleh et al., 2003). Otro requisito para la ejecución a través de ML es el tamaño de la muestra, el cual debe ser de 200 observaciones (Bollen, 1990). La falta de estos supuestos ocasionaría que los estimadores, los errores estándar y el ajuste del modelo resulten poco armoniosos y de bajo poder explicativo (Fabrigar et al., 1999; Schumacker y Lomax, 2010).

Estos requerimientos han sido analizados por diversos investigadores, que han establecido posturas distintas respecto al método de máxima verosimilitud. En primer lugar, en cuanto a la normalidad multivariante, algunos autores, como Muthen (1993), aseguran que niveles bajos de este supuesto no afectan de manera significativa el ajuste de los estimadores y siguen siendo consistentes. Al respecto, Levy-Manguin y Varela (2006) explican que distintos *softwares* estadísticos, como el AMOS, permiten la estimación del modelo SEM aun con la ausencia de normalidad multivariada, aunque existen otros métodos de cálculo que también pueden hacer los ajustes pertinentes.

En cuanto al uso de medidas ordinales mediante las escalas y medición de actitudes, como es común en las ciencias sociales, complican la estimación asintótica de los datos (Bollen, 1990) a través de máxima verosimilitud. Pell (2005) explica que es recomendable realizar una transformación como si se tratase de una variable cuantitativa; esta transformación puede ser por medio de un escalamiento o de las medidas de transformación más comunes como lo son el inverso, la transformación exponencial o la transformación logarítmica. La Tabla 4.1 presenta las principales transformaciones utilizadas para lograr la normalidad valiéndose de la disminución de problemas de asimetría y curtosis.

Tabla 4.1. Transformación no lineal de variables

Transformación	Función
Raíz cuadrada (raíz)	$\sqrt{x_i}$
Logaritmo neperiano (ln)	$\ln x_i$

Transformación	Función
Inversa (INV)	$1/X_i$
Cuadrática (exp)	X_i^2

Fuente: Basado en la escala de transformación de Tukey.

Sobre el uso de variables en escalas tipo Likert se ha escrito mucho, sin llegar a un consenso sobre si pueden ser tratadas como medidas de intervalo. No es el objetivo de este capítulo profundizar en ello, pero debe aclararse que la literatura hace alusión a una diferencia entre el uso de la terminología escala y formato de respuesta, considerando que en la mayoría de las ocasiones el objetivo en ciencias sociales es capturar un concepto abstracto a través de una medida que posteriormente es validada; por tanto, se refiere a un formato de respuesta que es sujeto a un tratamiento de medida de intervalo (Carifio y Perla, 2007; Knap, 1990), lo cual es reforzado dado que las opciones de respuesta son iguales o superiores a 5, que los indicadores de bondad de ajuste, los parámetros estimados y los errores estándar no presentan grandes distorsiones, por ende, es un método robusto (Finney y DiStefano, 2006).

Lo anterior obliga a preguntar ¿puede utilizarse una base de datos medida como formato de respuesta Likert en un modelo estructural resuelto por Máxima Verosimilitud? La respuesta es sí, siempre que cumpla con la normalidad multivariante o se demuestre su validez. También se puede optar por otros métodos menos exigentes con el supuesto de normalidad.

Otro de los elementos a considerar para la ejecución del método de Máxima Verosimilitud es el tamaño de la muestra. Debido a que se trata de estudios que buscan explicar causalidad, se sugiere que la muestra mínima sea de 200 observaciones (Hoelter, 1983). Otros investigadores, como Hair et al. (1999), recomiendan que para definir adecuadamente el tamaño de la muestra por este método, se verifique la identificación del modelo y se considere un tamaño de la muestra igual a 5 observaciones por ítem incluido en el instrumento de recolección, o bien por cada parámetro estimado en el modelo. En caso de que el tamaño de la muestra no sea el adecuado por ML, puede recurrirse al *bootstrapping*, que permite la estimación de los errores estadísticos (Efron, 1981); no obstante, no se debe descuidar que para muestras inferiores a 100 los estimadores obtenidos pierden robustez.

4.4.2. Mínimos cuadrados generalizados (GLS)

Este método es similar al de máxima verosimilitud, con la diferencia de ser menos restrictivo en cuanto a la normalidad multivariante de los datos. Se recomienda

cuando el modelo incluye indicadores o variables medidas de forma distinta a una medición cuantitativa continua (Medrano y Muñoz-Navarro, 2017). Este método conserva las características asintóticas del de Máxima Verosimilitud, por lo que los estimadores resultan consistentes e insesgados; el único requisito para ello es que los datos cuenten con una distribución univariante con simetría y curtosis próxima a 0 (cero), de modo que se observe la normalidad univariante de esos datos (Schumacker y Lomax, 2010). El GLS se desaconseja cuando la complejidad del modelo es alta, por lo que cuando el número de parámetros a estimar sea mayor a 30 deberá optarse por otro método de estimación (Muthén, 1989).

4.4.3. Mínimos cuadrados no ponderados (ULS)

Este método es de estimación estructural, aconsejable cuando las variables analizadas no siguen una distribución normal multivariante ni univariante, por lo que se recomienda cuando la medición de las variables se realizó con una mezcla de medidas, contado con variables de tipo continuo o intervalar y alguna medida categórica u ordinal. Lo que hace este tipo de método es transformar la matriz de covarianzas en una matriz policórica (Bollen, 1990).

Una de las ventajas principales de las estimaciones por mínimos cuadrados no ponderados es que permite una solución adecuada con muestras pequeñas o inferiores a 200, llegando a un valor mínimo de 100, y en especial cuando se poseen pocos elementos (parámetros) de estimación (Freiberg et al., 2013). Este procedimiento es aplicable a todo tipo de variables categóricas, aunque su mejor ajuste y estimación se logra cuando se trabajan ítems dicotómicos (Muthén, 1983).

4.4.4. Distribución Asintóticamente libre (ADF)

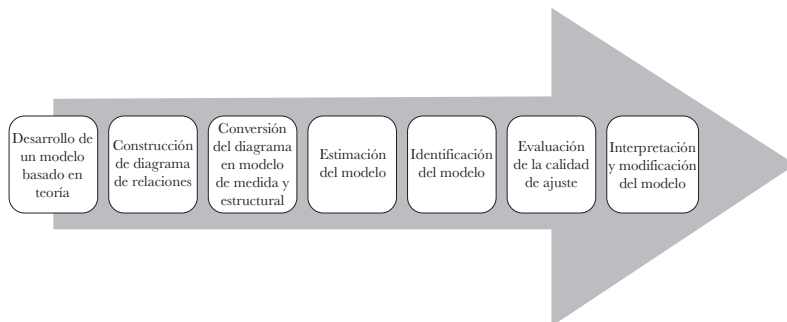
Este es el método de menor uso para la estimación de estructurales por covarianzas. Se aplica cuando no se logra la normalidad univariante ni multivariante, sobre todo cuando los datos presentan una distribución muy asimétrica. Su mayor limitación es que para que las estimaciones sean confiables se requiere una muestra grande, mayor de 500 observaciones (Hu et al., 1992).

4.5. Etapas para la construcción de SEM por covarianzas

El modelo de ecuaciones estructurales es un análisis multivariante que analiza la relación de variables latentes y observadas entre sí y sus relaciones cruzadas. Tiene mayor amplitud de análisis que un modelo de regresión puesto que permite calcular los efectos del error de medida sobre los coeficientes estructurales (Hair et al., 1999). La construcción de un modelo de ecuaciones estructurales trae consigo un proceso complejo que implica una adecuada fundamentación del modelo, su correcta

medición, la comprobación de las relaciones causales propuestas y la evaluación de su ajuste. El proceso seguido para el desarrollo de un modelo SEM por covarianzas consta, de acuerdo con Hair et al. (1999), de siete pasos, como se observa en la Figura 4.6.

Figura 4.6. Proceso SEM



Fuente: Hair et al. (1999).

No obstante, este proceso puede ser analizado a partir de tres grandes apartados, propuestos a continuación.

4.5.1. Paso 1. Especificación teórica del modelo de medida y del modelo estructural

Como ya se ha mencionado, los modelos SEM por covarianzas buscan comprobar relaciones causales a través de la identificación de aquellas hipótesis que tienen el soporte empírico para corroborar una teoría; por tanto, en primera instancia deberá contarse con antecedentes causales que apoyen una relación entre variables. Cuando no se tiene la certeza de la forma en que se da una relación, no se sugiere el uso de métodos de covarianza, puesto que realizar un Análisis Exploratorio que dé luz a los datos del investigador conllevaría la utilización de métodos basados en PLS (Hair et al., 1999). En el uso de modelos SEM por covarianza se debe tener cuidado de que no queden fuera variables predictivas importantes para explicar la variable endógena, ya que se tendría un error de especificación que generaría sesgo en los resultados.

Luego de haber revisado teóricamente las fundamentaciones del modelo, se procede a la elaboración del diagrama *path*, estableciendo gráficamente las relaciones a estudiar en el modelo empírico. En el proceso de modelado gráfico es importante que se tenga claro qué constructos son tratados como endógenos y cuáles como exógenos, pues las variables latentes endógenas deberán tener presente su error de

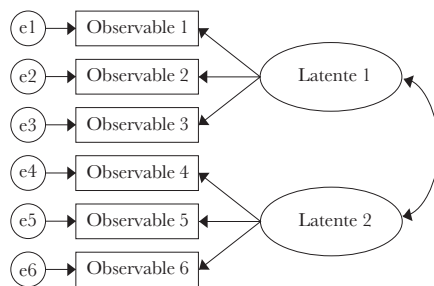
medición, el cual es agregado por considerar que hay una parte de la variación de la variable endógena que no es explicada por las variables independientes incluidas en el modelo (Cupani, 2012).

Debe recordarse que en los modelos SEM por covarianzas se pueden presentar tres tipos de relaciones: *i*) causales unidireccionales y bidireccionales, *ii*) covarianzas o correlaciones, *iii*) relaciones directas e indirectas, considerando a su vez que los modelos SEM trabajan solo bajo el supuesto de linealidad de las relaciones dibujadas (Loehlin, 1987).

4.5.2. Paso 2. Valoración del modelo de medida

Una vez definido el diagrama *path*, de manera teórica se revisa el modelo de medida, que consiste en la verificación de la fiabilidad y validez de la escala de medición utilizada para la operacionalización de las variables latentes. El objetivo del modelo de medida es identificar la idoneidad de los ítems o indicadores seleccionados para construir las variables latentes (Byrne, 2001), las cuales no pueden ser evaluadas de manera simple, a través de un solo indicador. Por lo general, cuando se trabaja con percepciones en las ciencias sociales, como pueden ser de satisfacción o análisis de desempeño, se requiere de más de un ítem para ser medidas, y pueden ser validadas a través de análisis factorial confirmatorio (Schreiber et al., 2006). En la Figura 4.5 se puede ver un modelo con 2 variables latentes con tres observables cada una.

Figura 4.5. Ejemplo de un modelo de medida



Fuente: Elaboración propia, con base en Schreiber et al. (2006).

Según Mulaik y Millsap (2000), el modelo de medida se analiza para validar las variables latentes del modelo estructural mediante un análisis factorial confirmatorio, que traza correlaciones entre los constructos latentes presentes en el modelo, como se observa en la figura 4.5.

El modelo de medida está basado en un modelo multifactorial (Lomax y Algina, 1979) en el que se incluyen todos los indicadores que se espera construyan un concepto. Para cada uno de los constructos se recomienda un mínimo de 3 indicadores, aunque generalmente oscilan entre 3 y 7 (Ding, Velicer y Harlow, 1995). Una práctica común dentro de la definición de los constructos consiste en fijar la fiabilidad de uno de los indicadores con el valor 1, y deberá ser aquel que ha sido utilizado previamente y del que se puede asegurar que cuenta con fiabilidad (Hair et al., 1999). Este indicador, por tanto, se considera que tiene una restricción en el parámetro. También es importante destacar que cada indicador posee un error propio, el cual no debe correlacionarse con los errores del resto de los indicadores; esta característica permite que las covarianzas resultantes de los indicadores expliquen de manera única el constructo generado.

En los modelos de medida es necesario revisar la fiabilidad y validez de los constructos, para asegurar su buena representación (Campbell y Fiske, 1959; Cole y Maxwell, 1985). Al respecto, la validez del constructo involucra evidencias de validez convergente y discriminante (Cole y Maxwell, 1985).

4.5.2.1. Fiabilidad interna

De acuerdo con Nunnally (1978), la fiabilidad es el grado en que la medición de un constructo arroja los mismos resultados al repetirse, es decir, implica la consistencia de la escala. Esto significa que el constructo debe ser medido de manera precisa. Frías-Navarro (2019) explica que se logra fiabilidad cuando el instrumento obtiene el mismo tipo de respuestas al ser aplicado en repetidas ocasiones a un mismo grupo de individuos. Se debe revisar la fiabilidad del ítem y la fiabilidad del constructo.

Fiabilidad del ítem

En el análisis factorial confirmatorio se hace la prueba de la fiabilidad de los indicadores observados; asimismo, se emplea el modelo de medida para explorar el análisis de las correlaciones y la covarianza entre las variables no observables, calculándose las cargas factoriales de los ítems para cada constructo. Cada vez que se elimina un ítem o se modifica una ruta se cambian estas cargas y los coeficientes de modificación. Por otro lado, este análisis permite medir la contribución de cada indicador en el constructo o factor.

Para el análisis de la fiabilidad interna de los ítems, se revisan los pesos de regresión de cada indicador y las correlaciones entre los elementos que componen la escala. En primer lugar, se revisan las correlaciones entre indicadores a través de la correlación corregida entre elementos; este indicador es conocido como el Índice de homogeneidad y evalúa la correlación entre el indicador y la puntuación

total de la escala. Se espera que exista una alta correlación entre los indicadores de un constructo, por lo que valores superiores a 0.2 podrían garantizar que el ítem analizado tendrá una puntuación alta en el factor; es decir, que su peso de regresión será alto (Cervantes, 2005).

Además, la fiabilidad individual de los indicadores se revisa a través del análisis factorial confirmatorio, verificando los valores de los pesos de regresión. De acuerdo con Hair et al. (1999), los pesos de regresión que no sobrepasen el 0.500 deben ser eliminados del modelo estructural, aunque debe considerarse que posteriormente esta medida será elevada al cuadrado (comunalidades), para revisar la validez discriminante, y que el peso mínimo obtenido de esta elevación cuadrática o comunalidad deberá lograr un valor superior a 0.700, por tanto, un peso de regresión inferior al 0.840 no alcanzará dicho valor.

Es conveniente distinguir entre los términos pesos de regresión y carga factorial; el primero se refiere a los pesos de las variables observadas, en tanto que las cargas factoriales aluden a los pesos de los constructos o variables latentes. Sin embargo, algunos *softwares* para modelado de SEM o de confirmatorios, como AMOS, no hacen una diferencia entre ambos y sólo utilizan el término pesos de regresión.

Fiabilidad del constructo

Para evaluar la similitud de las variables observables que conforman una variable latente, es decir, que efectivamente los ítems analizados pertenecen a cada constructo, se revisa la fiabilidad de la variable latente. Esto se verifica a través del coeficiente alfa de Cronbach (α) y el coeficiente omega (ω) (Heise y Bohrnstedt, 1970).

Cronbach (1949) propone el índice Alpha de Cronbach como una medida para explicar la varianza debida al factor común entre los ítems. Este coeficiente puede oscilar entre un valor de 0 y 1, por lo que valores próximos a la unidad implican una mayor correlación entre los ítems que constituyen un constructo.

La fórmula utilizada para su cálculo es:

$$\alpha = \frac{k}{k-1} \left[1 - \frac{\sum s_i^2}{s_T^2} \right]$$

Donde:

k = número de ítems

$\sum s_i^2$ = sumatoria de varianza de cada ítem

s_T^2 = Varianza total del constructo

Como se puede deducir de la fórmula, mientras más ítems tenga el constructo, mayor será el coeficiente Alpha obtenido. Esto, entre otros puntos, ha llevado a cuestionar la solidez de este tipo de indicador (Green y Yang, 2009; Huysamen, 2007; Dunn et al., 2013) al no cumplir con el supuesto de equivalencia propia de una escala, referente a que las cargas o pesos de regresión entre los ítems deben ser estadísticamente equivalentes; también, por la falla al supuesto de unidimensionalidad, el cual busca que los ítems midan un solo rasgo de la variable (Cho, 2016). Aun con estas críticas, el alfa de Cronbach es el indicador más utilizado a nivel internacional para la revisión de la fiabilidad de un constructo (Dunn et al., 2013).

Los rangos de aceptación o valores críticos para el coeficiente alfa de Cronbach, según Nunnally (1978), deberán ser superiores a 0.700. Al respecto, George y Mallery (2013) señalan que este valor es el mínimo aceptable y que valores inferiores se vuelven cuestionables e invalidan la seguridad de la escala. Sin embargo, el mismo Nunnally (1978) indica que en el caso de investigaciones aplicadas que implican per se fuertes fundamentos teóricos, el valor adecuado debería oscilar entre 0.900 y 0.950, en tanto que si son investigaciones avanzadas su valor mínimo es 0.800 y solo en caso de investigaciones exploratorias se podría tolerar un valor mínimo de 0.700 (Nunnally y Bernstein, 1994).

En 1999 McDonald presenta el coeficiente Omega (ω) o índice de fiabilidad compuesta, como una medida más sensible para la evaluación de la fiabilidad de la escala, considerando que su cálculo atendía el supuesto de tal equivalencia de una escala que el alfa de Cronbach no lograba, además de que el coeficiente Omega corría un menor riesgo de sobrestimar o subestimar la fiabilidad, siendo de esta forma más preciso (Graham, 2006; Dunn et al., 2013).

El coeficiente Omega es calculado de la siguiente forma:

$$\omega = \frac{(\sum \lambda)^2}{(\sum \lambda)^2 + \sum \sigma_{\epsilon}^2}$$

Donde:

$(\sum \lambda)^2$ = suma de los coeficientes estandarizados, elevada al cuadrado

$\sum \sigma_{\epsilon}^2$ = suma total de los errores de medida de los indicadores

En este cálculo los errores de medida son iguales a la unidad menos el coeficiente estandarizado cuadrático del indicador, que no es otro que la carga factorial o peso de regresión de cada ítem. Por lo que la principal diferencia de esta medida con la presentada por Cronbach es que su cálculo no se basa en las varianzas sino en las *lambdas* o cargas factoriales estandarizadas. Los valores que el coeficiente

Omega debe alcanzar para considerarse aceptables deberán ser superiores a 0.700 (Campo-Arias y Oviedo, 2008; Hair et al., 1999).

El uso del coeficiente de Cronbach y el Omega suelen presentarse de manera conjunta en los resultados de fiabilidad de una escala. Hoy en día los *softwares* estadísticos calculan de manera automática el alfa de Cronbach, pero no el Omega, por lo que debe calcularse de manera manual o a través de aplicaciones adicionales como el *plugin* de Gaskin y Lim (2016), diseñado para AMOS y que será abordado en el capítulo siguiente.

4.5.2.2. Validez

De acuerdo con Campbell y Fiske (1959), la validez se refiere al grado en que un *test* mide el concepto que pretende medir teóricamente; esto es, que auténticamente traduzca a una medida el concepto teórico para el que ha sido construido. Un instrumento que tiene validez implícitamente trae consigo fiabilidad, por lo que debe ser evaluada con rigurosidad. La validez de una escala comprende el análisis del contenido y del constructo.

La validez de contenido se refiere a que la escala recoge todos los elementos relevantes que componen el concepto o variable que se pretende estudiar; en otras palabras, que se incluyen las dimensiones que materializan la variable teórica. Por lo general, la revisión de este tipo de validez se realiza a través de un juicio de expertos o profesionales en el tema (Ding y Hershberger, 2002). Estos determinarán si el ítem o indicador aporta al concepto que se busca medir. Debido a que puede resultar como un juicio subjetivo, Lawshe (1975) propone una medida para que pueda valorarse de manera cuantitativa dicho juicio; el inconveniente es que los consultados deben ser un número alto de panelistas o expertos, lo que puede resultar complicado al momento de su selección. Por ello, Tristán-López (2008) propone una modificación a la medida de Lawshe que permite un número menor de expertos.

Respecto a la validez de constructo, se centra en verificar que el constructo o variable no observable teórica se haya generado a partir de los indicadores que efectivamente representan dicho constructo. Cronbach (1984) señala que es la parte más importante de la validación de una escala y que permite explicar y comprender el concepto teórico. La validez de constructo trae consigo a su vez dos tipos de validez, la convergente y la discriminante.

Validez convergente

Un constructo latente tiene validez convergente cuando los indicadores que lo componen se encuentran fuertemente correlacionados (Anderson y Gerbing,

1988). La forma de comprobar que un constructo cuenta con validez convergente es a través de la revisión de su Varianza Media Extraída (*Average Variance Extracted* -AVE-), medida propuesta por Fornell y Lacker (1981) que intenta identificar qué proporción de la varianza de los indicadores es explicada o atribuida al constructo. Para su cálculo se utiliza la siguiente expresión.

$$AVE = \frac{(\sum \lambda)^2}{(\sum \lambda^2) + \sum \sigma_\varepsilon^2}$$

Donde

$(\sum \lambda)^2$ = Suma de los coeficientes estandarizados elevadas al cuadrado

$\sum \sigma_\varepsilon^2$ = suma total de los errores de medida de los indicadores

Esta medida es muy similar al coeficiente de fiabilidad compuesto Omega, con la única diferencia de que en este caso las cargas factoriales o coeficientes se elevan al cuadrado y después se obtiene una suma total. Los valores aceptables del AVE para garantizar la validez convergente, de acuerdo con Hair (1999), deberán ser superiores a 0.500. Una característica del AVE, a diferencia del alfa de Cronbach, es que entre menor número de indicadores, mayor será el valor de su varianza media extraída.

Validez discriminante

Se refiere a la capacidad de unicidad del constructo, por lo que implica que no se correlaciona con otros constructos dentro de un mismo instrumento (Campbell y Fiske, 1959). La validez discriminante es esencial para garantizar que los parámetros estimados no se encuentran sesgados. La carencia de esta validez implicaría que las relaciones planteadas en el modelo son inadecuadas (Henseler et al., 2015), por lo que en cualquier publicación los resultados de estas pruebas deben ser revelados. Para la verificación de la validez discriminante la prueba más extendida es la del Criterio de Fornell y Lacker (1981), y recientemente se empieza a utilizar la prueba de HTMT.

-Criterio de Fornell-Larcker

Este análisis permite garantizar que los constructos son diferentes entre sí, y para ello se compara el AVE con las correlaciones al cuadrado del interconstructo, que deben ser menores al AVE de la variable latente analizada; de esta manera se puede garantizar la exclusividad de los ítems a cada constructo (Fornell y Larcker, 1981). La validez discriminante es calculada en el *software* de SEM por covarianza a través de *plugins* que facilitan además su interpretación (Gaskin y Lim, 2016).

-Heterotrait-monotrait ratio of correlations (HTMT)

Es una medida novedosa y alternativa propuesta por Henseler et al. (2015) con la finalidad de analizar la validez discriminante. Parte de la estimación de las correlaciones entre los constructos de un modelo. El HTMT usa la matriz de múltiples métodos y múltiples rasgos, más conocida como la MTMM, por sus siglas en inglés, y presentada por Campbell y Fiske (1959), aunque, por su complejidad de cálculo, resulta poco usada en pruebas de validez discriminante.

La HTMT fue planteada en un primer momento para los modelos de varianzas (PLS), pero por su forma de obtención puede aplicarse a los modelos por covarianzas. Al respecto, algunos paquetes informáticos que calculan SEM por covarianzas ya cuentan en sus últimas versiones con un *plugin* que permite obtener este indicador (Gaskin y James, 2019).

La lógica de este indicador descansa en el uso de dos matrices de correlación, la *Heterotrait correlations*, que calcula las correlaciones de los indicadores de un constructo con los indicadores de otro constructo, y la *Monotrait correlations*, que calcula los valores promedio de las correlaciones entre los indicadores de un constructo (Henseler et al., 2015). Donde el cociente entre estas matrices genera el HTMT, logrando un valor entre 0 y 1 propio de una correlación.

Se considerará que se cumple con el criterio de validez discriminante del constructo cuando los valores obtenidos en la matriz de correlación resultante sean menores de 0.850 (Clark y Watson, 1995) y de 0.900 (Gold et al., 2001).

4.5.3. Paso 3. Valoración del modelo estructural

Una vez valorado el modelo de medida a través del análisis factorial confirmatorio y la revisión de la fiabilidad y validez de los constructos latentes, se procede a la representación del modelo estructural para su valoración. El modelo estructural es el modelo gráfico que muestra las relaciones causales y las covarianzas que existen entre un conjunto de variables.

4.5.3.1 Identificación del modelo

Sin embargo, antes de su estimación, debe verificarse la identificación del modelo estructural, es decir, su capacidad de encontrar una solución a las ecuaciones planteadas. Se pueden presentar dos tipos de modelos, los *sobreidentificados*, que tienen mayor número de grados de libertad que el número de parámetros a estimar, y los modelos *infraestimados*, donde los grados de libertad son menores al número de parámetros (Rigdon, 1995; Hair et al., 1999). Para que un modelo pueda tener solución deberá ser un modelo sobreidentificado.

Los grados de libertad de un modelo se determinan con la siguiente función:

$$gl = \frac{(variables\ observadas) * (variables\ observadas + 1)}{2} - \text{párametros a estimar}$$

4.5.3.2 Estimación de parámetros del modelo

Ya que se ha concluido que se cuenta con un modelo sobreidentificado, debe estimarse el modelo estructural, que establece las relaciones entre las variables latentes además de las observables que fueron asignados a cada constructo. Estas relaciones reflejan hipótesis principales basadas en la revisión de la literatura (Gefen et al., 2000). De esta manera, este modelo sugiere una sucesión de ecuaciones estructurales muy parecido a un conjunto de ecuaciones de regresión (Kline, 2011). Una vez establecido el modelo estructural, se debe realizar el proceso de estimación de parámetros y de sus respectivos errores de medición; para el cálculo de los estimadores se emplea el método más conveniente, según lo explicado en el apartado 4.2.

Como se ha mencionado a lo largo de este capítulo, un modelo estructural permite identificar aquellas relaciones causales que alcanzaron apoyo empírico para ser consideradas como hipótesis válidas, lo cual se logra a través del cálculo de los parámetros de estimación y sus respectivos errores estándar (SE). El cociente del parámetro estimado y su error estándar permite obtener el valor del estadístico *t de student*, con el cual se comprueba si el parámetro resultante de la relación entre determinadas variables latentes es estadísticamente significativo y de esta forma contrastar la hipótesis planteada teóricamente.

4.5.3.3 Bondad de ajuste

Como es común en la evaluación de relaciones causales, no basta con estimar los parámetros y verificar su estadístico *t de student* para determinar significatividad, también es necesario analizar la bondad de ajuste del modelo o la consistencia de sus resultados, a fin de conocer si los parámetros estimados efectivamente representan a la población estudiada; para tal revisión existen diferentes indicadores, pero no hay un consenso en cuáles medidas son las que deben utilizarse (Bentler, 1990; Bollen, 1990). No obstante, los diversos indicadores de ajuste son agrupados en tres tipos de medidas: de ajuste absoluto, de ajuste incremental y de ajuste de parsimonia.

Medidas de ajuste absoluto

El ajuste absoluto determina la magnitud en que el modelo propuesto predice la matriz de covarianzas. Su medida por excelencia es la chi cuadrada (χ^2), al medir el grado de discrepancia entre la matriz de covarianzas y la muestra, tomando en

cuenta su tamaño (Hu y Bentler, 1999). En esta prueba se buscan valores de chi cuadrada bajos, lo cual indicaría una ausencia de significación estadística (valores de significación menores a 0.05), que implicaría que el modelo estructural propuesto se ajusta a la matriz de covarianzas observadas (Hair et al., 1999).

Otra de las medidas de ajuste absoluto utilizada es el índice de bondad de ajuste (GFI), propuesto por Jöreskog y Sörbom (1982). Este índice representa la proporción de la varianza explicada por el modelo, y puede tomar valores en un continuo de 0 a 1; de acuerdo con Hu y Bentler (1999), considerados los autores referentes clave en medidas de ajuste, los valores aceptables para el GFI deben ser superiores a 0.950.

Por último, dos medidas adicionales de ajuste absoluto se relacionan con el error del modelo: el Residuo cuadrático medio (RMSR) y el Error de aproximación cuadrática media (RMSEA). El RMSR representa la media de los residuos de la matriz de covarianza en términos de la muestra; como todo residuo, su valor debe ser próximo a 0 (cero) y se recomienda como valor de corte aceptable un 0.060, siempre y cuando se trate de una muestra superior a 500 observaciones; para muestras menores se recomienda un RMSR menor de 0.500 (Hu y Bentler, 1999). En cuanto al Error de aproximación cuadrática media, al igual que el RMSR, mide las discrepancias entre las correlaciones, pero a nivel poblacional (Hair et al., 1999); los valores críticos aceptables deberán ser inferiores al 0.060 (Hu y Bentler, 1999).

Medidas de ajuste incremental

Por otra parte, el ajuste incremental o coeficientes relativos (McDonald y Ho, 2002) son índices que ayudan a interpretar la chi cuadrada debido a que son muy dependientes del tamaño de la muestra (Miles y Shevlin, 2006), por lo que se considera que se deben emplear indicadores de ajuste que soporten estos valores. Los indicadores que se emplean en este ajuste son el Índice de ajuste normado (NFI), el de Ajuste comparativo (CFI) y el de Tucker Lewis (TLI).

El Índice de ajuste normado (NFI) es un valor obtenido al comparar dos modelos: el estimado y el nulo; además, es un índice que se ve afectado con muestras pequeñas. Los valores pueden estar entre cero y uno, pero el umbral adecuado es de 0.90 (Hu y Bentler, 1999; Hair et al., 1999). Se expresa de la siguiente forma:

$$NFI = \frac{(\chi^2_{nulo} - \chi^2_{propuesto})}{\chi^2_{nulo}}$$

El Índice de ajuste comparativo (CFI), al igual que el de Ajuste normado, emplea dos modelos, el estimado y el nulo, pero se basa en la comparación entre ambos

con base en la chi cuadrada. Los valores aceptables deben superar el 0.95 (Hu y Bentler, 1999).

El Índice de ajuste no normado (NNFI) o Índice de Tucker Lewis (TLI) se comporta mejor con muestras pequeñas; la desventaja es que, por no estar normado, los valores pueden superar el 1, lo que dificulta su interpretación. Los valores aceptables deben ser superiores a 0.9 (Hu y Bentler, 1999; McDonald y Ho, 2002). Su expresión matemática es la siguiente:

$$TLI = \frac{(x_{nulo}^2/gl_{nulo}) - (x_{propuesto}^2/gl_{propuesto})}{(x_{nulo}^2/gl_{nulo}) - 1}$$

Medidas de parsimonia

Debido a que las medidas de ajuste de tipo absoluto e incremental tienden a obtener un mejor resultado de ajuste cuando se trata de modelos simples con pocos parámetros a estimar (Mulaik et al., 1989), se recurre a las medidas de parsimonia, que corrigen considerando los grados de libertad del modelo como un factor de ponderación. Las principales medidas de parsimonia utilizadas son el Índice de ajuste normado de parsimonia y la chi cuadrada normada.

El Índice de ajuste normado de parsimonia (PNFI) está basado en el Índice de ajuste normal (NFI), el cual es multiplicado por el cociente de los grados de libertad para el modelo propuesto, y los grados de libertad para el modelo nulo. Sus valores oscilan entre cero y uno, aunque su umbral crítico debe ser próximo a la unidad (Hu y Bentler, 1999). Se representa de la siguiente forma:

$$PNFI = \frac{gl_{propuesto}}{gl_{nulo}} * NFI$$

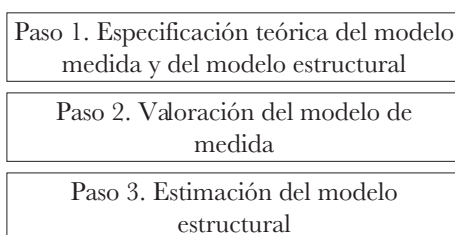
Finalmente, la chi cuadrada normada es la corrección de la chi cuadrada que se logra dividiendo su valor entre los grados de libertad a estimar en el modelo. Este índice de ajuste permite identificar modelos sobreajustados, cuando su valor es menor a 1, e infraajustados, cuando es superior a 3. Entonces, los valores críticos para esta medida se ubican en el rango entre 1 y 3 (Byrne et al., 1989). La falta de ajuste en el modelo no permitiría la aceptación global del mismo, por lo que requeriría su reespecificación. Este es el último punto del proceso de un modelo estructural SEM por covarianzas.

Capítulo 5. Aplicación práctica de ecuaciones estructurales con base en la covarianza

5.1. Introducción

El propósito de este capítulo es explicar la creación y valoración de un modelo estructural basado en la covarianza (CB-SEM) mediante un *software* especializado. Los temas a tratar incluyen la creación de un diagrama CB-SEM que permita explicar el proceso de estimación del modelo de medida, para posteriormente enfocarse en el procedimiento de evaluación del modelo estructural; todo ello mediante un ejemplo ilustrativo. La Figura 5.1 presenta los pasos que se deben seguir.

Figura 5.1. Pasos básicos a seguir para estimar un modelo CB-SEM



Fuente: Elaboración propia.

El primer paso consiste en que, con base en la teoría relacionada con el tema abordado, la especificación del modelo de medida indicará la forma en que se van a medir o calcular los constructos, por ejemplo, si se medirán con múltiples ítems o solo uno (Hair et al., 1999), mientras que la especificación del modelo estructural determinará el porqué de la utilización de determinadas variables, la defensa de su orden, su agrupación o relación propuesta entre ellas.

5.2. Construcción de un diagrama de secuencias

El presente apartado aborda la forma básica de crear un diagrama CB-SEM utilizando el *software* Amos versión 24, que ofrece una interfaz gráfica amigable e intuitiva. El ejemplo ilustrativo es la adaptación de un modelo compuesto por variables latentes que examinan el proceso de generación de conocimiento a través de la capacidad de absorción y su efecto en la innovación de las empresas del sector industrial llevada a cabo por los investigadores Vázquez et al. (2019). Se hipotetizará un nomograma que permita valorar las relaciones causales propuestas por estos autores. Para empezar, debemos asegurarnos de que la aplicación *Amos Graphics* se encuentre instalada en la computadora. Para ejecutarla, se debe localizar el archivo y dar doble clic en el ícono, lo que desplegará el programa; el resultado se puede ver en la Figura 5.2.

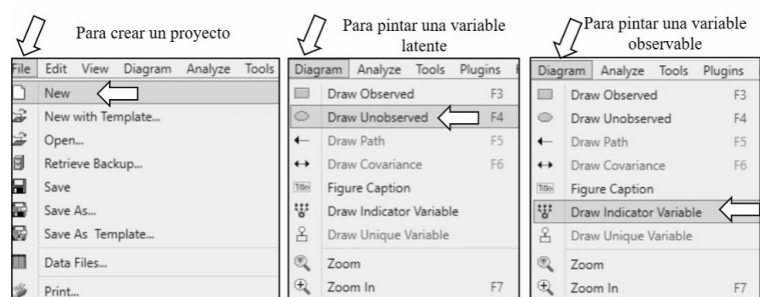
Figura 5.2. Entorno gráfico del software AMOS versión 24



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Situados en el entorno gráfico del paquete, el proceso para *crear un diagrama* es: dar clic en el botón **File**, lo que presentará un menú emergente, en él, elegir *New*, y quedará a la vista el área de modelización para pintar objetos. Una vez ahí, para pintar *una variable latente* en un diagrama la secuencia es: seleccionar *Diagram* y después, en el menú emergente, *Draw unobserved* (pintar variable latente). Enseguida, llevar el puntero al área de diseño; al seleccionarla aparecerá un círculo, el cual representa a la variable latente. (Figura 5.3.). Los pasos para pintar una variable latente se repiten según la necesidad del modelo (recuerde que se debe tener un boceto previo del mismo). Para *desactivar* la opción *Draw unobserved* solo tiene que elegir el mismo comando o cualquier otro ícono u opción del entorno gráfico del *software*.

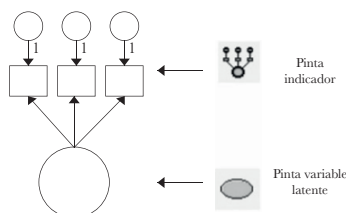
Figura 5.3. Pasos para crear un nuevo proyecto y pintar una variable latente.



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Si se desea *añadir indicadores* (variables observables) a una variable latente previamente creada, desde la opción *Diagram* del menú principal seleccionar *Draw indicator variable* (ver Figura 5.3). Al posicionarse con el ratón y dar clic sobre la variable se pinta un rectángulo (indicador) y un círculo (término de error) unidos mediante flechas. Este proceso se debe realizar para cada una de las variables latentes pintadas; el resultado se puede observar en la Figura 5.4. Para *desactivar* la opción *Draw indicator variable* solo tiene que elegir el mismo comando o cualquier otro ícono u opción del entorno gráfico del *software*.

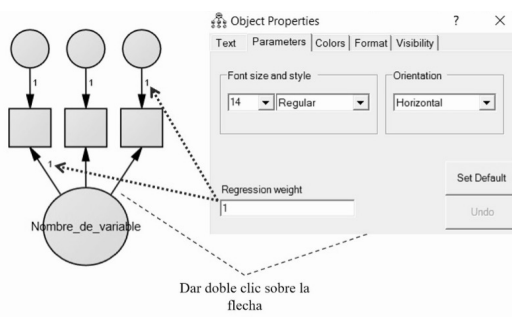
Figura 5.4. Resultado de pintar una variable latente y sus indicadores.



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Para restringir un parámetro, es decir, fijar el coeficiente de trayectoria en la unidad (1), el proceso es: en el área de dibujo o de modelización hacer doble clic sobre la flecha deseada; al hacerlo, aparece la ventana emergente *Object Properties*, que contiene varias pestañas (p. ej. *Text*, *Parameters*, *colors*, otras). Se da clic en la opción *Parameters* y en el cuadro de texto *Regression weight* se escribe 1. En el caso de que un conjunto de flechas apunten desde la variable latente a un grupo de indicadores, solo una de ellas deberá contener el valor configurado, ya que si ninguna o más de una tienen el valor 1, el *software* marcará error al intentar ejecutar la estimación del modelo. La Figura 5.5 ilustra lo comentado.

Figura 5.5. Pasos para fijar el coeficiente de trayectoria en la unidad

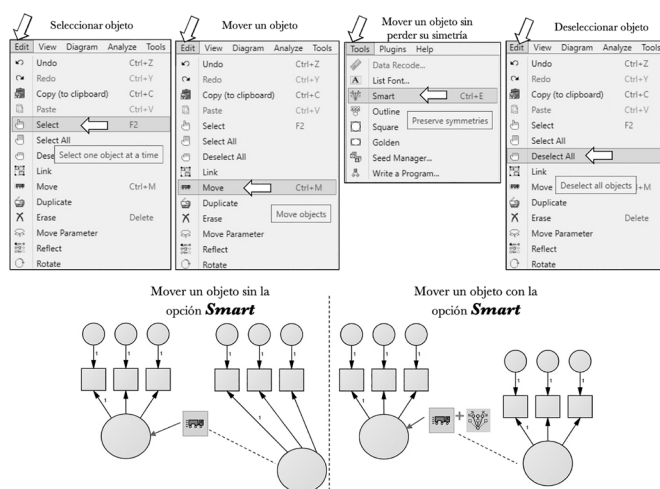


Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Si se desea *acomodar las variables o indicadores pintados* en el área de modelización, primero hay que seleccionar el objeto deseado y después moverlo. Para ello, se escoge la opción *Edit* del menú principal; en el submenú se elige *Select* (selecciona solo un objeto), como se muestra en la Figura 5.6. En el área de diseño se señala una variable latente o indicador (objeto) previamente creado (su contorno cambiará a color azul); posteriormente, ir de nuevo al menú principal, seleccionar *Edit* y luego *Move*. Enseguida, desplazarse al área de diseño, dar clic sobre el objeto previamente elegido y arrastrarlo a la posición deseada (repetir esta acción según la conveniencia del investigador). Para *desactivar Move* solo se tiene que usar el mismo comando o cualquier otro ícono u opción del entorno gráfico del *software*.

Si se quiere *mover la variable con todos sus indicadores*, se sigue este proceso: en el menú principal seleccionar *Edit*, posteriormente *Select* y escoger la variable latente a mover del área de modelización; a continuación, ir a menú principal, en la opción *Tools* escoger *Smart*, ir de nuevo a *Edit* y elegir *Move*. Ahora, desplazarse al área de diseño sobre la variable elegida, dar clic sobre ella con el botón izquierdo del ratón y sin dejarlo de presionar arrastrarla al lugar deseado; el resultado se observa en la parte baja de la Figura 5.6. Para *quitar la selección de uno o varios objetos*, ir al menú principal, opción *Edit* y luego dar clic en *Deselect All*; esto cambiará el (los) contorno(s) azul(es) del (los) objeto(s) a negro. Para *desactivar la opción Smart o deselect all* solo se tiene que elegir el mismo comando o cualquier otro ícono u opción del entorno gráfico del *software*.

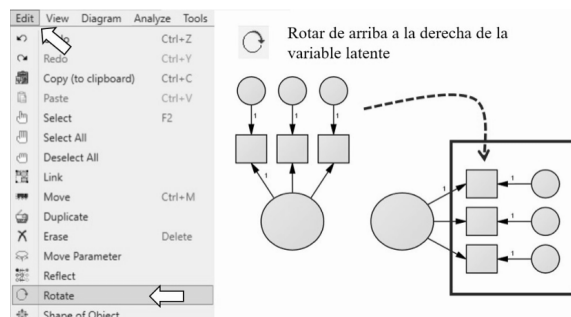
Figura 5.6. Pasos para seleccionar y mover un objeto en AMOS versión 24



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Si lo que se busca es *cambiar la orientación* de los indicadores de la variable latente de izquierda a derecha o de arriba hacia abajo para ajustar el acomodo de los objetos dentro del área de dibujo, los pasos a seguir son: en el menú principal elegir *Edit*, en el submenú dar clic en *Rotate*. Posteriormente, ir al área de dibujo/modelización y dar clic en la variable latente requerida y arrastrarla hasta que se ubique en la orientación deseada. La Figura 5.7 muestra el resultado de esta acción.

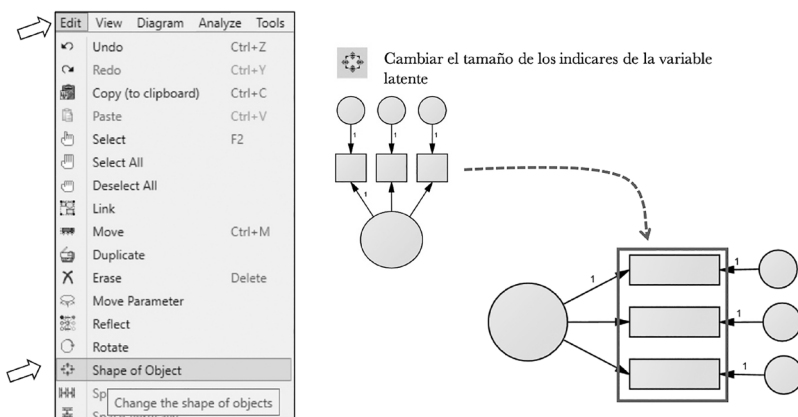
Figura 5.7. Pasos para cambiar la orientación de los indicadores de una variable latente



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Para *cambiar el tamaño o forma* de los objetos a fin de mejorar su aspecto en el diagrama, el proceso es: en el menú principal seleccionar *Edit*, después *Shape of Object*, ir al área de modelización, dar clic en el objeto seleccionado, sin soltar el botón izquierdo ampliar o reducir su tamaño, como se aprecia en la Figura 5.8.

Figura 5.8. Pasos Para cambiar la forma o tamaño de un objeto



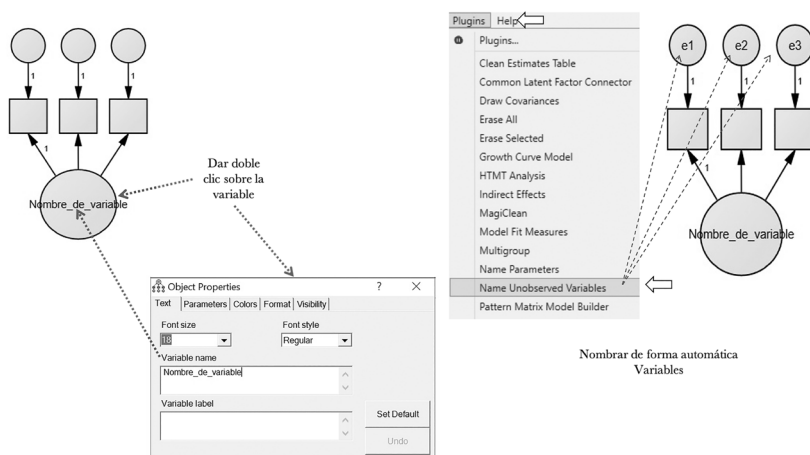
Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Antes de *cambiar el tamaño o forma* es conveniente verificar que otros objetos no estén seleccionados (esto se podrá visualizar si su contorno está en color azul); en caso de ser así, en el menú principal ir a Edit, dar clic en *Deselect All*, lo cual cambiará el (los) contorno(s) azul(es) del (los) objeto(s) a negro. Para desactivar cualquiera de las opciones mencionadas, dar clic en las mismas o en otra opción del entorno gráfico. Si lo que se necesita es *eliminar un objeto*, entonces, ir al menú principal, señalar *Edit*, luego *Delete*, ir al área de diseño, situarse sobre el objeto a eliminar y dar clic sobre él.

La secuencia *para nombrar una variable latente* -requisito indispensable para ejecutar más adelante los cálculos en un diagrama dentro del *software* Amos- es: dar doble clic sobre una de las variables latentes, aparecerá una ventana emergente con varias pestañas (como se ilustra en la Figura 5.9), las cuales habilitan una serie de opciones a realizar. Por ahora, nuestro interés se enfoca en la pestaña **Text**, que permite, entre otras acciones, cambiar el tamaño de la fuente y su tipo, pero, sobre todo, asignar a la variable un nombre y una etiqueta. Este proceso se debe realizar en cada una de las variables pintadas en su modelo.

Si se desea nombrar las variables pintadas o designadas en el diagrama para medir los términos de error (que a veces es una tarea pesada, por la gran cantidad de éstos), el *software* permite hacerlo de forma automática, para lo cual hay que localizar en el menú principal la opción *Plugins* y dar clic en ella, luego, en el comando *Name Unobserved Variables*. El resultado será que todas las variables tendrán ahora un nombre, pero si el requerimiento es modificarlo, el proceso es el mismo que para nombrarlas. La acción se puede ver en la Figura 5.9.

Figura 5.9. Proceso para nombrar una variable latente y sus indicadores

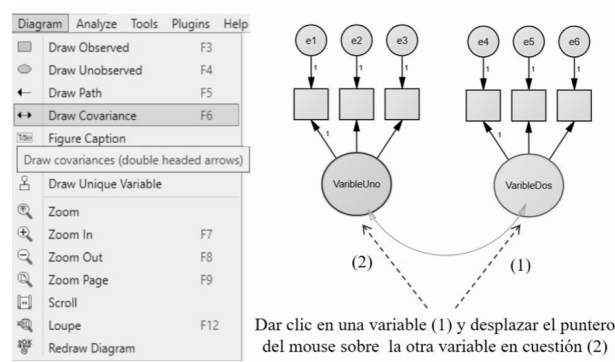


Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Sobre cómo *pintar las relaciones* entre las variables latentes creadas, en principio, la opción a elegir depende del tipo de relación; si se busca correlacionar dos o más variables (sin error de medida), éstas deberán ser unidas por una flecha con puntas en ambos extremos (*Draw covariances*), pero si se pretende relacionar dos variables (una exógena y otra endógena), la flecha será con una sola punta, que va pintada del constructo independiente o exógeno al constructo dependiente o endógeno. Veamos cada caso.

Para pintar una correlación, desde el menú principal dar clic en *Diagram*, luego en *Draw Covariance* (Figura 5.10); trasladarse enseguida al área de diseño del diagrama, para posicionarse, y dar clic con el botón izquierdo del ratón en una de las variables a relacionar; sin soltarlo, desplazar el puntero hasta estar encima de la otra variable en cuestión, para finalmente soltarlo. Recordar que ninguna de las dos variables latentes deberá tener pintado su término de error; en caso contrario, el *software* marcará equivocación al tratar de realizar la acción.

Figura 5.10. Proceso para pintar una relación de covarianza



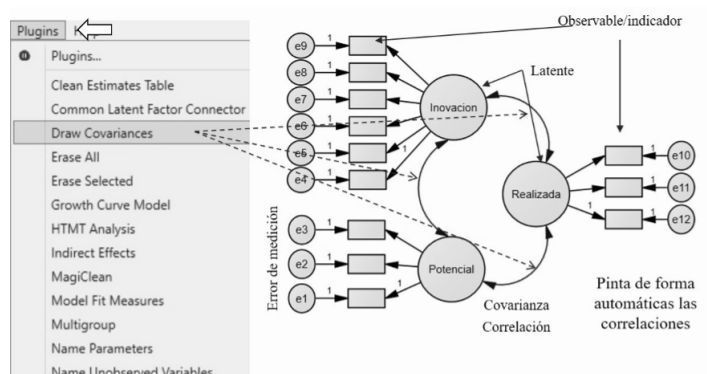
Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Si son muchas variables por correlacionar (este proceso no aplica para el pintado de rutas *path*), el *software* ofrece una acción que lo efectúa de forma automática. Para ello, en el menú principal hacer clic en *Edit*, luego en *select*; posteriormente ir al apartado de modelización y seleccionar todas aquellas variables endógenas que estarán involucradas, dando un clic sobre ellas (su contorno se pone azul); ir de nuevo al menú principal y dar clic en *Plugins*, luego en *Draw Covariances*, como se muestra en la Figura 5.11.

Si el proceso de pintar las correlaciones entre variables se realizó correctamente, el resultado será similar al que se ve en la Figura 5.11, que también presenta el

diagrama de causalidad completo que se usará como ejemplo ilustrativo para la parte de evaluación del modelo de medida. Es relevante mencionar que se debe crear un proyecto independiente para analizar y estimar un diagrama causal enfocado a la evaluación del modelo de medida y otro para valorar el modelo estructural.

Figura 5.11. Proceso para pintar varias relaciones de covarianza de forma automática



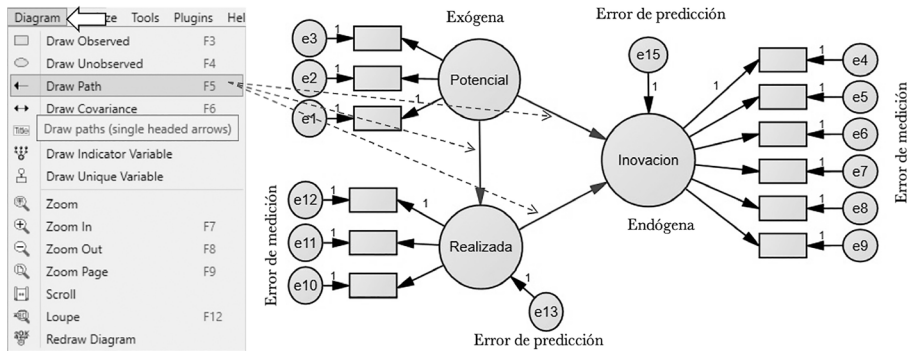
Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Para pintar una relación en un modelo causal como el que se muestra en la Figura 5.12, lo primero a considerar es que un diagrama causal o *path* puede contener ambos tipos de relaciones, de covarianza-correlaciones y de rutas, así como una variable endógena puede tener apuntada más de una flecha directa hacia ella, lo que ocasionará tener que pintarle un objeto que contenga el denominado error de predicción. Cabe señalar que lo descrito en los siguientes apartados para este punto no limita diseños más elaborados.

Retomando el proceso de pintar una relación *path* en un diagrama: localizar en el entorno gráfico de Amos la opción Diagrama y dar clic en ella, luego en el comando *Draw Path* del menú emergente; enseguida, trasladarse al área de diseño-modelización del diagrama y dar clic en una de las variables a relacionar (independiente o exógena) y sin soltar el botón izquierdo moverse hasta estar encima de la variable dependiente o endógena (soltar el botón del ratón).

Como puede verse, las relaciones siempre se dibujan de la variable independiente hacia la dependiente, y los pasos se deberán repetir tantas veces como relaciones o hipótesis se tengan en el modelo de investigación propuesto. En la Figura 5.12 se muestra el diagrama causal completo, el cual forma parte del ejemplo ilustrativo que se aplicará para explicar la valoración de un modelo estructural, que se abordará en el punto 5.3 de este mismo capítulo.

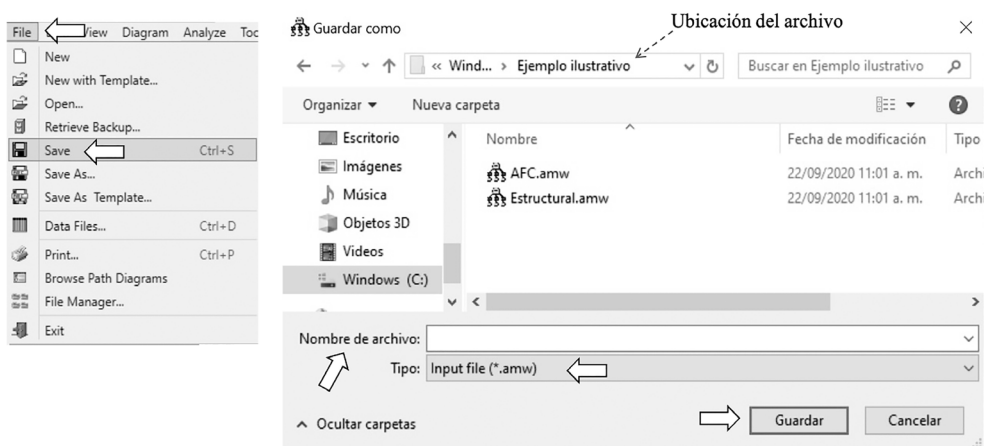
Figura 5.12. Pasos para pintar rutas path entre variables latentes



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Es importante señalar que se deben ir *almacenando las modificaciones o actualizaciones* que se hagan al archivo de diagrama, de lo contrario, se corre el riesgo de perder todo lo realizado. Para ello hay que ir a la opción *File* del menú principal y seleccionar *Save*. Si es la primera vez que se guarda se mostrará una pantalla para determinar la ubicación y el nombre del archivo (ver Figura 5.13). Ya almacenado, en lo sucesivo *Save* sólo lo actualizará con los cambios realizados. En el *software* Amos la extensión de los archivos de diagrama guardados es *amv*.

Figura 5.13. Proceso para guardar un diagrama en AMOS



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Ahora que ya se sabe cómo pintar y guardar un diagrama causal en Amos, lo que sigue es cargar el archivo de datos con las observaciones al proyecto; para ello, el procedimiento es: en el menú principal dar clic en *File*, luego en *Data Files*; esta acción abrirá una ventana que permite visualizar los archivos de observaciones con que cuenta el proyecto; en caso de ser la primera vez, hay que dar clic en el botón *File Name*, el cual abrirá una ventana que permite buscar el archivo de observaciones en la computadora; una vez localizado, hay que dar clic en el botón *Abrir* (ver Figura 5.14). Realizada esta acción, la ventana *Data File* indicará el nombre del archivo adjuntado. Cabe mencionar que el *software* Amos permite abrir una gran variedad de tipos de archivos, pero por *default* solicita aquellos con extensión **.sav**, es decir, del paquete informático SPSS.

Figura 5.14. Proceso para cargar-adjuntar archivo de observaciones al proyecto



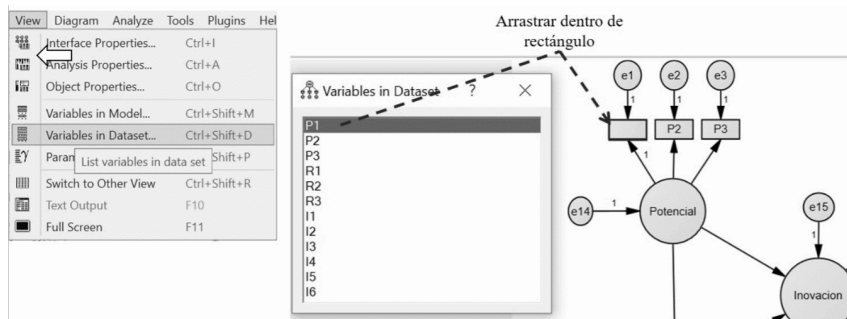
Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Una vez cargado el archivo de observaciones, el siguiente paso es *asignar las variables de éste a los indicadores* de una variable latente pintada en el diagrama causal (aplica tanto para el diseño de un modelo de medida como para uno estructural). Para ello hay que hacer clic en *View* del menú principal y después en *Variables in Data Set*. Esta acción despliega una ventana emergente con el nombre de las variables incluidas en el archivo de observaciones.

Posteriormente se tiene que seleccionar una de las variables mostradas dando clic en ella; sin soltar el botón izquierdo del ratón, desplazarse al área de modelización,

para posicionarse sobre el rectángulo del indicador correspondiente, y soltarlo. Para comprobar que esta acción se realizó con éxito, el indicador tomará el nombre de la variable que se seleccionó de la ventana *Variables in Data Set*. Estos pasos se tienen que repetir para cada indicador pintado en el diagrama. La Figura 5.15 indica lo señalado.

Figura 5.15. Proceso para asignar observaciones a los indicadores de una variable latente



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

El proceso de carga del archivo de observaciones y la asignación de éstas a los indicadores correspondientes se deberá realizar tanto para el proyecto que analiza el modelo de medida como para el que estima el modelo estructural. Esto quiere decir que el proceso se repetirá al menos dos veces, uno para el diagrama de covarianza y otro para el diagrama causal. Por último, todas las opciones abordadas se pueden realizar desde el menú de íconos que se ubica en la parte izquierda de la ventana de diseño; para saber qué acción realiza cada uno, solo hay que posicionar el puntero sobre ellos (ver Figura 5.2).

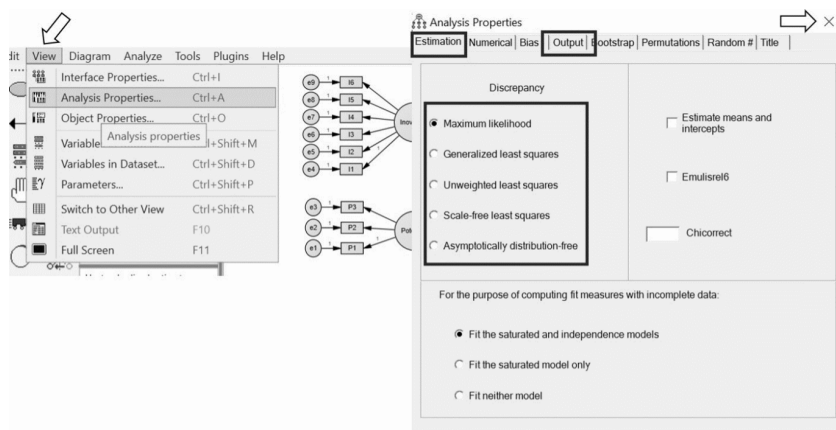
5.3. Configuración de parámetros para valoración de un modelo en el software AMOS

Antes de valorar la identificación de un modelo se deben establecer los parámetros de estimación en el *software* Amos. Este procedimiento permite ofrecer valores relacionados con errores estándar, pesos estandarizados y no estandarizados, estadísticos *t*, la significancia o valores *p*, estadísticos de ajuste (absolutos, comparativos, de parsimonia), coeficientes de determinación (R^2), pruebas de normalidad y valores atípicos, entre otros más, con los cuales el investigador puede tomar decisiones respecto a si se debe o no reespecificar el modelo propuesto.

Para configurar la forma en que se van a estimar los parámetros en el *software* Amos hay que localizar en el menú principal la opción *View*, de la cual se despliega

un menú; se da clic en *Analysis Properties*, lo que abrirá una ventana con varias pestañas, como se muestra en la Figura 5.16. Por el momento nos enfocaremos en *Estimation* y en *Output*.

Figura 5.16. Selección del método de estimación



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

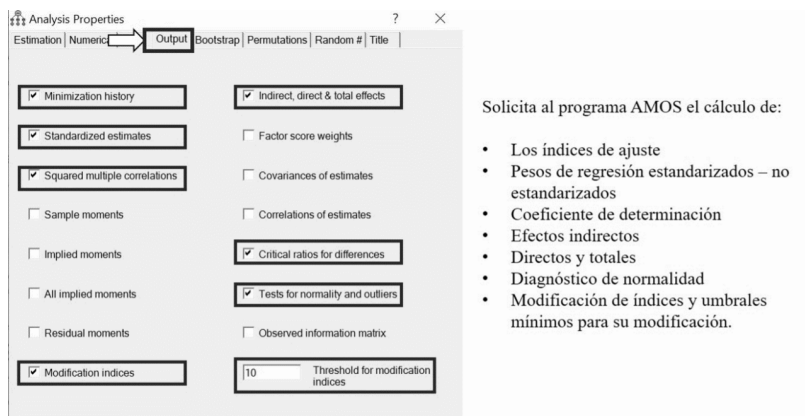
En la pestaña *Estimation* (Figura 5.16) aparecen las diferentes técnicas o algoritmos de estimación que Amos aplica, entre ellos, Máxima verosimilitud (*Maximum likelihood*), Mínimos cuadrados generalizados (*Generalized least squares*), Mínimos cuadrados no ponderados (*Unweighted least squares*) y Asintóticamente libre de distribución (*Asymptotically distribution-free*). La selección dependerá de varios factores, aunque el más aplicado es el de Máxima verosimilitud, ya que es capaz de proveer convergencia aun presentando los datos ligeras desviaciones en su distribución normal (Bentler, 1995); no obstante, es recomendable cumplir con los supuestos de normalidad multivariante establecidos, así como contar con un adecuado número de observaciones, situación similar para el de Mínimos cuadrados generalizados.

Por su parte, el método de Mínimos cuadrados no ponderados se orienta a la estimación de variables categóricas, mientras Asintóticamente libre de distribución puede aplicarse en casos donde la falta de normalidad extrema esté presente, pero exige para ello un gran número de observaciones ($N \geq 500$). Para más detalles se puede revisar el trabajo de Brown (2006). En nuestro ejemplo ilustrativo se seleccionó el de Máxima verosimilitud, ya que las observaciones cumplen con lo establecido para el mismo.

La pestaña *Output* (Figura 5.17) muestra las diferentes opciones de salida que pueden ser solicitadas (Índices de ajuste, Pesos de regresión, Coeficiente de

determinación, Efectos indirectos, Directos y totales, Diagnóstico de normalidad, Modificación de índices y umbrales mínimos para su modificación), pero la elección queda a criterio o necesidad del investigador. Después de hacer la selección se deben cerrar las ventanas de ambas pestañas, quedando con ello configurados y listos los parámetros de estimación que se ocuparán, ya sea para el modelo de medida o para el estructural.

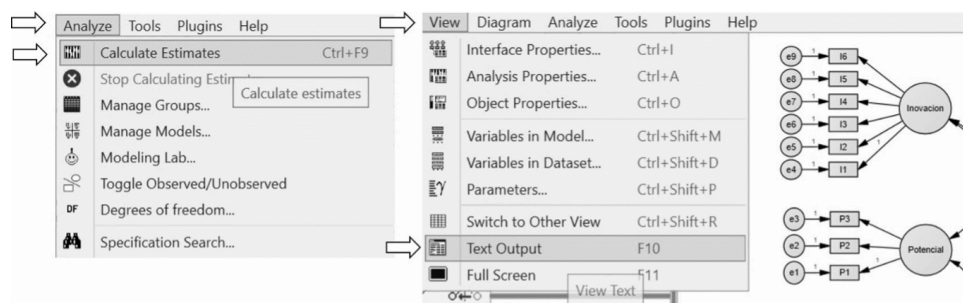
Figura 5.17. Selección de principales salidas para valorar un modelo en AMOS



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

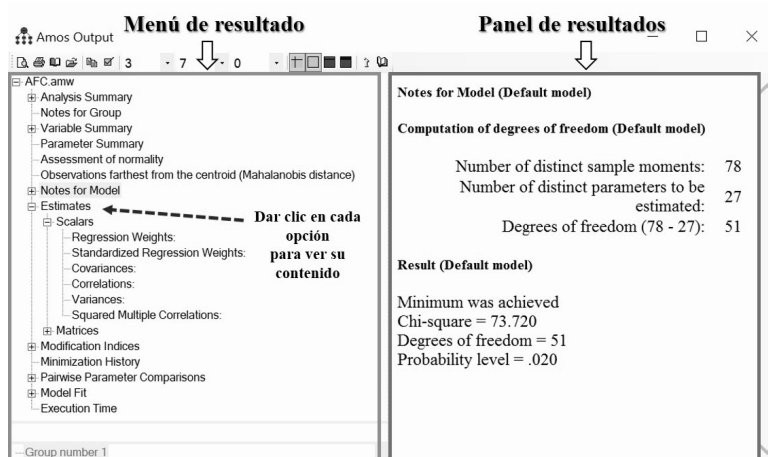
Para ejecutar lo configurado se debe dar clic en la opción *Analyze* del menú principal, luego en *Calculate Estimates*; enseguida se comienzan a correr los cálculos de las opciones seleccionadas. El alcance de este proceso se podrá visualizar con la opción *View* del menú principal (dar clic) en el *Text Output* del submenú, como se muestra en la Figura 5.18. Si todo se realizó correctamente, se desplegará una ventana como la que se observa en la Figura 5.19, en la cual, a través del menú que se muestra a la izquierda, se tendrá acceso a los diferentes resultados (panel derecho de la ventana) que presenta el *software*.

Figura 5.18. Pasos para visualizar ventana de principales resultados en AMOS



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Figura 5.19. Ventana de resultados obtenida mediante AMOS



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

5.4. Evaluación del modelo de medida

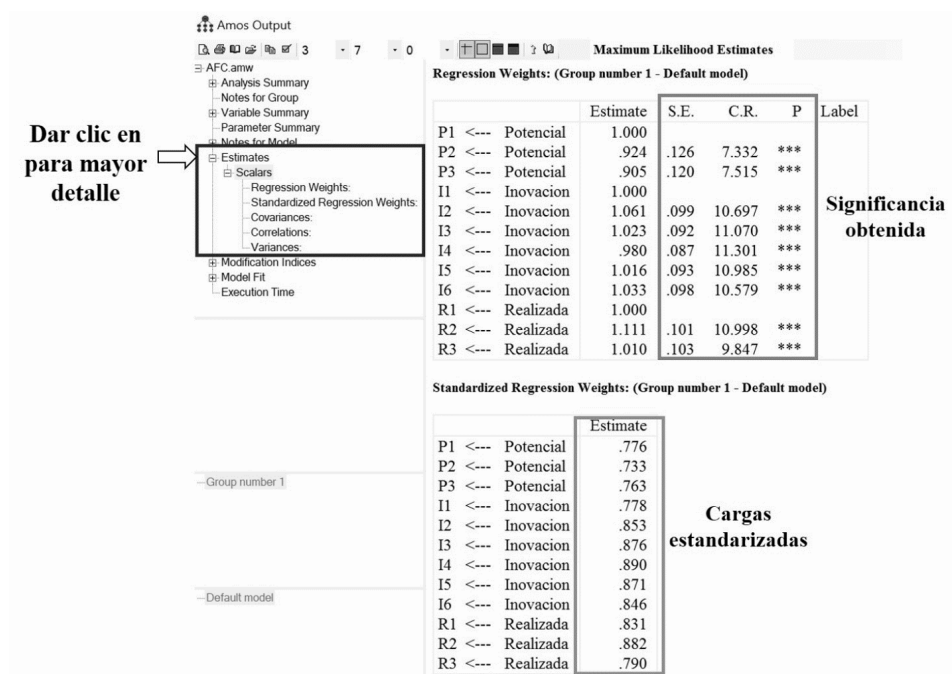
Para evaluar el modelo de medida el investigador tuvo que transitar de un análisis factorial exploratorio -en el cual se probaron las variables que posiblemente estimarían un constructo(s) o factor(es) específico(s)- a un proceso de análisis confirmatorio (CFA), donde se conocen o se controlan las variables (indicadores) que van a medirlo(s). La razón principal por la cual es conveniente realizar la valoración de un CFA es por su capacidad de evaluar si en realidad un constructo medido refleja al constructo latente teórico. Para ello, el investigador debe analizar la fiabilidad del constructo, la validez convergente (comprobar la unicidad del

conjunto de indicadores que lo conforman) y la validez discriminante (comprobar que las variables generadas no muestren relación con otras dentro del modelo).

La fiabilidad se obtiene de la valoración de las cargas factoriales o ponderaciones obtenidas por sus indicadores, así como por la consistencia interna del constructo. La validez convergente puede ser generada a partir de los valores conseguidos por el coeficiente denominado Varianza Extraída Media (AVE). La validez discriminante se evalúa mediante los criterios de Fornell y Larcker de 1981, el análisis de cargas cruzadas o la prueba Heterotrait-monotrait (HTMT).

El primer paso para la revisión de un modelo de medida será valorar las cargas factoriales estandarizadas de sus indicadores. Para ello, el investigador podrá tomar como criterio que los valores alcanzados deberán ser mayores o iguales a 0.5, considerando que lo ideal es que sean superiores a 0.7 (Hair et al., 2014); de lo contrario, se debe analizar la posibilidad de eliminar al indicador. Lo anterior se obtendrá de la sección *Estimate* que se muestra en la ventana de la Figura 5.20.

Figura 5.20. Obtenciones de valores estandarizados para las cargas de indicadores.



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Dando clic sobre esta opción se mostrarán los resultados de los diferentes estimados calculados, por ejemplo, pesos de regresión estandarizados y no estandarizados, valores de covarianza, de correlaciones, de varianzas y de correlaciones múltiples cuadradas.

No obstante, para determinar si los indicadores cumplen con los valores de las cargas factoriales estandarizadas hay que dirigirse a *Standardized Regression Weights*, donde los resultados requeridos se despliegan en forma de tabla, tal como se observa en la Figura 5.20. Para el ejemplo, se puede notar que todos los indicadores cumplen con lo establecido (cargas estandarizadas mayores a 0.707) y que además fueron significativas. Pero en caso contrario, se recomienda renunciar a ellos si esto mejora la fiabilidad compuesta del constructo (Henseler et al., 2009).

Por otra parte, para evaluar la fiabilidad del constructo se deberá aplicar el criterio de alfa de Cronbach, tomando como referencia las siguientes escalas: de 0.80, satisfactorio para investigaciones en etapas avanzadas, y de 0.70 para investigaciones en etapas iniciales (Nunnally & Bernstein, 1994), mientras que valores por debajo de 0.60 indican una falta de fiabilidad. Otra alternativa es el Índice de fiabilidad compuesta (P_c), donde valores de 0.60 a 0.70 son aceptables en investigaciones exploratorias, mientras que de 0.70 a 0.90 son óptimos (Henseler et al., 2009).

La estimación de la validez convergente se puede realizar mediante el cálculo de la varianza media extraída (AVE), donde el valor mínimo aceptado es de 0.5 (Fornell & Larcker, 1981). Mientras que la validez discriminante puede ser comprobada, en un primer término, con la comparación del AVE y las correlaciones al cuadrado de cada constructo, de forma que el AVE, al ser superior a las correlaciones al cuadrado, comprobaría la validez discriminante (Fornell & Larcker, 1981; Raykov & Penev 2010). También puede ser evaluada a través de la nueva metodología de Henseler et al. (2015) denominada *heterotrait-monotrait ratio of correlations* o HTMT.

AMOS, hasta la versión 24 (última revisada), no incluye por defecto todas las estimaciones mencionadas, por ello, la mayoría de los investigadores utilizan *softwares* independientes para realizarlas de una manera más eficiente y rápida. No obstante, en el presente libro se utilizarán nuevos complementos y estimados (*plugins*) que investigadores han desarrollado para compensar las limitaciones del AMOS, los que pueden ser descargados en <http://statwiki.kolobkreatiions.com/index.php?title=Plugins>, donde también se encuentran los tutoriales que explican su instalación para las diferentes versiones del *software* AMOS, necesario para realizar lo que a continuación se expone.

La ruta para obtener la fiabilidad compuesta, el AVE y la validez discriminante de los estimados antes mencionados es: en el menú principal dar clic en *plugins*, lo que despliega un submenú; seleccionar y dar clic en *Validity and Reliability Test*. Enseguida se abrirá una ventana del navegador, como se observa en la Figura 5.21,

con los resultados obtenidos para la fiabilidad compuesta (CR), el AVE (validez convergente) y para la valoración de la validez discriminante utilizando el criterio de Fornell y Lacker de 1981.

Figura 5.21. Ventana de resultados para la validez convergente y discriminante.

Model Validity Measures

Validez discriminante
Criterio Fornell y Larcker (1981)

	CR	AVE	MSV	MaxR(H)	Potencial	Inovacion	Realizada
Potencial	0.801	0.574	0.290	0.803	0.757		
Inovacion	0.941	0.728	0.412	0.944	0.539***	0.853	
Realizada	0.873	0.697	0.412	0.881	0.498***	0.642***	0.835

References
Significance of Correlations:
† p < 0.100
* p < 0.050
** p < 0.010
*** p < 0.001

Fuente: Elaboración propia a partir del AMOS Plugin de Gaskin y Lim (2016).

Para calcular la validez discriminante a través del nuevo método HTMT los pasos a seguir son: ir a la opción *plugins* en el menú principal y luego dar clic en *HTMT Analysis*, esto hará que el programa realice un proceso que abrirá una nueva ventana con los resultados obtenidos y mostrará sugerencias a seguir en caso de no obtener los valores deseados. En el ejemplo se puede apreciar que el modelo de medida cumple con los criterios solicitados por la validez convergente y por la validez discriminante, destacando que este último fue evaluado tanto con el criterio de Fornell y Larcker de 1981 (Figura 5.21) como por el de HTMT (Figura 5.22).

Figura 5.22. Ventana de resultados para la validez discriminante mediante el método HTMT

HTMT Analysis

	Potencial	Inovacion	Realizada
Potencial			
Inovacion	0.537		
Realizada	0.491	0.638	

Thresholds are 0.850 for strict and 0.900 for liberal discriminant validity.

-Henseler, J., C. M. Ringle, and M. Sarstedt (2015). A New Criterion for Assessing Discriminant Validity in Variance-based Structural Equation Modeling, *Journal of the Academy of Marketing Science*, 43 (1), 115–135.

-Gaskin, J. & James, M. (2019) HTMT Plugin for AMOS.

Warnings:
There are no warnings for this HTMT analysis.

Fuente: Elaboración propia a partir del AMOS Plugin de Gaskin y James (2019).

Por último, se debe comprobar que el conjunto de datos no presente sesgos originados por la aplicación de métodos únicos o comunes externos a ellos que provienen de una misma fuente (p. ej., una encuesta en línea) (Podsakoff et al., 2003); para ello se pueden realizar algunas de las siguientes pruebas: prueba de factor único de Harman (actualmente ya no es ampliamente aceptada y se considera un enfoque obsoleto e inferior), método de la latente de factor común o el método de la variable marcador. No obstante, su desarrollo y explicación mediante herramientas informáticas quedan fuera del alcance del presente libro.

5.5. Evaluación del modelo estructural

La explicación de cómo evaluar un modelo de ecuaciones estructurales con el *software* AMOS la hemos dividido en dos partes, una enfocada a la valoración de los criterios de calidad del modelo (estimación del ajuste del modelo) y otra relacionada con la interpretación de los resultados obtenidos a través de los coeficientes de determinación R^2 , valores beta y significancia de las rutas *path* propuestas.

5.5.1. Valoración de criterios de calidad

La valoración de la calidad de un modelo se realiza a través de los llamados Estadísticos de bondad de ajuste, los cuales, de acuerdo con la literatura, se agrupan en tres categorías: los que permiten estimar el ajuste absoluto, los de ajuste relativo, incremental o de comparación, y los que valoran el ajuste de parsimonia. Sin embargo, ninguna medida o categoría por sí sola puede proporcionar los datos requeridos para evaluar un modelo SEM; por ello, Hu y Bentler (1999) recomiendan combinarlas. La Tabla 5.1 muestra las categorías mencionadas, sus estadísticos más utilizados y los límites críticos sugeridos para ellos.

Tabla 5.1. Tipos de bondad de ajuste, sus estadísticos y criterios de referencia

Tipo	Estadístico	Abreviatura	Criterio
Ajuste absoluto			
	Chi-cuadrado	χ^2	Significación > 0.05
	Índice de bondad de ajuste	GFI	≥ 0.95
	Raíz del residuo cuadrático medio	RMR	Próximo a 0
	Raíz cuadrada media del error de aproximación	RMSEA	< 0.06
Ajuste comparativo/incremental			
	Índice de ajuste comparativo	CFI	≥ 0.95
	Índice de bondad de ajuste corregido	AGFI	≥ 0.90

Continúa Tabla 5.1

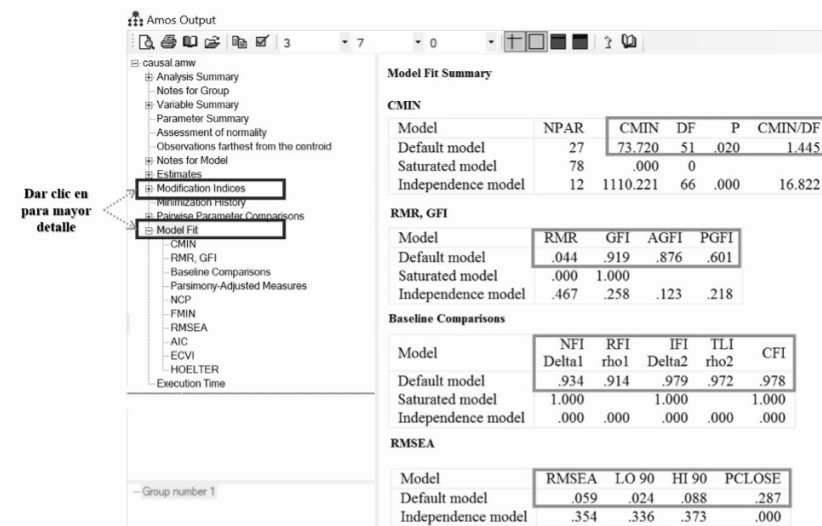
Tipo	Estadístico	Abreviatura	Criterio
	Índice de Tucker-Lewis	TLI	≥ 0.90
	Índice de ajuste normalizado	NFI	≥ 0.90
	Índice de ajuste incremental	IFI	≥ 0.90
Ajuste parsimonioso			
	Corregido por parsimonia	PNFI	Próximo a 1
	Razón Chi-cuadrado / grados de libertad	CMIN/gl	$1 < > 3$

Fuente: Elaborado a partir de Byrne (2001), Hooper, Coughlan y Mullen (2008).

Para obtenerlos del *software* AMOS previamente se debe haber configurado la forma en que se van a estimar los parámetros (ver punto 5.1.2). Luego, ubicar en el menú principal la opción *Analyze* y dar clic, después en *Calculate Estimates* del menú desplegable; enseguida se comienzan a correr o ejecutar los cálculos de las opciones seleccionadas.

El alcance del proceso se podrá conocer con la opción *View* del menú principal (dar clic) y dentro del submenú elegir *Text Output*, como lo ilustra la Figura 5.18. Si los procesos se realizaron correctamente, se desplegará una ventana similar a la que se muestra en la Figura 5.23, en la cual se tiene acceso a los diferentes resultados (panel derecho de la ventana), los cuales son muy completos y detallados. No obstante, solo nos enfocaremos en la sección *Model Fit* (ajuste del modelo); para ello, dar clic sobre ésta.

Figura 5.23. Ventana de resultados con los estadísticos básicos de ajuste



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Para su interpretación, basta con tener en cuenta los criterios especificados en la Tabla 1. De acuerdo con el ejemplo ilustrativo, se puede ver (Figura 5.22) que presenta un ajuste adecuado ($\chi^2 = 73-720$; $p = 0.20$; $CMIN/DF = 1.445$; $RMR = 0.44$; $GFI = 0.919$; $CFI = 0.978$; $RMSEA = 0.59$; $PClose = 0.287$). De no haber sido así, se tendrían que examinar los índices de modificación, puesto que dan información de los parámetros individuales y un análisis detallado de los residuos y su posible reducción. Pero se debe tener precaución, pues el investigador puede verse tentado a sacrificar valor teórico a cambio de un mejor ajuste; en caso de proceder, se señala que estas acciones deben ir en concordancia con la teoría que lo soporta.

5.5.2. Interpretación de resultados

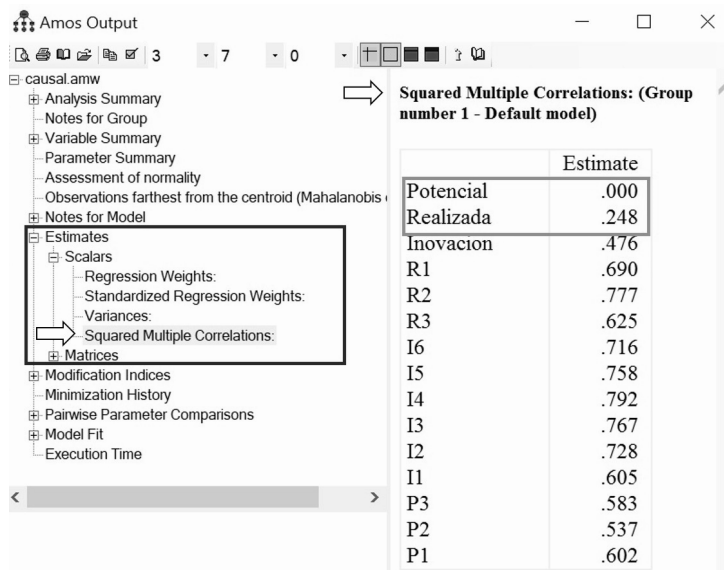
Continuando con el análisis de la valoración del modelo estructural, se debe proceder a determinar el valor explicativo del modelo. Para ello suele considerarse como indicador la proporción de varianza explicada indicada a través del coeficiente de determinación (R^2). Su cálculo se basa en los valores de las correlaciones cuadradas del valor real y del predicho del constructo dependiente (Rigdon, 2012). El criterio de aceptación de éste varía de 0 a 1, por consiguiente, entre más cerca de 1, más será su nivel de predicción. Sin embargo, para estimar un valor R^2 no existe una guía que permita determinar si es aceptable o no, ya que será el área de estudio en cuestión la que establezca los parámetros. Pero si se obtiene un poder explicativo pobre, el investigador deberá analizar la razón de éste; los valores bajos del R^2 se presentan cuando el modelo propuesto no incluyó constructos o factores explicativos de relevancia, o que los elementos evaluados en el mismo no eran los adecuados o pertinentes para el contexto aplicado.

Para obtener los valores del coeficiente de determinación (R^2) de un modelo mediante el *software* AMOS se requiere configurar la forma en que se van a estimar los parámetros (ver punto 5.1.2). Para ejecutar lo configurado, se escoge en el menú principal la opción *Analyze*, luego *Calculate Estimates* del menú que se despliega. Hecho esto, se comienzan a correr o ejecutar los cálculos de las opciones seleccionadas. El alcance de este proceso se podrá visualizar dando clic en *View* del menú principal y después en *Text Output*, como se muestra en la Figura 5.18.

Si todo se realizó correctamente se desplegará una ventana similar a la de la Figura 5.24, en la cual se mostrarán los valores. Mediante el menú desplegable se puede acceder a los diferentes resultados obtenidos (panel derecho de la ventana), muy completos y detallados, pero hay que estar atentos a la sección *Estimates*, para encontrar los valores R^2 . Para ello, dar clic sobre la misma, luego en *Scalars* y enseguida sobre *Squared multiple correlations*. Los resultados se muestran en formato de tabla e incluyen valores para las relaciones propuestas y para los indicadores, aunque, como

se puede observar, solo nos interesan los primeros. Para el caso del ejemplo ilustrativo la proporción de varianza fue moderado, dado que solo un 47% de la variabilidad del constructo innovación fue explicado por sus variables precedentes.

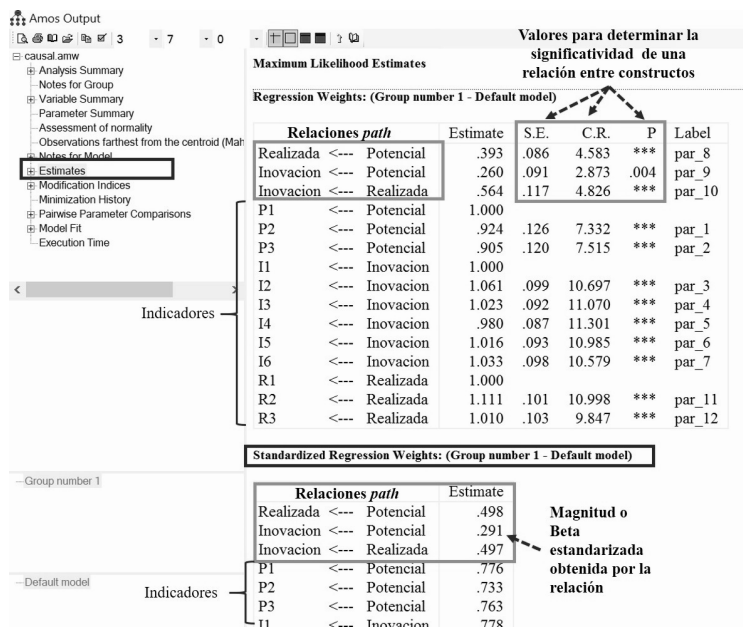
Figura 5.24. Ventana de resultados obtenidos para el coeficiente de determinación (R^2)



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

Finalmente, resta analizar la estimación de las rutas o coeficientes *path* entre las variables del modelo estructural y su significatividad, puesto que son los que se van a necesitar para testear las hipótesis planteadas. Para obtenerlos se deben seguir los mismos pasos que para el coeficiente de determinación, tal como se muestra en la Figura 5.25, pero en esta ocasión sólo hay que dar clic en *Estimates*; el resultado se puede visualizar en el panel derecho de la ventana en la sección *Standardized Regression Weights*.

Figura 5.25. Resultados de la magnitud y significación de los coeficientes path



Fuente: Elaboración propia a partir del software IBM SPSS AMOS 24.

La tabla que contiene esa sección muestra los valores beta (β) obtenidos para cada relación pintada en el modelo. Para una adecuada interpretación de los datos hay que tener presente que los coeficientes *path* son valores estandarizados que indican la fuerza de las relaciones entre las variables dependientes e independientes, por lo que pueden oscilar entre -1 y +1; entre más se acerquen a la unidad positiva, más fuertes serán las relaciones y, por el contrario, entre más se aproximen a 0, las relaciones serán más débiles (Johnson et. al, 2006).

Para determinar su nivel de significatividad se verifican los datos de la sección *Regression Weights* (Figura 5.25), prestando atención a los valores de las columnas Solución estandarizada (S.E., *solution standardized*), Valores críticos (C.R., *critical ratio*) y Significatividad (valor *P*). Para determinar si lo obtenido es estadísticamente significativo se debe comparar con un valor crítico, el cual puede ser medido a través de la prueba de una o la de dos colas. Para el caso de una cola, un valor *t* de 1.28 equivale a un nivel de significación de 10%, uno de 1.65 a 5% y de 2.33 a 1%. Para una de dos colas, un valor *t* de 1.65 equivale a un nivel de significación de 10%, uno de 1.96 a 5% y uno de 2.57 a 1%.

Como se puede observar en la Figura 5.25, las tres relaciones del ejemplo ilustrativo fueron aceptadas, puesto que en la relación entre la variables potencial y realizada el valor beta alcanzado fue de 0.498, con un nivel de significación de $p < 0.001$ (***), siendo además la más alta de las tres. Para la relación entre la variable realizada y el constructo innovación el valor obtenido fue de 0.497, con un valor $p = 0.004$, y la relación entre potencial e innovación alcanzó un coeficiente de 0.291, pero altamente significativa ($p < 0.001$).

Capítulo 6. Modelado de ecuaciones estructurales basado en la varianza

6.1. Introducción

Hasta la década de 1960 no existían los métodos de investigación con base en la inferencia estadística y herramientas de segunda generación. Sin embargo, en tiempos posteriores la tecnología de la información ha aportado avances significativos para el estudio de grandes cantidades de datos, incluso Hair et al. (2014) indican que las barreras metodológicas ya no son un obstáculo para investigadores que buscan evidencia con base en sus trabajos empíricos. En la actualidad, las técnicas de investigación del siglo XX para el análisis de datos en las ciencias sociales y administrativas están quedando en el olvido; herramientas como regresión simple, chi cuadrada, media aritmética, entre otras, solo se emplean como auxiliares, porque hacen uso de la correlación en lugar de una matriz de varianzas o covarianzas, ya que estas últimas evalúan si un ítem es confiable en su medición (Chin et al., 2003).

El análisis multivariante y los modelos causales han surgido, de acuerdo a Barclay et al. (1995), para: *i*) permitir la evaluación de las propiedades de medición de constructos en ambientes diversos, *ii*) manejar explícitamente la medición del error, *iii*) proporcionar beneficios fundamentados en esos análisis estadísticos superiores, no disponibles en las herramientas de primera generación, como son el análisis de componentes principales, análisis clúster, ANOVA, MANOVA, regresión múltiple o correlación canónica. Muchos de estos nuevos productos han favorecido principalmente a la mercadotecnia, la administración de empresas, la economía y los sistemas de información computacionales.

Con las nuevas herramientas de segunda generación, como el Modelado de Ecuaciones Estructurales (*Structural Equation Modeling* - SEM), es posible que los investigadores puedan examinar modelos complejos, que miden por medio de variables múltiples y que no solo evalúan la causalidad, sino que van más allá, en busca de la predicción, tan útil en el mundo organizacional globalizado, con el apoyo de las tecnologías de la información.

SEM es una de esas técnicas multivariantes que surge de la combinación de dos enfoques: la línea econométrica, que se orienta a la predicción, y la perspectiva psicométrica, que modela constructos latentes (no observados) que son indirectamente inferidos de diversas medidas observadas (ítems o variables manifiestas), lo que proporciona flexibilidad a la teoría del modelo. Dentro de las principales características del SEM, destacan:

- El uso de estándares de alta calidad para el análisis estadístico (Bagozzi y Fornell, 1982), que pretenden ayudar a enlazar datos y teoría de una manera flexible e interactiva (Fornell, 1982).

- Analiza en forma simultánea variables teóricas no observables, el criterio múltiple, constructos predictores y la utilización de las técnicas de una manera confirmatoria más que exploratoria.
- Considera el modelo propuesto como verdadero y busca la mejor estimación de parámetros.
- Los modelos estadísticos *testados* son descritos en forma gráfica y lenguaje sencillo.
- La replicación de la investigación se vuelve fácil, tomando en cuenta las matrices, tablas y otros datos generados en el estudio.
- Facilita la valoración de las mediciones multiítems en cuanto a la confiabilidad, validez y dimensionalidad (Diamantopoulos y Winklhofer, 2001).
- Permite el modelado de análisis *path* estandarizado con variables latentes.
- Es más útil cuando se pretende dar validez a un modelo de investigación, en lugar de buscar un modelo adecuado.

Así también, Fornell (1982) indica que el análisis multivariante destaca los elementos acumulativos del desarrollo de la teoría, aceptando conocimiento *a priori* que es incorporado en el análisis empírico y puede crearse de la teoría, de hallazgos empíricos o del diseño de la investigación. Se asume que la varianza y covarianza de los constructos independientes también son estimados, a diferencia del análisis de primera generación, que considera solo datos conocidos.

Una técnica derivada del SEM basada en las varianzas es la de Mínimos Cuadrados Parciales (*Partial Least Squares – PLS*), cuyo objetivo es la predicción de las variables latentes; se apoya en el análisis de componentes principales y en la estimación de mínimos cuadrados ordinarios (Cepeda y Roldán, 2004). Así mismo, es una herramienta para la comprobación de modelos complejos en las ciencias sociales y en la administración, donde los constructos no son observables directamente (por ejemplo: satisfacción, toma de decisiones, imagen organizacional, uso de los sistemas de información, entre otros), que nace principalmente en el campo de la mercadotecnia y que ha ido mejorando a través de los años en el desglose de sus resultados, debido a las críticas recibidas en forma de retroalimentación.

Nomograma de PLS

Los constructos son abstracciones de la realidad y de la teoría que no son captados directamente, es decir, su estimación es con base en una serie de ítems (variables manifiestas) que buscan medir lo que el investigador requiere. Después de su validación se convierten en variables para un proyecto en particular, lo que es realmente útil en los análisis de datos. Para Rigdon (2012), los constructos en un modelo son sólo

proxies cuyo objetivo es la modelización proyectada, por lo que existirá una brecha de validez entre los conceptos teóricos y estos *proxies*. Los constructos se representan en el nomograma a través de círculos, óvalos o hexágonos (compuestos).

Para la operacionalización del SEM es necesario el nomograma de PLS, por ello Barclay et al. (1995) explican los conceptos esenciales:

a. Constructo teórico, variable latente o no observable. En el nomograma se representa con un círculo. Conocidos también como exógenos (ξ), que proceden como variables predictoras o *causales* de constructos endógenos (η). Un constructo exógeno se asocia con una variable independiente (X) y un endógeno se relaciona con la idea de dependiente (Y).

b. Ítems (indicadores), variables manifiestas u observables. Se simbolizan en el gráfico por medio de cuadrados y existen de dos tipos:

- Ítem formativo: involucra que el constructo se exprese como una función de los ítems, un causante o predecesor del constructo; es decir, son variables observadas que causan una variable latente (Diamantopoulos y Winklhofer, 2001).

- Ítem reflectivo: los ítems observados son expresados en función del constructo, en otras palabras, que reflejan o son manifestaciones del constructo (Sellin, 1995), lo que explica la varianza (Fornell y Bookstein, 1992).

Por lo anterior, los modelos de investigación utilizados en PLS reciben el nombre de nomogramas, que son una representación gráfica de las relaciones entre las variables dependientes e independientes establecidas en forma hipotética, las cuales son de importancia por su facilidad de lectura e interpretación de los resultados obtenidos en los análisis realizados por el investigador.

Las flechas del nomograma manifiestan las relaciones supuestas, a comprobar entre los constructos, y éstos a su vez con los ítems asignados teóricamente. Hasta hoy, las flechas son exclusivamente unidireccionales, lo que se interpreta como un tipo de causalidad, donde un constructo independiente X manifiesta un efecto positivo/negativo sobre una variable dependiente Y, pero no realiza análisis de Y sobre X. No obstante, se espera que en el corto tiempo, con el desarrollo de PLS consistente (PLSc), se permita establecer relaciones en ambas direcciones en forma simultánea.

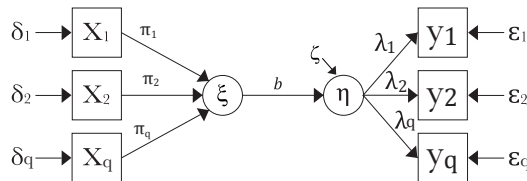
Otro elemento dentro del nomograma es el error (ζ), siempre presente en todo análisis estadístico de alto impacto; es representado en PLS como la varianza no explicada cuando se estima un modelo de investigación. A continuación, se presenta la simbología de los elementos que forman parte de un nomograma:

- x: Variables x (formativas), medidas o mejor conocidas como ítems formativos.
- y: Variables y (reflectivas), medidas o mejor conocidas como ítems reflectivos.
- b: Coeficiente de regresión simple entre ξ y η , también llamado coeficiente *path* estandarizado (β).

- δ : Residuos provenientes de las regresiones.
- ε : Términos de error.
- ζ : Residuo en el modelo estructural.
- π : Pesos de regresión (solo ítems formativos).
- λ : Cargas (solo ítems reflectivos).
- ξ : Constructo exógeno o predictor o causal (independiente).
- η : Constructo endógeno (dependiente).

La Figura 6.1 presenta gráficamente los símbolos y conceptos descritos previamente y utilizados en el desarrollo de un modelo en PLS-SEM.

Figura 6.1. Configuración para PLS-SEM



Fuente: elaboración propia

6.2. Fundamentos de Mínimos Cuadrados Parciales (PLS)

El modelado PLS fue desarrollado por Herman Wold, mentor de Karl Jöreskog, fundador de SEM, con el propósito de mostrar las condiciones presentes en las investigaciones de las ciencias sociales, las que incluyen a las del comportamiento. El objetivo de PLS es la predicción de las variables dependientes, en un esfuerzo por maximizar la varianza explicada de las variables independientes (en ambos tipos de variable pueden ser latentes o manifiestas), por lo que sus parámetros están basadas en la idea de empujear las variables residuales de las variables endógenas (Cepeda y Roldán, 2004); por su parte, Wold (1980) señala que la meta de PLS es manifestar las condiciones teóricas y empíricas en la investigación social, donde son habituales las teorías poco consolidadas y escasa información disponible.

A esta modelización también se le conoce como flexible (*soft modeling*), al no realizar suposiciones relativas a niveles de medidas, distribución de datos y el tamaño de la muestra (que puede ser pequeña o mediana; no obstante, es un tema a tratar con mucho cuidado), y su algoritmo consiste en una secuencia de regresiones en términos de vectores de peso (Henseler et al., 2009) que, obtenidos en la convergencia, satisfacen las ecuaciones de punto fijo. De esta manera, en las ciencias sociales PLS no requiere forzosamente una normalidad

en los datos ni que los ítems sean independientes unos de otros (Chin, 2010); así mismo, evita dos problemas serios (Fornell y Bookstein, 1982): *i*) soluciones inadmisibles o inapropiadas (ejemplo: estimaciones negativas de la varianza de los ítems y correlaciones mayores a 1), y *ii*) indeterminación de factores, al definir explícitamente los compuestos (variables latentes) y con ello la disponibilidad de las puntuaciones de dichas variables latentes.

Desde su aparición, PLS ha cambiado para bien la forma de validar los constructos y el modelo de investigación. Sus múltiples bondades han sido la principal causa de su aplicación porque, sin duda, el número de constructos formativos y reflectivos que se manejan es un elemento clave y base para la elección de una herramienta estadística, en la que el investigador debe especificar qué ítems son los asociados con cada constructo y proponer una dirección causal para las relaciones entre ellos (Fornell y Larcker, 1981; Barclay et al., 1995).

El algoritmo PLS básico, como lo sugiere Lohmöller (1989), incluye las siguientes tres etapas:

- Etapa 1: la estimación iterativa de las puntuaciones de las variables latentes consiste en un procedimiento de 4 pasos, que se repite hasta que se obtiene la convergencia o se alcanza el número máximo de iteraciones.
- Etapa 2: estimación de pesos/cargas externas y coeficientes *path* estandarizados.
- Etapa 3: estimación de los parámetros de localización.

Algunas de las principales características y ventajas de PLS se resumen a continuación:

- Analiza constructos multiítems con dirección directa e indirecta.
- No define valores mínimos en la escala de medida, que puede ser nominal, ordinal, por intervalos o ratios (Sosik et al., 2009); no obstante, es preciso mantener cierta precaución con las variables dependientes.
- Estima modelos complejos y errores con procedimiento de remuestreo (*bootstrapping*) (Chin et al., 2003).
- Barclay et al. (1995) lo recomiendan para la definición de modelos de investigación que induzcan hacia la predicción, en busca del desarrollo de una nueva teoría.
- Para el tamaño de la muestra se han mencionado algunas bases erróneas que sugerían que con 30 casos es suficiente. Es preciso encontrar otras maneras de determinarla, y para ello se reconoce que su tamaño muestral no es de dimensiones adecuadas para asegurar su potencial.
- Impide dos problemas graves: las soluciones inadmisibles y la indeterminación del factor (Fornell y Bookstein, 1982).
- Estima los modelos *path*, que comprenden constructos latentes que son medidos

indirectamente por los múltiples ítems (Mathieson et al., 2001).

- Ajusta simultáneamente el modelo estructural y de medida.

Permite el uso de ítems formativos y reflectivos, sin problema de identificación (Chin, 2010).

- Es aplicable cuando la intención es la búsqueda de información y la teoría no está totalmente desarrollada (Fornell y Bookstein, 1992).
- Analiza cada ítem de manera separada y cada uno puede variar la cantidad de impacto en la estimación del constructo (Chin et al., 2003), por lo tanto, ítems con relaciones pobres contienen pesos bajos.
- Se busca más la idea de predictibilidad que de causalidad.
- Las estimaciones de PLS se dirigen hacia los parámetros reales de la población por medio del incremento del número de ítems y del tamaño muestral (Chin et al., 2003).
- Los datos no necesariamente cuentan con una distribución normal o modelo específico para testar.

Indudablemente, es un hecho la generalización de la utilización de PLS a nivel mundial; desafortunadamente muchos tienen ideas erróneas de lo que realmente es esta herramienta estadística. Rigdon (2016) expone los argumentos incorrectos para su uso o rechazo:

- Argumentos incorrectos para apoyar su uso: permite muestras pequeñas, distribución de los datos no normales, presencia de indicadores formativos y solo aplicado a investigación de carácter exploratorio.
- Argumentos incorrectos para rechazar su uso: estima con sesgo los parámetros, no modela variables latentes (no es SEM), no cuenta con medidas de ajuste del modelo y no está libre de los errores de medida.

A fin de diferenciar y separar los valores de los ítems utilizados en PLS y proporcionar un orden a los índices, la modelación se realiza en dos etapas: Modelo de Medida y Modelo Estructural, que se presentan en los siguientes apartados.

6.2.1. Modelación PLS en el Modelo de Medida

Analiza si los constructos están medidos correctamente por medio de los ítems (variables) observados. Su análisis es a través de los atributos de validez (medir lo que se pretende medir) y fiabilidad (resultados estables y consistentes en mediciones repetidas). Asimismo, esta etapa implica el análisis de la *i*) fiabilidad individual del ítem, *ii*) la consistencia interna, *iii*) la validez convergente y *iv*) la validez discriminante.

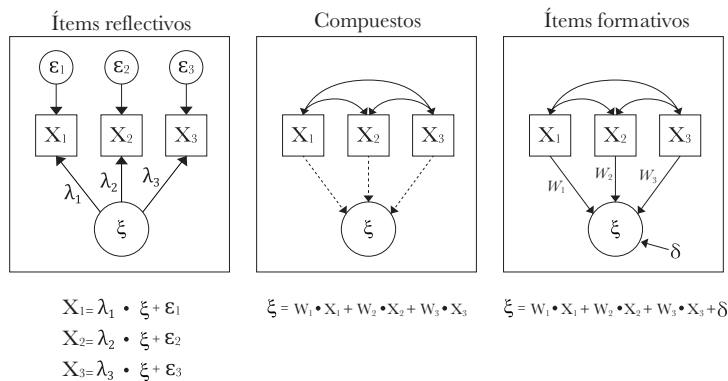
El Modelo de Medida representa las relaciones epistémicas que describen la relación entre una variable latente y sus ítems al mostrar cómo se mide el constructo

por medio de un conjunto de ítems (Roldán y Cepeda, 2017). Dichos ítems pueden ser reflectivos, formativos o compuestos y su fundamentación debe considerarse con base en la naturaleza del constructo (Henseler, 2017). Sí se estudian:

- Constructos comportamentales (*behavioral constructs*). Se organizan como modelos de factor común (ítems reflectivos) y si además se cuenta con ítems causales o formativos se puede aplicar el *modelo de medida causal-formativo*. Ejemplo: actitud de consumidores.
- Constructos de diseño o artefactos (*design constructs*). Se modelan como compuestos o componentes (*composite models*) que son justificados teóricamente. Ejemplo: eficiencia administrativa, compuesta de ahorro en costos, imagen corporativa y atracción de clientes.

La Figura 6.2. muestra en detalle los tipos de ítems que se utilizan en el desarrollo de un análisis con PLS.

Figura 6.2. Tipos de Modelo de Medida



Fuente: Henseler, Ringle y Sarstedt (2016a).

El Modelo de Medida, en general, se divide en tres formas de operar y que al final serán la pauta principal para realizar los análisis pertinentes. Se ilustran en la figura 6.2. y enseguida se describe su proceso de uso.

a. Modelo de Medida Reflectivo (Modelo de Factor Común)

En el nomograma las flechas indican una dirección del constructo hacia los ítems. En este modelo se sostiene que la varianza de un conjunto de indicadores (X_i) está explicada por un constructo (el factor común) (ξ) y sus errores (ϵ_i) siempre presentes (Roldán y Cepeda, 2017); del mismo modo, se asume que los errores (ϵ_i) que cada

ítem/medida posee no están relacionados con otros constructos o variables o con errores en el modelo. Como se ha establecido, la variable latente no se encuentra en algún problema o fenómeno social, solo se apoya en las correlaciones de sus ítems; por tal motivo comparten un tema o problema común que se espera correlacionen. Al ser un constructo no observable directamente, sus ítems pueden intercambiarse o eliminarse/modificarse mientras muestren los valores mínimos aceptados, no habiendo alteración en el significado del constructo bajo análisis.

Ejemplos: Constructo: Uso de Computadoras

- Cuento con las habilidades necesarias para operar una computadora.
- Obtengo el mejor provecho de la información proporcionada por la computadora.
- La computadora es una herramienta en mi vida laboral.

b. Modelo de Medida Formativo (Modelo de Indicadores Causales)

En el nomograma la dirección de las flechas ocurre de manera contraria a la medida reflectiva, es decir, en el formativo las flechas van de los ítems hacia el constructo. En este modelo se sostiene que la varianza de un constructo (ξ) está explicada por un conjunto de ítems (x_i) y su término de error (δ); este último elemento indica que el constructo no ha sido completamente medido y evaluado por sus indicadores, que no se espera que estén correlacionados. De la misma forma, los ítems se pueden referir a diferentes temas y eliminar uno o más, lo que incrementa el error del constructo en su conjunto; por lo tanto, los análisis de confiabilidad y consistencia interna no son necesarios para esta relación de ítem-constructo, solo es un requisito atender la aparición de la multicolinealidad. Posteriormente se detallan los índices básicos para su evaluación.

Ejemplo: Hornear un pastel

- Cuento con la harina suficiente para hacer un pastel.
- La mantequilla proporciona mejor sabor a los pastelillos.
- El horno cuenta con diferentes niveles de cocción.
- Los huevos para batir son de calidad comprobada.
- La levadura posee las características de esponjar sin quemar el pastel.
- El azúcar proporciona la dulzura pertinente.
- La leche proporciona la consistencia requerida.

c. Modelo Compuesto (Composite)

Este modelo compuesto en PLS cuenta con la variante de ser estimado de dos formas (Roldán y Cepeda, 2017):

- Modo A (*correlation weights*): los ítems se espera que estén correlacionados, también se puedan aplicar medidas de consistencia interna (ejemplo: pastel compuesto de harina, huevos, mantequilla, etc.)

- Modo B (*regression weights*): los ítems se espera que no estén correlacionados, representando una idea similar al modelo de medida formativo.

En la Tabla 6.1. se muestra una comparación de las diferencias de los tres modelos descritos anteriormente:

Tabla 6.1 Diferencias entre los modelos

Factor	Factor común (reflectivo)	Causal-Formativo	Compuesto
Relaciones entre el constructo y los ítems	El constructo causa sus ítems	Los ítems causan el constructo	Los ítems componen el constructo
Patrón correlacional esperado entre los ítems	Se esperan altas correlaciones	No hay razón para esperar que las medidas correlacionen	Altas correlaciones son comunes, pero no requeridas
Validez de la puntuación de la escala	No representa de forma adecuada al constructo	No representa de forma adecuada al constructo	Representa de forma adecuada al constructo
Tratamiento del error de medida	El error de medida se considera a nivel de ítem	El error de medida se considera a nivel de constructo	No incluye error de medida
Consecuencias de eliminar un constructo	No altera el significado del constructo	Incrementa el error de medida a nivel de constructo	Altera al compuesto y puede cambiar el significado
Red nomológica	Se requiere que los ítems tengan los mismos antecedentes y consecuencias	No se requiere que los ítems tengan los mismos antecedentes y consecuencias	Se requiere que los ítems tengan las mismas consecuencias

Fuente: Henseler (2017).

6.2.2. Modelación PLS en el Modelo Estructural

Este modelo es posible subdividirlo en Modelo Interno (*inner model*), en el cual se aprecian las relaciones (*paths*) entre los constructos basados principalmente en una teoría definida y estructurada claramente o la experiencia ganada por el investigador. Dichos *paths* indican cómo es que se relacionan los constructos planteados. En esta relación de constructos se distinguen aquellos definidos como exógenos (ξ), que son predictores o causales de otros constructos en el nomograma, y los constructos endógenos (η), que son explicados en el modelo (Roldán y Cepeda, 2017). En la Figura 6.1. se mostró gráficamente su representación en el nomograma de PLS.

En cuanto al Modelo Externo (*outer model*), representa las relaciones entre los constructos y sus ítems (indicadores) que, igual que en el modelo estructural, se establecen mediante una teoría de medida, porque se especifica cómo las variables latentes se miden, ya sea en Modo A (reflectivas), con el peso de correlación (*correlation weights*) y en la cual las flechas direccionales van del constructo a cada uno de sus indicadores. O puede ser en Modo B (formativas), medidas con el peso de la regresión (*regression weights*), donde las flechas direccionales inician en los ítems para terminar en los constructos.

6.2.3. Estimaciones PLS, PLSc y PLSpredictivo

Hasta este punto se ha escrito acerca de **PLS tradicional** en Modo A para los constructos reflectivos y Modo B para los formativos, y en ambos casos solo son aproximaciones (Sarstedt et al., 2016) y adaptados para la modelación de los compuestos (*composites*), donde sus estimaciones están sesgadas cuando se utilizan con la modelación con factores comunes (cargas sobrevaloradas y coeficientes *path*-relación-infravalorados); por lo tanto, surge el **PLS Consistente** (PLSc) para los Modelos de Factor Común (Modelados Reflectivos) que, de acuerdo a Dijkstra y Henseler (2015), PLSc hace uso de una corrección para generar estimaciones robustas en modelos donde los constructos subyacentes están modelados a través de factores comunes.

Sarstedt et al. (2016) determinaron que se produce sesgo cuando se emplea PLS tradicional (composite) para estimar modelos de factor común, y con el uso de SEM, con fundamento en las covarianzas (CBSEM), para estimar modelos basados en compuestos. De la misma manera, cuando se desconoce la naturaleza de los datos de un modelo de factor común o modelo composite es preferible utilizar PLS tradicional (Roldán y Cepeda, 2017).

Por otra parte, PLS se ha introducido como un enfoque *causal-predictivo*, diseñado para superar la aparente dicotomía entre explicación y predicción. Sin embargo, mientras que los investigadores que utilizan PLS-SEM rutinariamente enfatizan la naturaleza predictiva de sus análisis, el modelo se basa exclusivamente en métricas diseñadas para evaluar el poder explicativo de *coeficientes path*. Investigaciones recientes han propuesto PLSpredict, un procedimiento basado en muestras reservadas que genera predicciones a nivel de caso en un elemento o en un nivel de constructo (Shmueli et al., 2016; Shmueli et al., 2019).

6.2.4. Tamaño de la Muestra

Desde su creación, los investigadores han pensado que PLS se maneja de manera óptima con muestras pequeñas, lo cual es un error que puede traer consigo inconsistencias en los resultados de los análisis. No obstante, Roldán y Cepeda

(2017) indican que sí se pueden hacer estimaciones con muestras pequeñas, ya que el algoritmo de PLS aplica un procedimiento de segmentación, en donde un modelo complejo se divide en subconjuntos; sin embargo, dicha pequeñez dependerá del tamaño del universo.

Existen diversas formas de obtener el tamaño muestral: *i)* algunos creen que el 25% del universo es adecuado, *ii)* otros usan fórmulas matemáticas e incluso existen páginas Web que ayudan a determinar este valor, *iii)* se puede revisar el *software* G*Power presentado por Faul et al. (2009), *iv)* otros investigadores recomiendan que la muestra sobrepase los 100 casos que permitan desarrollar el potencial del *software* PLS y que los resultados sean de una calidad aceptable; de igual modo, algunos editores de *Journals* manifiestan que este número es una regla no escrita como primer punto de contacto de aceptar a revisión un artículo científico, *v)* además de lo anterior, existe otra manera: en una idea inicial, Barclay et al. (1995) indican que es preciso encontrar la regresión múltiple en el modelo propuesto, y con base en el nomograma determinar cuál es la mayor entre a) el bloque con más número de ítems formativos, o b) el número mayor de relaciones *path* dirigidos a un constructo latente dependiente, y con ello aplicar una regla heurística: el resultado multiplicarlo por 10 (se recomienda no usar esta regla), y *vi)* otra forma más robusta y precisa, que además es una validación interna, es especificar el *tamaño del efecto* (*effect size*) por cada relación presente y al mismo tiempo consultar la aproximación a las tablas de potencia de Jacob Cohen de 1988 o las desarrolladas por Green (1991) (Tabla 6.2):

- Efecto del tamaño (*effect size*) (f^2): determina el grado en que el fenómeno bajo análisis existe en la población objetivo.
- Potencia (*power*): probabilidad de rechazar correctamente la hipótesis nula cuando sea necesario.
- Alfa: probabilidad de presentar un falso positivo, la prueba presenta significancia estadística cuando no la tiene.

Tabla 6.2. Tamaño de la muestra requerida. Power=0.80, Alfa=0.05

Número de Predictores	Tamaño de la muestra basada en el poder de análisis		
	Tamaño del efecto		
	Pequeño	Medio	Grande
1	390	53	24
2	481	66	30
3	547	76	35

Continúa Tabla 6.2.

Número de Predictores	Tamaño de la muestra basada en el poder de análisis		
	Tamaño del efecto		
	Pequeño	Medio	Grande
4	599	84	39
5	645	91	42
6	686	97	46
7	726	102	48
8	757	108	51
9	788	113	54
10	844	117	56
15	952	138	67
20	1066	156	77
30	1247	187	94
40	1407	213	110

Fuente: Green (1991).

6.2.5. Software PLS

La masificación de PLS fue con el surgimiento de PLS Graph, desarrollado por Wynne W. Chin de la Universidad de Houston, en Estados Unidos de América. Hoy en día existe una amplia variedad de *softwares* listos para instalarse, pero indudablemente el de mayor envergadura y más usado a nivel mundial, reflejado en las publicaciones científicas de alto impacto, es el *software* SmartPLS.

A continuación se proporciona una lista de los *softwares* más utilizados en la investigación científica y que cuentan con descargas gratuitas, con algunas restricciones de operación:

- ADANCO (<https://www.composite-modeling.com/>)
- PLSpredict (<https://github.com/ISS-Analytics/pls-predict>)
- SmartPLS (<https://www.smartpls.com/>)
- VisualPLS (<http://www2.kuas.edu.tw/prof/fred/vpls/>)
- WarpPLS (<http://www.warppls.com/>)
- XLSTAT (<https://www.xlstat.com/es/soluciones/funciones/aproximacion-pls>)

6.3. Evaluación de la Modelación con PLS

PLS se ha caracterizado por la robustez en sus análisis, pero a la vez en su sencillez para generarlos e interpretarlos. Como se ha dicho, el constructo regularmente se compone de varios ítems, que representan cada uno de los aspectos o condiciones de un concepto o variable. Para ello es necesario distinguir dos:

- Fiabilidad (*reliability*). Representa la consistencia de un ítem, que suele ser fiable cuando se obtienen resultados consistentes (casi iguales) en condiciones parecidas de investigación.
- Validez (*validity*). Indica el nivel con que los ítems de un determinado constructo miden en grupo lo que deben medir.

Considerando los dos puntos anteriores, enseguida se detallan los índices básicos para la evaluación del Modelo de Medida Estimado en PLS en sus dos versiones: Modo A (reflectivos) y Modo B (Formativos).

6.3.1. Modelos de Medida Estimados en Modo A (Reflectivos)

- Fiabilidad individual del ítem. Son los ítems o medidas de las variables observables (manifiestas) con relación a sus variables latentes correspondientes. Es valorada analizando las cargas (λ) o correlaciones simples. El ítem debe alcanzar al menos un valor de 0.707 (50% de la varianza explicada) (Chin, 1998), implicando que la varianza compartida entre el constructo y sus ítems sea mayor que la varianza del error (Sha'ri y Aspinwall, 2000), pero aún más estricto (0.8) para investigaciones básicas; sin embargo, algunos practicantes e investigadores creen que esta regla heurística no debería ser tan severa en las etapas iniciales de desarrollo de escalas o en contextos diferentes (Chin, 1998). Cuando se obtienen resultados no satisfactorios en los ítems de cargas, deben ser eliminados y volver a ejecutar PLS.
- Fiabilidad del constructo (consistencia interna), también conocida como fiabilidad de la escala, tomando en cuenta que establecer niveles aceptables de la validación de los constructos es fundamental cuando se usa SEM. Su fin es determinar si los ítems de un constructo lo están midiendo realmente, debiendo ser similares sus puntuaciones. Las medidas utilizadas son: i) el alfa de Cronbach, que mide la consistencia interna de los factores basada en el promedio de las correlaciones entre los ítems (Nunnally, 1978), y se requiere un coeficiente superior a 0.7 para considerarlo bueno, o ii) la fiabilidad compuesta (*composite reliability*), que algunos autores manifiestan que es mejor que el alfa de Cronbach y por ello su mayor uso en el mundo. Sobre esto, Fornell y Larcker (1981) argumentan que han mejorado este índice al encontrar en sus investigaciones un valor mínimo de 0.707. Del mismo modo, Nunnally y Bernstein (1994) sugieren un 0.7 como nivel adecuado para investigaciones tempranas y un 0.8 o 0.9 para investigaciones avanzadas. En últimas fechas se ha utilizado el indicador rho_A de Dijkstra y Henseler (2015), que también requiere un mínimo de 0.7.
- Validez convergente. Considera que un grupo de ítems está representado por un único constructo subyacente, demostrado con la unidimensionalidad (Henseler et al., 2009). Analizada por la Varianza Media Extraída (*Average Variance Extracted*

-AVE), que mide la cantidad de varianza que un constructo captura de sus ítems relativa a la varianza contenida en el error de medición y debe ser mayor al cuadrado de las correlaciones entre los constructos, sus valores deben estar por encima de 0.50 (50% de la varianza del constructo es conformada por sus ítems asignados) (Fornell y Larcker, 1981) o el valor de la *t-student (statistic)* significativo. Este criterio solo puede aplicarse a ítems reflectivos.

- Validez discriminante. Con esta valoración se detecta si un determinado constructo es diferente a otro dentro del modelo (Barclay et al., 1995). La idea esencial es que ningún ítem de un constructo debe de *cargar* más en otro que no sea el que está midiendo (deben existir correlaciones débiles entre estos constructos) y, a la vez, cada constructo debe de *cargar* más sobre los ítems que lo componen. Existen dos índices para obtener esta valoración: *i*) criterio de Fornell y Larcker (1981), interpretada como la cantidad de varianza que los ítems proporcionan a su constructo y que debe ser mayor a la varianza que comparte con otro dentro de un modelo PLS. Para facilitar este proceso se puede obtener la raíz cuadrada de AVE (la matriz que se genera proporciona estos valores), que debe ser mayor a la correlación entre los constructos (Hair et al., 2019), o *ii*) por el estadístico HTMT, que representa el promedio de las correlaciones *heterotrait-heteromethod* con relación al promedio de las correlaciones *monotrait-heteromethod* (Henseler et al., 2015), que deben ser menores a 0.85. En la Tabla 6.3 están representados cinco constructos, y para comprobar la validez discriminante su coeficiente (se aprecia en la intersección marcada en negrita) debe ser mayor al resto de los coeficientes que se encuentran en su fila y su columna (inter-constructo). Y en la Tabla 6.4. (para HTMT), donde el valor de cada uno de los índices cumple satisfactoriamente con los valores permitidos.

Tabla 6.3 Ejemplo de Matriz de Validez Discriminante

Constructo	Cap	Mot	Cal	TD	Sat
Capacitación (Cap)	0.836				
Motivación (Mot)	0.483	0.737			
Calidad (Cal)	0.508	0.521	0.808		
Toma Decisiones (TD)	0.287	0.398	0.343	0.755	
Satisfacción (Sat)	0.720	0.555	0.494	0.503	0.830

Fuente: Elaboración propia.

Tabla 6.4 Ejemplo de correlación de variables: Heterotrait-Monotrait Ratio (HTMT)

	A	B	C	D	E
A					
B	.460				
C	.510	.630			
D	.597	.589	.600		
E	.629	.549	.615	.848	

Fuente: Elaboración propia.

6.3.2. Modelos de Medida Estimados en Modo B (Medidas Formativas)

Con la presunción de que los ítems formativos no requieren estar correlacionados y libres de error, de acuerdo con Bagozzi (1994), la evaluación de fiabilidad y validez no es aplicable. Por lo anterior, su evaluación se lleva a cabo en dos niveles (Chin, 2010): ítem y constructo.

Nivel Ítem

- Multicolinealidad. Este concepto se refiere a las interrelaciones lineales que se encuentran entre dos o más ítems formativos. Una alta colinealidad indica estimaciones inestables, al no poder identificar el efecto de cada ítem en el constructo; por lo tanto: *i*) aumenta el error estándar y *ii*) estimación de pesos errónea, que inclusive puede cambiar los signos. Esta validación se lleva a cabo por medio de pruebas estadísticas, la más utilizada es la propuesta por Diamantopoulos y Siguaw (2006), denominada Factor de Inflación de la Varianza (FIV), cuyo valor superior a 3.3 indica alta multicolinealidad; no obstante, pueden existir factores bajos que no incluyan a todas las variables independientes y, por lo tanto, no son detectadas por el FIV.
- Pesos (*weights*) y su significación. Permite entender la estructura o composición de las variables latentes. Los pesos indican la contribución individual a la formación de un constructo, y a medida que aumenta el número de ítems, disminuye la media de los pesos y por consiguiente pueden existir pesos no significativos. El máximo valor alcanzable para un indicador formativo es $1/n^{1/2}$, siendo n el número de ítems. Roberts y Thatcher (2009) consideran que no obstante que un ítem contribuye poco a la varianza explicada del constructo, se debe incluir, porque al eliminarlo se elimina una parte del constructo latente.

Nivel Constructo

- Validez externa. Se necesita que exista un grupo de ítems reflectivos validados con anticipación para el mismo constructo. Chin (2010) estima que se realiza por medio de un modelo de redundancia de dos bloques: *i*) por un lado, para el

mismo concepto teórico se cuenta con un grupo de medidas reflectivas validadas y, por otro, un grupo con medidas formativas, y *ii*) la relación entre los dos constructos debe ser alta, recomendando que sea mayor a 0.8 para alcanzar la validez externa (Mathieson et al., 2001). En caso de contar solo con medidas formativas, considerar lo expuesto para el primer bloque.

- Validez nomológica. Requiere disponer del constructo formativo en un modelo teórico en el lugar donde en análisis pasados estuvo incluido el constructo, pero de forma reflectiva (Chin, 2010), esperando que se mantengan las relaciones establecidas, ahora del modo formativo.
- Validez discriminante. Como se ha explicado en el modo reflectivo, lo discriminante indica el grado en que es diferente un constructo de otro dentro de un modelo PLS. Para esta valoración, si las correlaciones entre el constructo formativo y los demás constructos están a un valor inferior a 0.7, se entiende que los constructos son diferentes entre sí (Urbach y Ahlemann, 2010).

6.3.3. Evaluación del Modelo Estructural

a. Valoración de Colinealidad

Los coeficientes *path* estandarizados se generan con base en regresiones de mínimos cuadrados ordinarios (OLS), como en una regresión múltiple, por lo que es necesario evitar la multicolinealidad entre las variables que las preceden y los constructos endógenos; para ello, los valores FIV deben ser menores a 3.3, de no ser así, reflejaría un problema importante de colinealidad en el constructo; por tanto, se debe considerar su eliminación, fusión de varios o bien la creación de constructos de orden superior (Hair et al., 2017).

b. Evaluación del signo algebraico, magnitud y significancia estadística de los coeficientes path

Los coeficientes *path* estandarizados muestran el grado de las relaciones propuestas entre los constructos, representando las hipótesis planteadas entre ellos.

- Aquella relación que obtenga un signo (+ o -) contrario al establecido en su hipótesis implica que dicha hipótesis deberá ser rechazada.
- Para la magnitud, los coeficientes *path* estandarizados indican el número de probabilidad entre +1 y -1. A mayor valor, son más fuertes las relaciones entre los constructos y, por el contrario, cuanto más se acerca a 0 (cero) es más débil la relación y, de acuerdo con la investigación en la que se esté trabajando, sería rechazada la hipótesis (Roldán y Cepeda, 2017).
- La significancia se obtiene por medio del *bootstrapping* (remuestreo) que indicará si las relaciones en evaluación son significativamente diferentes a cero. Esta técnica se refiere precisamente a generar muestras aleatoriamente para reposición de

la muestra original (Hair et al., 2011), que, a la fecha, los expertos sugieren un valor de remuestreo de 5000. Sus resultados proporcionan valores tales como los errores estándar, los estadísticos t (*t-student*) y los intervalos de confianza, que permitirán en su conjunto agregar un valor más en la evaluación de las hipótesis. En este apartado el índice más utilizado es el *t-student*, cuyos valores mínimos a aceptar en un remuestreo de 5000 de una cola (Hair, 2019) son:

$t(0.05; 4999) = 1.645$, que representa $* p < 0.05$

$t(0.01; 4999) = 2.327$, que representa $** p < 0.01$

$t(0.001; 4999) = 3.092$, que representa $*** p < 0.001$

6.3.4. Otros índices para valoración en PLS

- Coeficiente de determinación (R^2), también conocido como varianza explicada. Representa una medida de poder predictivo, al proveer un indicativo de la habilidad de previsión de las variables independientes; es decir, manifiesta la cantidad de varianza explicada del constructo por su variable predictora. Para R^2 , está dentro de los parámetros de la probabilidad que oscile entre 0 y 1, y entre más se acerque a la unidad, cuenta con un mayor poder de predicción. En los tiempos actuales, para las ciencias sociales se siguen los parámetros de Chin (1998), quien señala que R^2 a un nivel de 0.67 tiene un efecto sustancial, 0.33 moderado y 0.19 débil y, junto con Sellin (1995), que se requiere determinar la relevancia predictiva de la varianza explicada por medio del estadístico de *Stone – Geisser test* (Q^2), que precisa ser mayor a 0.0 (cero punto cero) para que sea significativo.

- Coeficientes *Path* Estandarizados (β). Representan la relación (R) entre dos constructos. Muestran, en un número de probabilidad de 0 a 1, la existencia o no de una correlación. Son identificados en el nomograma de PLS por medio de las flechas que enlazan a los constructos en el modelo interno (estructural). Se obtiene como en la regresión múltiple; sobre ello, Chin (1998) indica que β , para ser considerado significativo y útil, debería alcanzar al menos un valor de 0.2 e idealmente situarse por encima de 0.3.

- *T-statistic* (*t-student*). Muestra la significancia de una relación establecida como hipótesis en PLS, que representa la relación de la salida del valor estimado de un parámetro, desde su valor hipotético hasta su error estándar. Para ser significativo, este índice debe ser menor a 0.05 ($p < 0.05$), es decir, al menos un 95% de confianza.

- Q^2 de *Stone-Geisser*. Mide la importancia de la predicción de los constructos dependientes. Este parámetro debe ser mayor a cero para que el constructo tenga validez predictiva (Chin, 1998). En contraste, los valores menores de cero indican una falta de validez predictiva (Hair et al., 2017).

Capítulo 7. Aplicación práctica de ecuaciones estructurales con base en la varianza

7.1. Introducción

El objetivo de este capítulo es conducir al lector a la creación y valoración de un modelo estructural basado en la varianza (PLS-SEM) a través de un *software* especializado. Los temas por abordar incluyen cómo crear un nomograma PLS-SEM que permita explicar el proceso de estimación de modelos de medida de tipo reflectivos y formativos, y el procedimiento de evaluación del modelo estructural; todo ello mediante un ejemplo. La Figura 7.1 presenta los pasos que se deben seguir.

Figura 7.1. Pasos a seguir para estimar un modelo PLS-SEM

Paso 1. Especificación del modelo propuesto
Paso 2a. Valoración del modelo de media reflectivo (Fiabilidad individual, de constructo, validez convergente y discriminante)
Paso 2b. Valoración del modelo de media formativo (Valores de cargas, Valores VIF de pesos y nivel de significancia)
Paso 3. Estimación del modelo estructural (Ajustes del modelo, colinealidad, coeficientes path (β) - significación, Varianza explicada R^2)

Fuente: Elaboración propia.

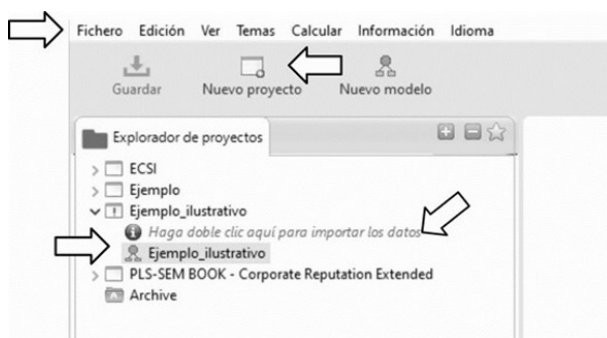
Previo a la descripción del procedimiento, es importante indicar la relevancia que tiene la especificación del modelo propuesto (paso 1, Figura 1), ya que su propósito es argumentar, con base en la teoría existente relacionada con el tema abordado, el porqué de la utilización de ciertas variables, su orden, agrupación o relación planteada, entre otras, por lo que los autores deberán justificarlo de la mejor manera, dado que ello sustentará su diseño. La herramienta estadística con la cual se explicarán los pasos 2 y 3 de la Figura 1 será el software SmartPLS versión 3.3.2. de Ringler et al. (2015), que ofrece una interfaz gráfica amigable e intuitiva. Se tomará como ejemplo ilustrativo un modelo del comercio electrónico en época de contingencia sanitaria que hipotetizará un nomograma que permita valorar el modelo de medida y las relaciones causales propuestas.

7.2. Creación de un nomograma en el software SmartPLS

En principio, hay que descargar el *software* SmartPLS del sitio <http://www.smartpls.com>, instalarlo e introducir la licencia correspondiente (el sitio ofrece un

licenciamiento de prueba con funcionalidad completa de forma gratuita por 30 días¹). Hecho esto, abrir el *software*, ir a la opción *Fichero* y del menú que se despliega seleccionar *crear un proyecto*; aparecerá una ventana emergente, la cual permite asignarle un nombre; por último, dar clic en el botón *ok*; el resultado de esta acción se muestra en la Figura 7.2.

Figura 7.2. Resultado de la creación de un proyecto en SmartPLS



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Antes de comenzar a pintar un nomograma PLS-SEM se debe adjuntar dentro de SmartPLS el archivo que contendrá el conjunto de observaciones por analizar en la investigación; para ello, hacer clic en el espacio designado para el proyecto en la opción *Haga doble clic aquí para importar datos* (ver Figura 7.2); aparecerá una ventana que permite buscar el archivo en la computadora; aceptarlo una vez localizado. Esto abrirá una nueva pestaña del lado derecho de la pantalla, la cual proporciona información de los datos cargados, tales como las variables o indicadores existentes en el archivo de observaciones, su estadística descriptiva básica, el tamaño de la muestra y otras (Figura 7.3). La extensión de los archivos a utilizar para este proceso deberá ser *.csv* (archivo delimitado por comas) o bien *.txt* (archivo de texto). También hay que considerar que el archivo de observaciones no deberá contener ningún elemento, variable o columna en formato de texto o cadena (por ejemplo, preguntas abiertas).

¹ Existe una versión para estudiantes gratuita, pero tiene la limitante de solo manejar un conjunto de datos de 99 observaciones. Para mayor detalle de los diferentes tipos de licenciamiento, visite el sitio <http://www.smartpls.com>.

Figura 7.3. Resultado de adjuntar archivo de observaciones al proyecto

FicheroEdiciónVerTemasCalcularInformaciónIdioma

Guardar

Nuevo proyecto

Nuevo modelo

Agregar datos de grupo

Generar grupos de datos

Borrar grupos de datos

Explorador de proyectos

> ECSI

> Ejemplo

> Ejemplo_Ilustrativo

> Ejemplo_Ilustrativo

> Observaciones [145 registros]

> PLS-SEM BOOK - Corporate Reputation Extended

> Archive

Observaciones.txt

Delimitador: Coma

Codificación: UTF-8

Volver a analizar

Ver en aplicación externa

Delimitador de texto: Ninguno

Tamaño de la muestra: 145

Formato de número: EE.UU. (por ejemplo: 1.000,23)

Indicadores: 15

Marcador de valor perdido: Ninguno

Valores perdidos: 0

Indicadores

Correlaciones de los indicadores

Fichero de datos originales

Copiar en el portapapeles

	Nº.	Perdido	Media	Mediana	Min	Max	Desviación estándar	Kurtosis excesiva	Asimetría
PU1	1	0	4.221	5.000	1.000	5.000	1.200	1.559	-1.595
PU2	2	0	4.348	5.000	1.000	5.000	1.166	1.865	-1.658
PU3	3	0	4.483	5.000	1.000	5.000	0.918	3.618	-1.975
PU4	4	0	4.483	5.000	1.000	5.000	0.983	4.694	-2.238
FP1	5	0	4.083	5.000	1.000	5.000	1.273	0.147	-1.170
FP2	6	0	4.262	5.000	1.000	5.000	1.180	1.436	-1.565
FP3	7	0	4.021	4.000	1.000	5.000	1.189	0.363	-1.110
FP4	8	0	4.269	5.000	1.000	5.000	1.059	1.369	-1.474
ATT1	9	0	4.552	5.000	1.000	5.000	0.862	5.497	-2.312
ATT2	10	0	4.434	5.000	1.000	5.000	0.916	3.038	-1.791
ATT3	11	0	4.124	5.000	1.000	5.000	1.264	0.582	-1.333
ATT4	12	0	4.441	5.000	1.000	5.000	0.975	4.000	-2.069

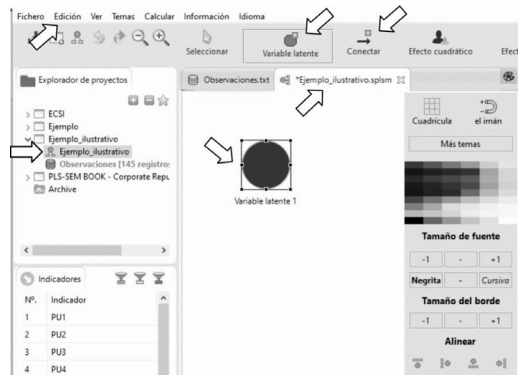
Indicadores

Numero de indicadores a mostrar.

Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Ya insertado el conjunto de datos en el proyecto, se puede comenzar a pintar el modelo (puede haber más de un nomograma o archivos de observaciones por proyecto). Para crear un nomograma hay que dar doble clic en el ícono que tiene el nombre del proyecto, en este caso se llama *ejemplo ilustrativo* y, de forma automática, del lado derecho de la pantalla se abre una pestaña vacía (ventana de modelización), la cual habilita la posibilidad de pintar variables latentes; para hacerlo, hay que ir al menú principal, opción *Edición* y luego a *añadir variable latente*; esto permitirá que cada vez que se dé clic con el botón izquierdo del *mouse*, se pueda dibujar un constructo en el área de modelización (ver Figura 7.4).

Figura 7.4. Ventana de modelización de nomogramas

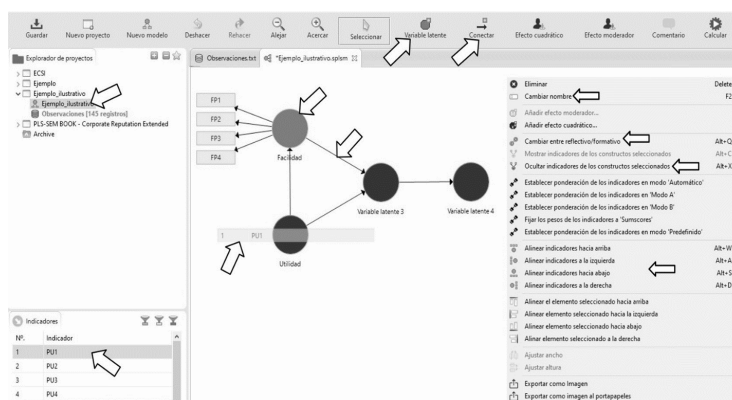


Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Una vez pintadas las variables latentes en la ventana de modelización, se puede personalizar cada una (por ejemplo, eliminarla, cambiar su nombre, alinear sus indicadores, entre otras); para ello hay que seleccionar una variable (o con clic en el ícono de *seleccionar* de la barra de menú) y después dar un clic con el botón derecho del ratón sobre ésta, logrando con ello desplegar un menú emergente con las opciones antes mencionadas (ver Figura 7.5).

Para dibujar las relaciones causales (flechas o rutas) entre las variables latentes, ir al ícono *conectar* que se ubica en la misma barra de herramientas (Figura 7.5), seleccionar una variable que se considera independiente o exógena y mover el cursor hasta estar encima de la variable latente endógena o dependiente. Esta acción se repite hasta haber pintado todas las relaciones causales estimadas en la investigación.

Figura 7.5. Menú emergente de opciones para una variable latente en SmarPLS



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Para asignar las variables manifiestas o indicadores a los constructos pintados en la ventana de modelización hay que localizar en la parte inferior del panel izquierdo de la pantalla la pestaña *indicadores*, que contiene las variables manifiestas que proceden del archivo de observaciones previamente cargado. Para adjuntar un indicador a una variable, solo se debe seleccionar uno del panel de indicadores y arrastrarlo hasta (sobre) el constructo deseado, y soltar; si la operación se realizó correctamente, el indicador aparece en la ventana de modelización como un rectángulo amarillo ligado a la variable mediante una flecha, pero también cambiará el color del constructo de rojo a azul. Al término de estos procedimientos con todos los indicadores y variables, la ventana será similar a la que se expone en la Figura 7.5. Es de notar que si se dejan variables latentes en color rojo, el *software* rechazará

el realizar algún cálculo estadístico. Por último, se debe guardar el proyecto, para ello, ir a *fichero* en el menú principal, y *guardar*.

7.3. Evaluación de los modelos de medida

Visto ya el diseño básico de un nomograma PLS, ahora se expondrá el procedimiento para la estimación de la calidad de lo obtenido, con base en sus datos. Se inicia con la valoración de las medidas empíricas existentes en las relaciones entre las variables manifiestas (indicadores) y las variables latentes. Debido a que los constructos pueden ser de dos tipos: reflectivo y formativo (pueden coexistir ambos en un modelo), deben analizarse separadamente. Empezaremos por la evaluación de los reflectivos.

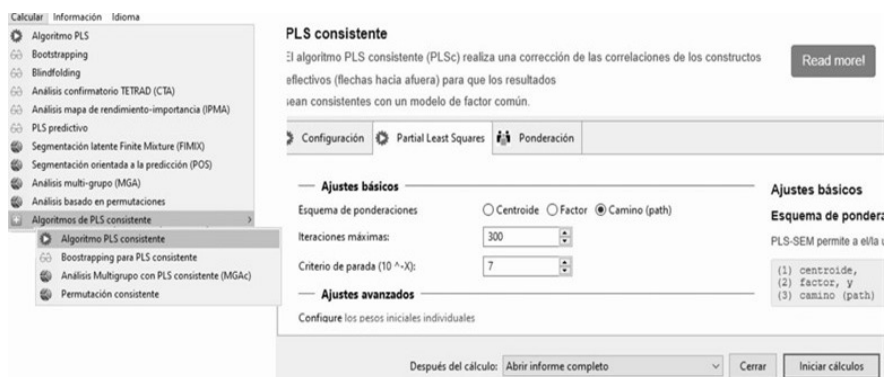
7.3.1. Estimación de constructos de tipo reflectivo

Este análisis permite la valoración de la fiabilidad y validez de los indicadores de un constructo y comprende la evaluación de la fiabilidad compuesta de los indicadores (consistencia interna), la validez convergente y la validez discriminante. La consistencia interna puede ser evaluada mediante el criterio de alfa de Cronbach, el cual aplica las siguientes escalas: 0.80 es satisfactorio para investigaciones en etapas avanzadas y 0.70 para investigaciones en etapas iniciales (Nunnally y Bernstein, 1994), mientras que un valor por debajo de 0.60 indica una falta de fiabilidad. Otra alternativa es el Índice de fiabilidad compuesta (P_c), donde valores de 0.60 a 0.70 son aceptables en investigaciones exploratorias, mientras que de 0.70 a 0.90 son óptimos (Henseler et al., 2009).

Los criterios de confiabilidad interna de un modelo PLS se obtienen con este proceso: seleccionar *calcular* del menú principal, enseguida, *Algoritmo PLS* si el modelo está formado de variables de tipo compuestos; de lo contrario, seleccionar *Algoritmo PLS_c*, adecuado para modelos combinados o solo con variables de factor común. En el ejemplo ilustrativo los constructos son de factor común, por ello se selecciona *Algoritmo PLS_c*. En ambos casos aparecerá una ventana como la que se visualiza en la Figura 7.6. Se recomienda dejar los valores por defecto y dar clic en *iniciar cálculos*.

Una vez ejecutado el Algoritmo PLS se despliega una ventana con una nueva pestaña que se divide en dos partes (Figura 7.7), una donde se muestran los resultados de un indicador o criterio (parte superior) y otra (inferior) que agrupa los resultados por categorías (ejemplo: resultados finales, criterios de calidad, resultados provisionales, base de datos). Primero, ubicar la columna *criterios de calidad* (parte inferior de la pestaña) y, como se observa, se mostrará una serie de opciones; elegir *fiabilidad y validez de constructo*.

Figura 7.6. Ventana de configuración del algoritmo PLS consistente



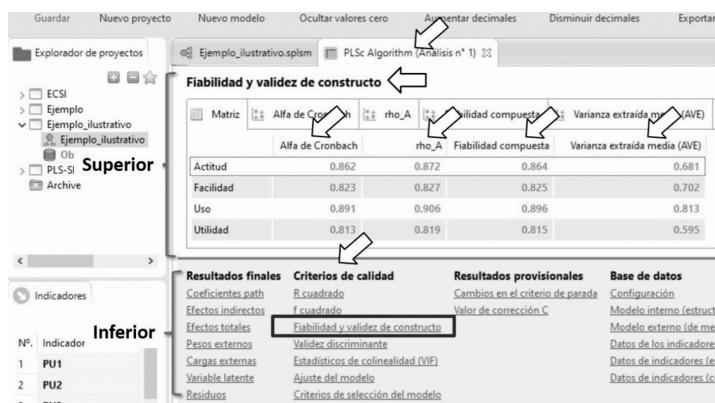
Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Lo anterior lleva a que en la parte superior de la pestaña aparezcan los resultados con el mismo título y, como se ve en la Figura 7.7, los valores obtenidos se muestran a detalle y en formato de tabla², apreciándose que se calculan para cada uno de los constructos presentes en el modelo. Es relevante mencionar que para los constructos de solo un ítem o formativos el valor presentado siempre será la unidad (1), por lo que no se toman en cuenta para el reporte escrito. Para el ejemplo ilustrativo, los valores obtenidos tanto para la fiabilidad como para la validez convergente del modelo cumplen con lo establecido.

En cuanto a la validez convergente, se aplica la técnica PLS para comprobar la correlación entre los indicadores de un constructo; para ello, la literatura indica que se debe aplicar la varianza media extraída (AVE) y las cargas externas de los indicadores. Para el AVE, según Chin (1998), su valor debe ser mayor a 0.50, ya que un valor por debajo significa que la varianza del error es mayor que la varianza explicada y, de ser así, debe ser rechazado. El proceso para obtenerlo del SmartPLS es el mismo que para los criterios de confiabilidad que se muestran en la misma pestaña (ver Figura 7.7).

² El software SmartPLS en su versión profesional permite exportar los datos a otras aplicaciones; para ello, dar clic en exportar en formato de Microsoft Excel o para el paquete R.

Figura 7.7. Ventana de resultados para fiabilidad y validez de constructos reflectivos



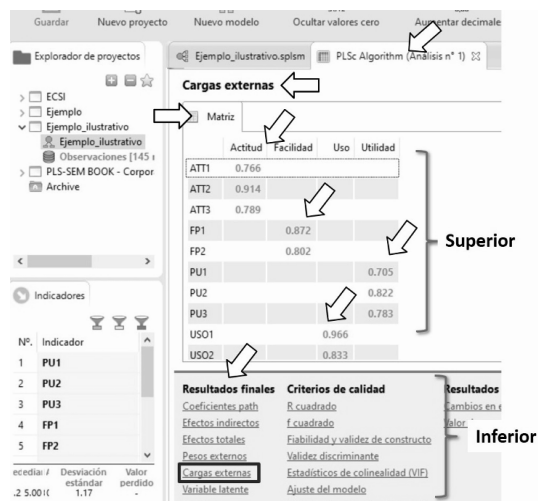
Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Las cargas externas (pesos estandarizados) en PLS-SEM significan la aportación total de un indicador en la determinación de un constructo. Su valor va de 0 a 1, pero entre más cercano a la unidad, la magnitud de la relación será más fuerte. La literatura menciona que en modelos de tipo reflectivo las cargas de un indicador deben obtener un valor mayor a 0.707 o, lo que es igual, al 50% de la varianza explicada por su factor. Pero si tienen valores estimados entre 0.400 y 0.707 se recomienda renunciar a ellos solo si esto mejora la fiabilidad compuesta del constructo (Henseler et al., 2009).

El proceso para obtenerlos es: elegir la opción *Algoritmo PLS* si su modelo está formado de variables de tipo compuesto; en caso contrario, seleccionar *Algoritmo PLS*, que es el adecuado para modelos combinados o solo con variables de factor común. En el caso ilustrativo, los constructos son de factor común, por eso se selecciona *Algoritmo PLS*. Sin embargo, en ambos casos aparecerá una ventana igual a la observada en la Figura 7.6. Se recomienda dejar los valores por defecto e iniciar cálculos.

Una vez terminados los cálculos, se despliega una ventana con una nueva pestaña, como se muestra en la Figura 7.8. En ella, ubicar la columna *criterios de calidad* (parte inferior de la pestaña) y luego elegir *cargas externas*. Entonces, en la parte superior aparecen los resultados con el mismo título. Los valores obtenidos se muestran a detalle y en formato de tabla, donde se puede apreciar que se estiman para cada uno de los indicadores presentes en el modelo, y solo hay que verificar y reportar aquellos que son considerados de tipo reflectivo. Para el caso del ejemplo ilustrativo que se sigue en este capítulo, se puede observar en la Figura 7.8 que los valores obtenidos por los indicadores, en sus cargas cumplen con lo establecido.

Figura 7.8. Ventana con los resultados de cargas externas de constructos reflectivos

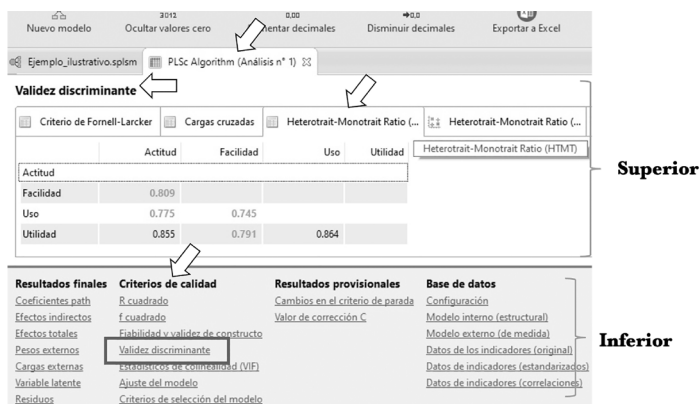


Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

En cuanto a la validez discriminante en la técnica PLS-SEM, ésta se debe aplicar para determinar el grado en que un constructo es distinto de los demás, es decir, que los constructos que no deben tener ninguna relación entre ellos no la tengan. En PLS también se puede aplicar para su evaluación la técnica ofrecida por Fornell y Larcker de 1981, pero las críticas de Henseler et al. (2015) llevaron a éstos a desarrollar y proponer una metodología para evaluarla denominada *heterotrait-monotrait ratio of correlations* o HTMT, en donde valores menores a 0.90 son adecuados, pero menores a 0.85 son óptimos, lo que significa que con ambas herramientas se aprueba la validez discriminante para los constructos analizados.

Para obtener los valores que permitan evaluar la validez discriminante del modelo es preciso seguir esta ruta: seleccionar *calcular* del menú principal y dar clic en *Algoritmo PLS*, o *Algoritmo PLSc* (según sea el tipo de constructos por analizar). En cualquiera de ellos aparecerá una ventana como la observada en la Figura 7.6. Dejar los valores por defecto e iniciar cálculos. Una vez ejecutado el algoritmo, se despliega una ventana con una nueva pestaña dividida en dos partes, como se muestra en la Figura 7.9; en una se detallan los resultados de un indicador (ítem) o criterio (parte superior) y en la otra se agrupan por categorías. Por lo tanto, primero es necesario ubicar la columna *criterios de calidad* (parte inferior de la pestaña), para posteriormente elegir *validez discriminante*.

Figura 7.9. Resultados de validez discriminante con el criterio HTMT



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

En la parte superior aparecen los resultados con el mismo título (Figura 7.9). La ventana muestra los resultados para el criterio de Fornell y Larcker, el de cargas cruzadas y el HTMT. En el ejemplo que se analiza y por las razones antes comentadas se elige el criterio HTMT; para ello, elegir la pestaña correspondiente. El *software* marca en verde si se cumple con valores menores a 0.85, en negro si son menores a 0.90 y en rojo si rebasan este último. En el ejemplo ilustrativo los valores obtenidos están dentro de los rangos establecidos, dado que todos son menores a 0.90.

7.3.2. Estimación de constructos de tipo formativo

El proceso para la valoración de los constructos reflectivos no es aplicable para los formativos, ya que, entre otras cosas, se asume que los constructos formativos están libres de error (Diamantopoulos, 2008). Para su estimación se recurre a la técnica *bootstrapping*, que ayuda a probar la significancia de los pesos de los indicadores, pero también se debe de probar su colinealidad.

En primer lugar, se deberá determinar la colinealidad. Una medida relacionada es el factor de inflación de la varianza (VIF, *Variance Inflation Factor*), que permite establecer niveles de tolerancia o valores críticos, con los que se puede detectar la existencia o no de potenciales problemas de colinealidad. Niveles de tolerancia ≤ 0.20 y un valor VIF ≤ 3.3 son indicativos de problemas con la colinealidad (Diamantopoulos y Sigauw, 2006), por lo que, de ser así, es necesario meditar su eliminación del modelo propuesto, dado que todos en su conjunto conforman la estimación de la variable y su eliminación alteraría dicho efecto.

Para obtener los valores VIF de un modelo de medida, el proceso es: seleccionar *Calcular* del menú principal y luego *Algoritmo PLS* si el modelo está formado de variables de tipo compuesto; en caso contrario, elegir *Algoritmo PLS_c*, adecuado para modelos combinados o solo con variables de factor común. De cualquier manera, en ambos casos aparecerá una ventana similar a la de la Figura 7.6; se recomienda dejar los valores por defecto e iniciar cálculos.

Hecho esto, aparecerá una pestaña que está dividida en dos partes, una donde se muestran los resultados de criterio o indicador seleccionado (parte superior) y otra (inferior) en la que se agrupan los resultados por categorías (resultados finales, criterios de calidad, resultados provisionales, base de datos). Se ubica la columna *criterios de calidad* (parte inferior de la pestaña) y, como se observa, aparecen varias opciones; elegir *estadísticos de colinealidad (VIF)*.

Enseguida, en la parte superior aparecen los resultados a detalle y en formato de tabla con el mismo título, pero para determinar si los criterios se aceptan para el modelo de medida, ir a la subpestaña *Valores VIF del modelo externo (modelo de medida)* y verificar que lo estimado para cada una de las variables formativas esté en los rangos aceptables, aunque, al igual que en los resultados del *bootstrapping*, el *software* marca en colores lo obtenido, para su mejor comprensión.

En cuanto a la magnitud de los indicadores y su significatividad, se deben examinar sus pesos y cargas externas; para ello, debe aplicarse el procedimiento *bootstrapping*, que permitirá valorar su significancia estadística. Si se presenta, entonces se cuenta con apoyo empírico para su retención; en caso contrario, hay que verificar su carga correspondiente, que debe ser mayor a 0.50. Pero si en ambos casos el indicador no es estadísticamente significativo o su carga es menor a 0.50, se debe eliminar (Hair et al., 2017), con los problemas teóricos que conlleva.

Es relevante mencionar que la opción *bootstrapping* aplica solo para modelos con variables de tipo compuesto, pero si el modelo es una combinación de compuestos y de variables de factor común o solo de factor común, se deberá optar por *bootstrapping para PLS Consistente* (Sarstedt et al., 2016); de no hacerlo, los resultados obtenidos podrían estar sesgados (Rigdon et al., 2017). En ambos casos aparecerá una ventana semejante a la de la Figura 7.11 y, como se aprecia, ofrece una serie de alternativas, para configurar según las necesidades de cada investigación.

Para ejecutar el procedimiento *bootstrapping* y visualizar sus resultados se siguen los siguientes pasos: en el menú principal, escoger *Calcular* y luego *bootstrapping*, o *bootstrapping para PLS Consistente*, conforme a lo comentado anteriormente. En la pantalla que se despliega se recomienda introducir esta configuración: para la opción *submuestras* se pide utilizar un mínimo de 5000 y seleccionar *bootstrapping completo*, con el fin de mostrar todos los estadísticos requeridos. En *ajustes avanzados*

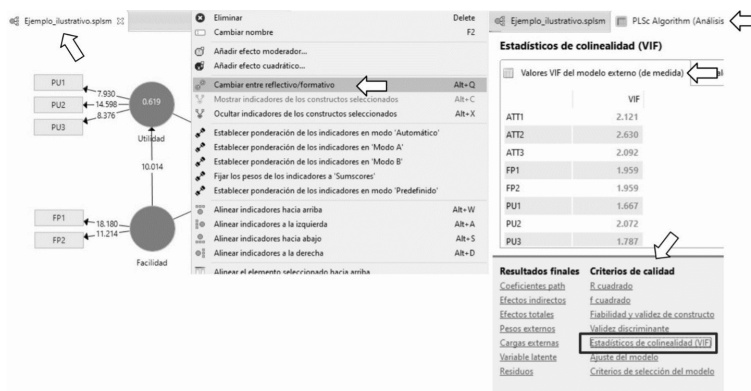
seleccionar *Bias-Corrected and Accelerated (BCa) Bootstrap* (Bootstrap con sesgo corregido y acelerado (BCa)). Para *Tipo de test* se opta por *Una cola* si se conoce la dirección de la relación (+ o -); en caso contrario, será de dos colas. Por último, se selecciona el nivel de significancia de los intervalos de confianza; el valor por *default* es de 0.05, o un 5% de nivel de significación. Luego, *iniciar cálculos*.

El resultado de ejecutar *bootstrapping* es el despliegue de la pestaña *bootstrapping*, o *bootstrapping c*, según la opción seleccionada. Esta pestaña se divide en dos partes, en una se presenta el detalle de los resultados de un indicador o criterio en particular (parte superior) y en la otra (parte inferior) se agrupan los resultados por categorías (resultados finales, criterios de calidad, ajustes del modelo, histogramas, base de datos).

Para mostrar los resultados para este punto, primero hay que ubicar la columna *Resultados finales* (parte inferior de la pestaña) y acto seguido *pesos externos*. Entonces, en la parte superior aparecen los resultados a detalle y en formato de tabla con el mismo título; lo que sigue es verificar si los valores se aceptan conforme lo previamente explicado para esta actividad, pero al igual que en los resultados del *bootstrapping*, el *software* marca en colores los resultados, para su mayor entendimiento, donde verde indica que se cumple.

Por último, en la figura 7.10 se muestra cómo cambiar un indicador de reflectivo a formativo en la ventana de modelización; también, la parte de localización de la opción para visualizar los valores de los pesos externos calculados por el procedimiento *bootstrapping* y, finalmente, la opción para localizar los valores VIF de los indicadores de tipo formativo, que se obtienen después de correr el Algoritmo PLS o el Algoritmo PLSc.

Figura 7.10. Ventana para configurar y obtener valores de indicadores formativos



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

7.4. Evaluación del modelo estructural

A partir de este punto, la orientación es explicar el proceso de evaluación básico de un modelo estructural en PLS, necesario para interpretar las hipótesis a comprobar. El procedimiento se compone de varios pasos o etapas: primero, determinar las medidas de ajuste del modelo, luego, valorar su colinealidad; enseguida, estimar la significancia y relevancia de las relaciones propuestas y, por último, definir los coeficientes de determinación. A continuación, se abordan estos pasos.

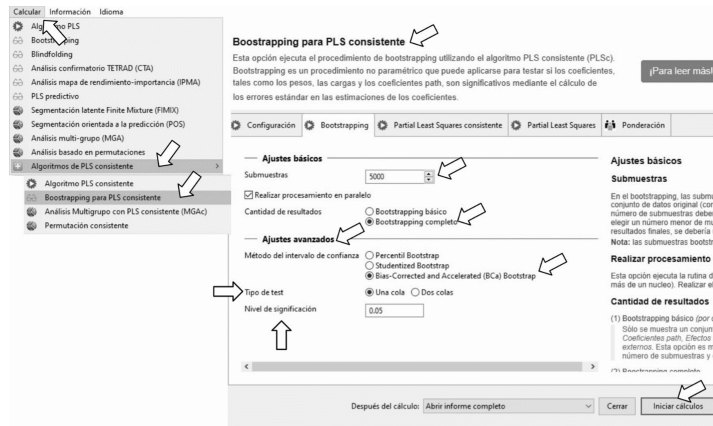
7.4.1. Evaluación del modelo global

Esta evaluación se efectúa como una pauta general para determinar si los datos obtenidos se ajustan al modelo hipotetizado y, por ende, identificar si existen especificaciones erróneas, es decir, medir la diferencia entre lo observado y lo predicho (Hair et al., 2012). Sin embargo, actualmente es un tema de debate la aplicación en la técnica PLS de los indicadores de ajuste propuestos, dado que se encuentran en etapas iniciales, y por el hecho de que PLS-SEM fue diseñado originalmente para fines de predicción, por lo que se recomienda tener reserva al momento de informar el resultado (Henseler, 2018).

SmartPLS ofrece una serie de medidas de ajuste, entre ellas el estadístico SRMR o Residual cuadrático medio estandarizado de Henseler et al. (2014) y las medidas de ajuste exacto como d_{ULS} y d_G (distancia euclídea y geodésica) (Dijkstra y Henseler, 2015b), por lo que se expondrá desde un enfoque procedimental cómo se obtienen desde el *software*. Para ello, seleccionar *calcular* en el menú principal y luego *bootstrapping*. Hay que recordar que SmartPLS aplica este procedimiento para ofrecer, entre otros indicadores, los relacionados con el error estándar, los estadísticos t , los intervalos de confianza de los parámetros —que sirven para mostrar información adicional sobre la estabilidad de los valores de SRMR, d_{ULS} y d_G (Henseler, Hubona y Ray, 2016a)— y otros que en su conjunto permitirán probar las hipótesis planteadas.

Es preciso recordar que la opción *bootstrapping* aplica solo para modelos con variables de tipo compuesto, pero si un modelo es una combinación de compuestos y de variables de factor común o solo de factor común, se deberá seleccionar *bootstrapping para PLS Consistente* (Sarstedt et al., 2016); de no ser así, los resultados obtenidos podrían estar sesgados (Rigdon et al., 2017). En ambos casos aparecerá una ventana semejante a la de la Figura 7.11 que, como se observa, ofrece una serie de alternativas para su configuración, de acuerdo con las necesidades de cada investigación.

Figura 7.11. Ventana de configuración del procedimiento Bootstrapping



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Para el ejemplo ilustrativo de este capítulo, el modelo analizado cuenta solo con variables de factor común, por lo que se aplicará la opción *bootstrapping para PLS Consistente*, con la siguiente configuración (Figura 7.11): en *submuestras* se pide utilizar un mínimo de 5000, ya que es la cantidad recomendada para presentar resultados finales de una investigación (Hair et al., 2014). Asimismo, seleccionar *bootstrapping completo*, para que se muestren todos los estadísticos requeridos. En *ajustes avanzados* seleccionar *Bias-Corrected and Accelerated (BCa) Bootstrap* (Bootstrap con sesgo corregido y acelerado (BCa)).

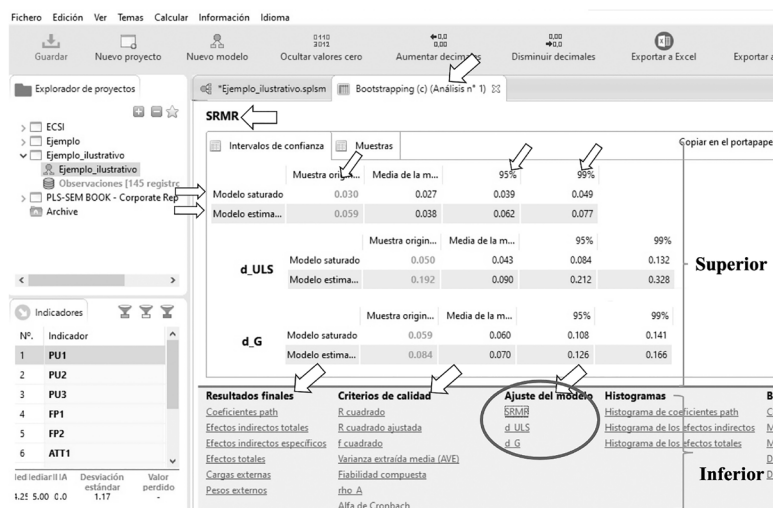
Para *Tipo de test* se escoge *Una cola*, ya que si se emplean hipótesis que especifican la dirección de la relación (+ o -), la distribución *t Student* será de 1 cola con $n-1$ grados de libertad (en el ejemplo ilustrativo todas las relaciones son positivas). Luego, se selecciona el nivel de significancia de los intervalos de confianza; el valor por *default* es de 0.05 o un 5% de nivel de significación. Finalmente, ir al botón *iniciar cálculos*.

El resultado del procedimiento *bootstrapping* se muestra en la Figura 7.12 y, como se nota, se despliega la pestaña *bootstrapping*, o *bootstrapping c*, según la opción seleccionada. Ésta se divide en dos partes, en una se presenta el detalle de los resultados de un indicador en particular (parte superior) y en la otra (parte inferior) se agrupan por categorías (resultados finales, criterios de calidad, ajustes del modelo, histogramas, base de datos). Para visualizar los resultados de los índices de ajuste proporcionados, primero, ubicar la columna *ajustes del modelo* (parte inferior de la pestaña) y, como se advierte, el *software* SmartPLS ofrece tres indicadores.

Ahora bien, para determinar si los criterios d_{ULS} y d_G cumplen con lo sugerido por la literatura hay que hacer clic en cada uno; los resultados aparecen en la parte superior en una subpestaña nombrada según el criterio seleccionado, en formato de tabla. Se verifica el valor estimado de la columna *muestra original*, tanto para el modelo saturado como para el estimado, y se compara si son inferiores a los valores de los intervalos establecidos (columnas denominadas 95% y 99%). Si se cumple en ambas columnas, se puede concluir que el modelo es verdadero y, por consiguiente, no puede ser rechazado (Dijkstra y Henseler, 2015a; Henseler et al., 2016b). De cualquiera manera, el *software* muestra los resultados en colores: el verde significa que se coincide con los criterios establecidos, el negro, que están en el límite sugerido y el rojo, que no son aceptables.

La estimación del índice SRMR es similar a lo expuesto anteriormente (comparar los valores estimados de los intervalos de confianza con los valores obtenidos tanto del modelo saturado como del estimado), pero con la salvedad de que un valor menor de SRMR, de 0.10 o 0.08 (versión más exigente) en ambos modelos (saturado y estimado) se considera un buen ajuste en general (Hu y Bentler, 1998; Williams et al., 2009). En el caso del ejemplo ilustrativo los dos modelos (saturado y estimado) cumplen (ver Figura 7.12), pero, como ya se mencionó, hay que tener cautela con lo obtenido para estos criterios, puesto que todavía están en etapas iniciales de comprobación. Continuando con la exposición del proceso de estimación del modelo estructural, ahora se abordará lo relacionado con la valoración de la colinealidad.

Figura 7.12. Valores estimados de los índices de ajuste



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

7.4.2. Valoración de la colinealidad

Para la valoración de los estadísticos de colinealidad (VIF) del modelo estructural se aplica la misma lógica que en la estimación de los modelos de medida formativos; para ello, es preciso inspeccionar la colinealidad en cada uno de los predictores del modelo estructural (Diamantopoulos y Siguaw, 2006). El límite establecido por la literatura para cada constructo exógeno deberá ser mayor a 0.20 pero menor a 5; de lo contrario, considerar su eliminación, fusión con el constructo donde presenta el inconveniente o creación de variables de orden superior (Hair et al., 2017).

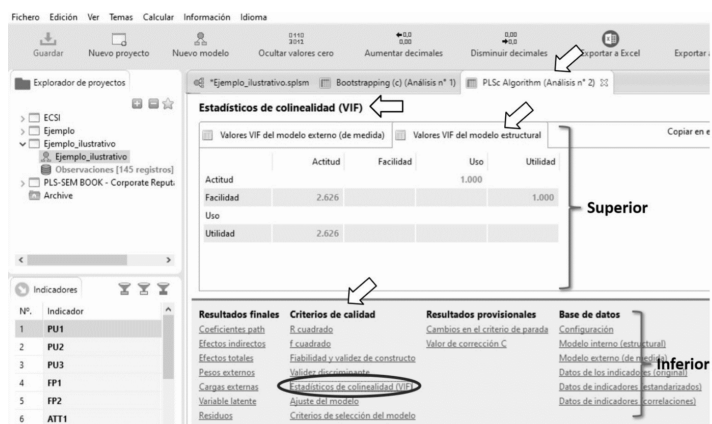
Para obtener los valores VIF se sigue el siguiente proceso: seleccionar *calcular* del menú principal y luego *Algoritmo PLS*, o *Algoritmo PLSc*, según sea el caso. En cualquiera de las dos opciones aparecerá una ventana similar a la que se muestra en la Figura 7.6. Dejar los valores por defecto e iniciar cálculos, lo que desplegará la pestaña *PLS algorithm* o *PLSc algorithm*, de acuerdo con la opción seleccionada previamente. Esto da como resultado una nueva pestaña, la cual está compuesta de dos partes (Figura 7.13), una en la que se muestran los resultados (parte superior) y en la otra (inferior) se agrupan por categorías (resultados finales, criterios de calidad, resultados provisionales, base de datos). Primero es necesario ubicar la columna *criterios de calidad* (parte inferior de la pestaña) y, como se aprecia, SmartPLS ofrece una serie de opciones; elegir *estadísticos de colinealidad (VIF)*.

En la parte superior aparecen los resultados a detalle y en formato de tabla con el mismo título (Figura 7.13), pero para determinar si los criterios se aceptan para el modelo estructural, ir a la subpestaña *Valores VIF del modelo estructural* y verificar que lo estimado para cada una de las variables esté en los rangos previamente indicados, aunque, al igual que en los resultados del *bootstrapping*, el *software* marca en colores lo obtenido, para su mayor entendimiento. En el ejemplo ilustrativo, los valores VIF para el modelo estructural consiguen lo establecido para este criterio.

7.4.3. Valoración de las rutas del modelo estructural

Continuando con la valoración del modelo estructural, ahora corresponde analizar la estimación de las rutas, relaciones o coeficientes *path* entre las variables del modelo estructural. Para ello, se requiere ejecutar el procedimiento de *bootstrapping*, que permitirá al *software* obtener los intervalos de confianza y el error estándar de las relaciones propuestas, lo que se traduce en el cálculo de los valores *t* empíricos y la significancia o valor *p*.

Figura 7.13. Resultados VIF del modelo estructural



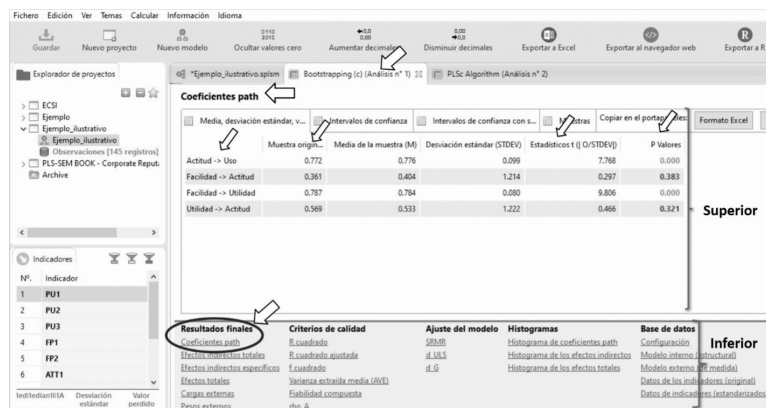
Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Para conseguirlo, ir a *calcular* del menú principal, luego a *bootstrapping*, o a *bootstrapping para PLS Consistente*, según sea el tipo de constructos por analizar. En cualquier caso, aparecerá una ventana semejante a la visualizada en la Figura 7.11, entonces: para *submuestras*, utilizar un mínimo de 5000 y la opción *bootstrapping completo*. En la parte que indica *ajustes avanzados*, seleccionar *Bias-Corrected and Accelerated (BCa) Bootstrap* (Bootstrap con sesgo corregido y acelerado (BCa)). Para la opción *Tipo de test*, se opta por *Una cola* si se conoce el sentido de la relación (+ o -); en caso contrario, será de *Dos colas*. Por último, en el nivel de significancia dejar el valor por *default* de 0.05 si éste aplica a la investigación. Una vez configurado, dar *clik* en *iniciar cálculos*.

En la Figura 7.14 se muestra la pantalla con los resultados y, como puede verse, se despliega la pestaña *bootstrapping*, o *bootstrapping c*, según la opción seleccionada. Ésta se divide en dos partes, en una se presenta el detalle de los resultados (parte superior) y en la otra se agrupan los resultados por categorías (resultados finales, criterios de calidad, ajustes del modelo, histogramas, base de datos). Hay que ubicar la columna *resultados finales* (parte inferior de la pestaña), posteriormente, elegir *coeficientes path*, con lo cual en la parte superior se mostrarán *media*, *desviación estándar*, *valores t* y *p valores*, con los resultados en forma de tabla.

Una vez localizada la tabla en la pestaña, prestar atención a los valores de las columnas *muestra original*, *estadísticos t* y *p valores*, dado que son los que se van a ocupar para testar las hipótesis. Para mayor facilidad, el SmartPLS marca en color verde aquellas relaciones que resultaron significativas y en rojo las que no, mientras que la columna *muestra original* señala el valor obtenido para la fuerza entre las relaciones trazadas (valores beta - β). La Figura 7.14 lo ilustra gráficamente.

Figura 7.14. Resultados de la magnitud y significación de los coeficientes path



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

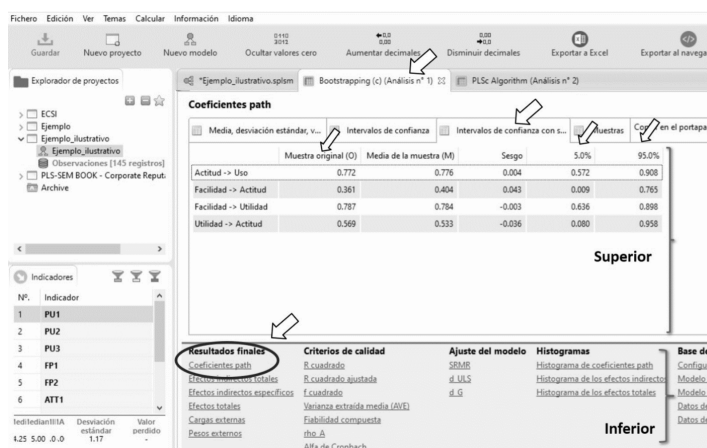
Para una adecuada interpretación de los datos es necesario recordar que los coeficientes *path* son valores estandarizados que indican la fuerza de las relaciones entre las variables dependientes e independientes, por lo que sus valores pueden oscilar entre -1 y +1; mientras más cercano a la unidad, más fuertes serán las relaciones y, al contrario, entre más cerca de 0, las relaciones serán más débiles (Johnson et al., 2006).

Se dice que existe significancia estadística en una relación si un valor *t* obtenido por una ruta *path* es mayor a un valor crítico (Hair et al., 2017). Lo anterior, con base en que un valor crítico puede ser medido a través de la prueba de una o dos colas. Para el caso de una cola, un valor *t* de 1.28 equivale a un nivel de significación de 10%, uno de 1.65, a 5%, y de 2.33, a 1%, mientras que para una de dos colas un valor *t* de 1.65 equivale a un nivel de significación de 10%, uno de 1.96, a 5%, y uno de 2.57, a 1%. Otro dato por tomar en cuenta es el análisis de la relevancia de las relaciones significativas, ya que permite una correcta interpretación de los resultados y, por ende, de conclusiones, ya que, según Chin (1998), para que se acepte como relevante un coeficiente *path* en una investigación debe alcanzar al menos un valor de 0.2.

Por último, se requiere comprobar que la estimación de los coeficientes *path* sea estable; para ello, se aplican los intervalos de confianza que proporciona la técnica *bootstrapping*, que se localiza en la subpestaña *intervalos de confianza con sesgo corregido*, en forma de tabla (ver Figura 7.15). Para comprobar lo obtenido, es preciso analizar que en las columnas 0.5% y 0.95% su intervalo estimado no incluya un cero, es decir, que no cambie de signo; de ser así, la hipótesis debe ser rechazada

(Henseler et al., 2009). Para el caso del ejemplo ilustrativo se puede notar que dos hipótesis se aceptan y dos se rechazan, lo que daría pie a la discusión de lo obtenido.

Figura 7.15. Coeficientes path e intervalos de confianza



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

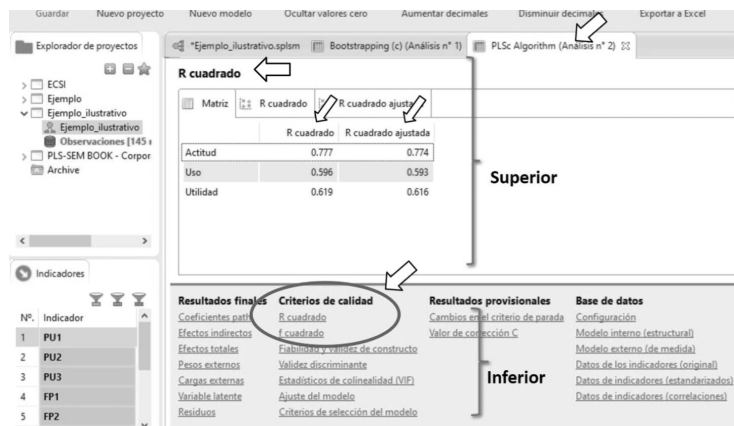
7.4.4. Estimación del coeficiente de determinación (R^2)

Este indicador permite estimar el poder predictivo del modelo y representa la cantidad de varianza de un constructo dependiente que es explicada por uno independiente. Su cálculo se basa en los valores de las correlaciones cuadradas del valor real y del predicho del constructo dependiente (Rigdon, 2012). Su valor varía de 0 a 1; por consiguiente, entre más cerca de 1, mayor será su nivel de predicción. Sin embargo, para estimar un valor R^2 no existe una guía que permita determinar si es aceptable o no, ya que será el área de estudio en cuestión la que establezca los parámetros; por ejemplo, Hair et al. (2011) consideran que para el área de marketing, valores de R^2 de 0.75, 0.50 y 0.25 se visualizan, respectivamente, como relevantes, moderados o débiles, mientras que Chin (1998) comenta que un valor de 0.67 se puede considerar como sustancial, un 0.33 como moderado y un 0.19 como débil.

Para obtener los valores R^2 de un modelo PLS el proceso es: seleccionar *calcular* del menú principal, luego *Algoritmo PLS*, o *Algoritmo PLS* (según sea el tipo de constructo por analizar). Una vez terminados los cálculos se despliega una ventana con una nueva pestaña que se divide en dos partes, una donde se muestran los resultados de un indicador y otra donde se agrupan por categorías, por ejemplo: resultados finales, criterios de calidad, resultados provisionales, base de datos. (Figura 7.16).

Enseguida, ubicar la columna *criterios de calidad* (parte inferior de la pestaña) y aparecerá una serie de opciones, elegir *R cuadrado* (Figura 7.16). Entonces, en la parte superior aparecen los resultados a detalle y en formato de tabla, con el título *R cuadrado*; para su análisis, seguir lo propuesto anteriormente en este mismo capítulo. El *software* también proporciona los resultados para, en caso de ser necesario, analizar los valores de R^2 ajustada, que controla la complejidad del modelo cuando existe un número diferente de variables latentes exógenas o conjuntos de datos con diferentes números de observaciones (Hair et al., 2014). En el caso del ejemplo ilustrativo, los valores R^2 obtenidos están en rangos de moderados a relevantes.

Figura 7.16. Ventana de resultados del coeficiente de terminación (R^2)



Fuente: Elaboración propia a partir del software SmartPLS v3.3.2.

Bibliografía

- Albaum, G. (1997). "The Likert scale revisited". *International Journal of Market Research*, 39(2), 1-21. <https://doi.org/10.1177/147078539703900202>
- Anderson, J. C., y Gerbing, D. W. (1988). "Structural equation modeling in practice: A review and recommended two-step approach". *Psychological bulletin*, 103(3), 411-423.
- Arbuckle, J. L. (2003). *AMOS user's guide*. SmallWaters
- Aron, A. y Aron, E. (2001). *Estadística para Psicología*. Bs.As. Pearson Education.
- Badii, M.H. y J. Castillo, J. (2007). *Técnicas Cuantitativas en la Investigación*. Universidad Autónoma de Nuevo León.
- Bagozzi, R. y Fornell, C. (1982). "Theoretical concepts, measurement, and meaning". En Fornell, C. (Ed.), *A Second Generation of Multivariate Analysis*, (pp. 5-21). Praeger Publishers.
- Bagozzi, R. (1994). "Structural equation models in marketing research: Basic principles". En Bagozzi, R.P. (Ed.), *Principles of Marketing Research*, (pp. 317-385). Blackwell.
- Barclay, D., Higgins, C. y Thompson, R. (1995). "The Partial Least Squares (PLS) approach to causal modeling: Personal computer adoption and use as an illustration". *Technology Studies. Special Issue on Research Methodology*, 2(2), 285-309.
- Baron, R. M. y Kenny, D. A. (1986). "The moderator–mediator variable distinction in social psychological research: Conceptual, strategic, and statistical considerations". *Journal of personality and social psychology*, 51(6), 1173-1182.
- Beltrán, L., y Espejel Blanco, J. E. (2016). "Análisis del estudio de las relaciones causales en el marketing". *Innovar Journal*, 26(62), 79-94. <https://doi.org/10.15446/innovar.v26n62.59390>.
- Bentler, P. (1995). *Structural Equations Program Manual*. Multivariate Software, Inc.
- Bollen, K. A. (1990). "Overall fit in covariance structure models: Two types of sample size effects". *Psychological bulletin*, 107(2), 256-259. <https://doi.org/10.1037/0033-2909.107.2.256>
- Bollen, K. A., y Pearl, J. (2013). "Eight myths about causality and structural equation models". En *Handbook of causal analysis for social research*, (pp. 301-328). Springer, Dordrecht.
- Bollen, K.A. y Davis, W.R. (2009). "Two rules of identification for structural equation models". *Structural Equation Modeling: A Multidisciplinary Journal*, 16, 523-533.
- Bollen, K.A. (1989). *Structural equations with latent variables*. John Wiley y Sons.
- Brace, I. (2018). *Questionnaire design: How to plan, structure and write survey material for effective market research*. Kogan Page Publishers.
- Bradley, J.V. (1978). Robustness? *British Journal of Mathematical and Statistical Psychology*, 31, 144-152.
- Brown, T. (2006). *Confirmatory factor analysis for applied research*. The Guilford Press.
- Byrne, B. (2001). *Structural equation modeling with AMOS: Basic concepts, applications, and programming*. Lawrence Erlbaum

- Byrne, B. (2016). *Structural equation modeling with AMOS Basic concepts, applications, and programming (Multivariate Applications Series)*. Routledge
- Byrne, B. M., Shavelson, R. J., y Muthén, B. (1989). “Testing for the equivalence of factor covariance and mean structures: the issue of partial measurement invariance”. *Psychological Bulletin*, 105(3), 456–466. <https://doi.org/10.1037/0033-2909.105.3.456>
- Cain, M.K., Zhang, Z. y Yuan, K. (2017). “Univariate and multivariate skewness and kurtosis for measuring nonnormality: Prevalence, influence and estimation”. *Behavior Research Methods*, 49, 1716–1735. <https://doi.org/10.3758/s13428-016-0814-1>
- Cameron, A. y Trivedi, P. (2009). *Microeconometrics using stata*. Stata Press.
- Campbell, D. T. y Fiske, D. W. (1959). “Convergent and discriminant validation by the multitrait-multimethod matrix”. *Psychological Bulletin*, 56, 81-105
- Campo-Arias, A., y Oviedo, H. C. (2008). “Propiedades psicométricas de una escala: la consistencia interna”. *Revista de Salud Pública*, 10, 831-839.
- Carifio, J., y Perla, R. J. (2007). “Ten common misunderstandings, misconceptions, persistent myths and urban legends about Likert scales and Likert response formats and their antidotes”. *Journal of Social Sciences*, 3(3), 106-116.
- Castro, M. C. B., Carrión, G. A. C., y Roldán, J. L. (2007). “Investigar en Economía de la Empresa ¿Partial Least Squares o modelos basados en la covarianza? En el comportamiento de la empresa ante entornos dinámicos”. *XIX Congreso anual y XV Congreso Hispano Francés de AEDEM* (p. 63). Asociación Española de Dirección y Economía de la Empresa (AEDEM).
- Cepeda, G. y Roldán, J. (2004). “Aplicando en la práctica la técnica PLS en la administración de empresas”. *Congreso de la ACEDE*, sep. 19, 20 y 21, Murcia, España.
- Cervantes, V. H. (2005). “Interpretaciones del coeficiente alpha de Cronbach”. *Avances en medición*, 3(1), 9-28.
- Chin, W. y Todd, P. (1995). “On the use, usefulness, and ease of use of Structural Equation Modeling in MIS research: A note of caution”. *MIS Quarterly*, 19(2), 237-246. <https://doi.org/10.2307/249690>.
- Chin, W. (1998). “Issues and Opinion on Structural Equation Modeling”. *MIS Quarterly*, 22(1), VII-XVI. <http://www.jstor.org/stable/249674>
- Chin, W. (2010). “How to write up and report PLS analyses”. En: Esposito Vinzi V., Chin W., Henseler J., Wang H. (Eds.), *Handbook of Partial Least Squares*. Springer Handbooks of Computational Statistics. Springer. https://doi.org/10.1007/978-3-540-32827-8_29.
- Chin, W., Marcolin, B. y Newsted, P. (2003). “A Partial Least Squares latent variable modeling approach for measuring interaction effects: Results from a Monte Carlo Simulation study and an electronic-mail emotion / adoption study”. *Information Systems Research*, 14(2), 189-217. <https://doi.org/10.1287/isre.14.2.189.16018>.

- Cho, E. (2016). "Making Reliability Reliable: A Systematic approach to reliability coefficients". *Organizational Research Methods*, 19(4), 651-682. <https://doi.org/10.1177/1094428116656239>
- Chomeya, R. (2010). "Quality of psychology test between Likert scale 5 and 6 points". *Journal of Social Sciences*, 6(3), 399-403.
- Clark, L. A., y Watson, D. (1995). "Constructing validity: basic issues in objective scale development". *Psychological Assessment*, 7(3), 309-319.
- Clavijo, F. y Gómez, O. (1979). "Parámetros e interdependencias en la economía mexicana. Un análisis econométrico". *El Trimestre Económico*, 46(182), 311-376.
- Cole, D. A., y Maxwell, S. E. (1985). "Multitrait-multimethod comparisons across populations: A confirmatory factor analytic approach". *Multivariate Behavioral Research*, 20(4), 389-417.
- Cronbach, L.J. (1949, 1960, 1970, 1984, 1990). *Essentials of psychological testing* (1ª-5ª edición). Harper.
- Cupani, M. (2012). "Análisis de Ecuaciones Estructurales: conceptos, etapas de desarrollo y un ejemplo de aplicación". *Revista tesis*, 2, 186-199.
- DeCarlo, L. (1997). "On the meaning and use of kurtosis. On the meaning and use of kurtosis". *Psychological methods*, 2(3), 292-307. <https://doi.org/10.1037/1082-989X.2.3.292>.
- Diamantopoulos, A. y Siguaw, J. (2006). "Formative versus reflective indicators in organizational measure development: A comparison and empirical illustration". *British Journal of Management*, 17, 263-282. <https://doi.org/10.1111/j.1467-8551.2006.00500.x>.
- Diamantopoulos, A. y Winklhofer, H. (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.38.2.269.18845>.
- Diamantopoulos, A. (2008). "Formative indicators: Introduction to the special issue". *Journal of Business Research*, 61(12), 1201-1202. <https://doi.org/10.1016/j.jbusres.2008.01.008>
- Díaz, M. (2000). "Introducción a los modelos de ecuaciones estructurales". Publicaciones del INICO, 43.
- Dijkstra, T. y 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>.
- Dijkstra, T. y Henseler, J. (2015b). "Consistent and asymptotically normal PLS estimators for linear structural equations". *Computational Statistics y Data Analysis*, 81, 10-23. <https://doi.org/10.1016/j.csda.2014.07.008>
- Ding, C. y Hershberger, S. (2002). "Assessing content validity and content equivalence using structural equation modeling. Structural Equation Modeling". *A Multidisciplinary Journal*, 9 (2), 283-297.

- Ding, L., Velicer, W. F., y Harlow, L. L. (1995). "Effects of estimation methods, number of indicators per factor, and improper solutions on structural equation modeling fit indices. Structural Equation Modeling". *A Multidisciplinary Journal*, 2(2), 119-143.
- Dunn, T. J., Baguley, T., y Brunsden, V. (2013). "From alpha to omega: A practical solution to the pervasive problem of internal consistency estimation". *British Journal of Psychology*, 105(3), 399-412. <https://doi.org/10.1111/bjop.12046>
- Efron, B. (1981). "Censored data and the bootstrap". *Journal of the American Statistical Association*, 76(374), 312-319.
- Escobedo, M., Hernández, J., Estebané, V. y Martínez, G. (2016). "Modelos de ecuaciones estructurales: Características, fases, construcción, aplicación y resultados". *Ciencia y trabajo*, 18(55), 16-22.
- Everitt, B. y Dunn, G. (1991). *Applied multivariate data analysis*. John Wiley y Sons.
- Fabrigar, L. R., Wegener, D. T., MacCallum, R. C., y Strahan, E. J. (1999). "Evaluating the use of exploratory factor analysis in psychological research". *Psychological Methods*, 4(3), 272-299.
- Falk, R. y Miller, N. (1992). *A primer for soft modeling*. Akron, OH: The University of Akron Press, USA.
- Faul, F., Erdfelder, E., Buchner, A. y Lang, A. (2009). "Statistical power analyses using G*Power 3.1: Tests for correlation and regression analyses". *Behavior Research Methods*, 41(4), 1149-1160. <https://doi.org/10.3758/brm.41.4.1149>.
- Finney, S. J., y DiStefano, C. (2006). "Non-normal and categorical data in structural equation modeling. Structural equation modeling": *A second course*, 10(6), 269-314.
- Fornell, C. y Bookstein, F. (1982). "Two structural equation models: LISREL and PLS applied to Consumer Exit-Voice Theory". *Journal of Marketing Research*, 19(4), 440-445. <https://doi.org/10.2307/3151718>.
- Fornell, C. y Larcker, D. (1981). "Evaluating structural equation models with unobservable variables and measurement error". *Journal of Marketing Research*, 18(1), 39-50. <https://doi.org/10.2307/3151312>.
- Fornell, C. (1982). "A second generation of multivariate analysis: An overview". En Fornell, C. (Ed.), *A Second Generation of Multivariate Analysis*, (pp. 1-5). Praeger Publishers.
- Freiberg Hoffmann, A., Stover, J. B., de la Iglesia, G., y Fernández Liporace, M. (2013). "Correlaciones policóricas y tetracóricas en estudios factoriales exploratorios y confirmatorios". *Ciencias psicológicas*, 7(2), 151-164.
- García-Ferrer, A. (2008). *Causalidad y Econometría*. Edición de Juan Carlos García-Bermejo, 225.
- Gaskin, J. y James, M. (2019). HTMT Plugin for AMOS.
- Gaskin, J. y Lim, J. (2016). *Master Validity Tool*. AMOS Plugin.

- Gefen, D., Straub, D. y Boudreau, M. (2000). "Structural Equation Modelling and regression: Guidelines for research practice". *Communications of the Association for Information Systems*, 4(7), 1-76. <https://doi.org/10.17705/1cais.00407>.
- George, D. y Mallery, M. (2003). *Using SPSS for Windows Step by Step: a simple guide and reference*. Allyn y Bacon.
- Gold, A. H., Malhotra, A., y Segars, A. H. (2001). "Knowledge management: an organizational capabilities perspective". *Journal of Management Information Systems*, 18(1), 185-214.
- Graham, J. M. (2006). "Congeneric and (essentially) tau-equivalent estimates of score reliability –What they are and how to use them". *Educational and Psychological Measurement*, 66, 930-944. <https://doi.org/10.1177/0013164406288165>
- Graham, J.W. (2009). "Missing data analysis: making it work in the real world". *Annual Review of Psychology*, 60, 549-76.
- Green, S. (1991). "How many subjects does it take to do a regression analysis?" *Multivariate Behavioral Research*, 26(3), 499-510. https://doi.org/10.1207/s15327906mbr2603_7.
- Gujarati, D. y Porter, D. (2010). *Econometria*. Mc Graw Hill.
- Hair J., Black, W., Babin, B. y Anderson, R. (2014). *Multivariate data analysis* (7a ed.). Pearson Education Limited.
- Hair, J., Anderson, R., Tatham, R. y Black, W. (1999). *Análisis multivariante*. Prentice Hall.
- Hair, J., Hult, G., Ringle, C. y Sarstedt, M. (2017). *A primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*. SAGE Publications.
- Hair, J., Hult, T., Ringle, C., Sarstedt, M., Castillo, J., Cepeda, G. y Roldán, J. (2019). *Manual de Partial Least Squares Structural Equation Modeling (PLS-SEM)*. SAGE Publishing.
- Hair, J., Ringle, C. y Sarstedt, M. (2011). "PLS-SEM: Indeed, a silver bullet". *Journal of Marketing Theory and Practice*, 19(2), 139-15. <https://doi.org/10.2753/MTP1069-6679190202>.
- Hair, J., Sarstedt, M., Hopkins, L. y Kuppelwieser, V. (2014). "Partial least squares structural equation modeling (PLS-SEM): An emerging tool in business research". *European Business Review*, 26(2), 106-121. <https://doi.org/10.1108/EBR-10-2013-0128>
- Hair, J., Sarstedt, M., Pieper, T. y Ringle, C. (2012). "The use of partial least squares structural equation modeling in strategic management research: A review of past practices and recommendations for future applications". *Long Range Planning: International Journal of Strategic Management*, 45(5-6), 320-340. <https://doi.org/10.1016/j.lrp.2012.09.008>
- Hair, J., Sarstedt, M., Ringle, C. y Gudergan, S. (2018). *Advanced issues in partial least squares structural equation modeling*. SAGE publications.
- Heise, D. R., y Bohrnstedt, G. W. (1970). "Validity, invalidity, and reliability". *Sociological methodology*, 2, 104-129.

- Henseler, J. (2017). "Bridging design and behavioral research with Variance-Based Structural Equation Modeling". *Journal of Advertising*, 46(1), 178-192. <https://doi.org/10.1080/00913367.2017.1281780>.
- Henseler, J. (2018). "Partial least squares path modeling: Quo vadis?" *Quality y Quantity*, 52(1), 1-8. <https://doi.org/10.1007/s11135-018-0689-6>
- Henseler, J., Dijkstra, T., Sarstedt, M., Ringle, C. M., Diamantopoulos, A., Straub, D., Ketchen, D., Hair, J., Hult, G., y Calantone, R. J. (2014). "Common Beliefs and Reality about Partial Least Squares: Comments on Rönkkö y Evermann" (2013). *Organizational Research Methods*, 17(2) 182-209. <https://doi.org/10.1177/1094428114526928>
- Henseler, J., Hubona, G. y Ray, P. (2016b). "Using PLS path modeling in new technology research: updated guidelines". *Industrial Management y Data Systems*, 116 (1), 2-20. <https://doi.org/10.1108/IMDS-09-2015-0382>
- Henseler, J., Ringle, C. y 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://doi.org/10.1007/s11747-014-0403-8>.
- Henseler, J., Ringle, C. y Sarstedt, M. (2016a). "Testing measurement invariance of composites using Partial Least Squares". *International Marketing Review*, 33(3), 405-431. <https://doi.org/10.1108/imr-09-2014-0304>.
- Henseler, J., Ringle, C. y Sinkovics, R. (2009). "The use of partial least squares path modeling in international marketing. In New challenges to international marketing". *Emerald Group Publishing Limited*, 277-319. [https://doi.org/10.1108/S1474-7979\(2009\)0000020014](https://doi.org/10.1108/S1474-7979(2009)0000020014)
- Hoelter, J. W. (1983). "The analysis of covariance structures: Goodness-of-fit indices". *Sociological Methods y Research*, 11(3), 325-344.
- Hooper, D., Coughlan, J. y Mullen, M. (2008). "Structural Equation Modelling: Guidelines for Determining Model Fit". *Electronic Journal of Business Research Methods*, 6(1), 53–60.
- Hu, L. y Bentler, P. (1998). "Fit indices in covariance structure modeling: Sensitivity to underparameterized model misspecification". *Psychological Methods*, 3(4), 424–453. <https://doi.org/10.1037/1082-989X.3.4.424>
- Hu, L. y Bentler, P. (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>.
- Hu, L. T., Bentler, P. M., y Kano, Y. (1992). "Can test statistics in covariance structure analysis be trusted?" *Psychological bulletin*, 112(2), 351.
- Huysamen, G. (2007). "Coefficient alpha: Unnecessarily ambiguous; unduly ubiquitous. SA". *Journal of Industrial Psychology*, 32, 34–40.

- Im, K. y Grover, V. (2004). "The use of Structural Equation Modeling in IS research: Review and recommendations". En Whitman, M.E. y A.B. Woszczyński (Eds.): *The Handbook of Information Systems Research*, (pp. 44-63). Idea Group Publishing. <https://doi.org/10.4018/978-1-59140-144-5.ch004>.
- Johnson, M., Herrmann, A. y Huber, F. (2006). "The evolution of loyalty intentions". *Journal of Marketing*, 70(2), 122-132. <https://doi.org/10.1509/jmkg.70.2.122>
- Jöreskog, K. G., y Sörbom, D. (1982). "Recent developments in structural equation modeling". *Journal of marketing research*, 19(4), 404-416.
- Jöreskog, K. G., y Wold, H. O. (1982). *Systems under indirect observation: Causality, structure, prediction*. North Holland.
- Joshi, A., Kale, S., Chandel, S. y Pal, D. (2015). "Likert scale: Explored and explained". *Current Journal of Applied Science and Technology*, 7(4), 396-403. <https://doi.org/10.9734/BJAST/2015/14975>.
- Kahn, J. H. (2006). "actor analysis in Counseling Psychology research, training, and practice": Principles, advances and applications. *The Counseling Psychologist*, 34, 1-36.
- Klein, R. (2011). *Principles and practice of Structural Equation Modeling* (3a ed.) The Guilford Press, New York, NY.
- Knapp, T.R., 1990. "Treating ordinal scales as interval scales: an attempt to resolve the controversy". *Nursing Research*, 39, 121-123.
- Kronmal, R. A. (1993). "Spurious correlation and the fallacy of the ratio standard revisited". *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 156(3), 379-392.
- Kutner, M., Nachtsheim, C. y Neter, J. (2004). *Applied Linear Regression Models* (4a ed.). McGraw-Hill Irwin.
- Lahura, E. (2003). *El coeficiente de correlación y correlaciones espúreas* (Vol. 218). Pontificia Universidad Católica del Perú, Departamento de Economía.
- Lawshe C. H. (1975). "A quantitative approach to content validity". *Personnel Psychology*. 28(4), 563-575.
- Lee, J., Jones, P., Mineyama, Y. y Zhang, X. (2002). "Cultural differences in responses to a Likert scale". *Research in Nursing y Health*, 25(4), 295-306.
- Lévy-Mangin, J. P., y Varela, J. (2006). *Modelización con estructuras de covarianzas en ciencias sociales. Temas esenciales, avanzados y aportaciones especiales*. Netbiblo.
- Likert, R. (1932). "A technique for the measurement of attitudes". *Archives of Psychology*, 140, 44-53.
- Loehlin, J. C. (1987). *Latent variable models: An introduction to factor, path, and structural analysis*. Lawrence Erlbaum Associates, Inc.
- Lohmöller, J. (1989). *Latent variable path modeling with Partial Least Squares*. Springer.
- Lomax, R. G., y Algina, J. (1979). "Comparison of two procedures for analyzing multitrait multimethod matrices". *Journal of Educational Measurement*, 177-186.

- Manzano, P. (2017). "Introducción a los modelos de ecuaciones estructurales". *Investigación en Educación Médica*, 7(25). <https://doi.org/10.1016/j.riem.2017.11.002>
- Mathieson, K., Peacock, E. y Chin, W. (2001). "Extending the Technology Acceptance Model: The influence of perceived user resources". *The DATA BASE for Advances in Information Systems*, 32(3), 86-112. <https://doi.org/10.1145/506724.506730>.
- McDonald, R. P., y Ho, M. H. R. (2002). "Principles and practice in reporting structural equation analyses". *Psychological methods*, 7(1), 64.
- Medrano, L. A. y Muñoz-Navarro, R. (2017). "Aproximación Conceptual y Práctica a los Modelos de Ecuaciones Estructurales". *Revista Digital de Investigación en Docencia Universitaria*, 11(1), 219-239. <http://dx.doi.org/10.19083/ridu.11.486>
- Miles, J., y Shevlin, M. (2007). "A time and a place for incremental fit indices". *Personality and Individual Differences*, 42(5), 869-874. <http://dx.doi.org/10.1016/j.paid.2006.09.022>
- Miller, B. (2001). *Beyond statistics: A practical guide to data analysis*. Allyn y Bacon.
- Morales, P., y Rodríguez, L. (2016). "Aplicación de los coeficientes correlación de Kendall y Spearman". Universidad Centroccidental Lisandro Alvarado.
- Mulaik, S. A., y Millsap, R. E. (2000). "Doing the four-step right". *Structural equation modeling*, 7(1), 36-73.
- Mulaik, S. A., James, L. R., Van Alstine, J., Bennett, N., Lind, S., y Stilwell, C. D. (1989). "Evaluation of goodness-of-fit indices for structural equation models". *Psychological bulletin*, 105(3), 430.
- Muthén, B. O. (1989). "Latent variable modeling in heterogeneous populations." *Psychometrika*, 54(4), 557-585.
- Muthén, B. O. (1993). "Goodness of Fit with Categorical". En Bollen, K. y Long, S. (Eds.), *Testing structural equation models*, (pp. 154-169). Sage Publications
- Nunnally, J. y Bernstein, I. (1994). *Psychometric theory* (3rd ed.). McGraw-Hill.
- Nunnally, J. (1978). *Psychometric Theory*. McGraw Hill Editorial.
- Pell, G. (2005), "Uses and misuses of Likert scales", *Medical Education*, 39(9), 970.
- Pérez, C. (2004). *Técnicas de análisis multivariante de datos*. Pearson Educación.
- Pérez, E. R., y Medrano, L. A. (2010). "Análisis factorial exploratorio: bases conceptuales y metodológicas". *Revista Argentina de Ciencias del Comportamiento (RACC)*, 2(1), 58-66.
- Pérez, E., Medrano, L. A., y Rosas, J. S. (2013). "El Path Analysis: conceptos básicos y ejemplos de aplicación". *Revista Argentina de Ciencias del Comportamiento*, 5(1), 52-66.
- Podsakoff, P., MacKenzie, S., Lee, J. y Podsakoff, N. (2003). "Common Method Biases in Behavioral Research: A Critical Review of the Literature and Recommended Remedies". *Journal of Applied Psychology*, 88(5), 879-903. <https://doi.org/10.1037/0021-9010.88.5.879>
- Raykov, T. y Penev, S. (2010). "Testing multivariate mean collinearity via latent variable modeling". *The British Journal of Mathematical and Statistical Psychology*, 63, 481-490.

- Rigdon, E. (2012). "Rethinking partial least squares path modeling: In praise of simple methods". *Long range planning*, 45(5-6), 341-358. <https://doi.org/10.1016/j.lrp.2012.09.010>
- Rigdon, E. E. (1995). "A necessary and sufficient identification rule for structural models estimated in practice". *Multivariate behavioral research*, 30(3), 359-383.
- Rigdon, E., Sarstedt, M. y Ringle, C. (2017). "On comparing results from CB-SEM and PLS-SEM: Five perspectives and five recommendations". *Marketing Zfj*, 39(3), 4-16. <https://doi.org/10.15358/0344-1369-2017-3-4>
- Ringle, C., Wende, S. y Becker, J. (2015). *SmartPLS version 3.3.2*. <http://www.smartpls.com>
- Roberts, N. y Thatcher, J. (2009). "Conceptualizing and testing formative constructs: Tutorial and annotated example". *ACM SIGMIS Database: the DATABASE for Advances in Information Systems*, 40(3), 9-39. <https://doi.org/10.1145/1592401.1592405>.
- Roldán, J. y Cepeda, G. (2017). *Modelos de ecuaciones estructurales basados en la varianza: Partial Least Squares (PLS) para investigadores en ciencias sociales*. Centro de Formación Permanente, Universidad de Sevilla.
- Rubin, D.B. (1976). "Inference and Missing Data". *Biometrika*, 63(3), 581-592.
- Ruiz, M. A., Pardo, A., y San Martín, R. (2010). "Modelos de ecuaciones estructurales". *Papeles del Psicólogo*, 31(1), 34-45.
- Sarstedt, M., Hair, J., Ringle, C., Thiele, K. y Gudergan, S. (2016). "Estimation issues with PLS and CBSEM: Where the bias lies!". *Journal of Business Research*, 69(10), 3998-4010. <https://doi.org/10.1016/j.jbusres.2016.06.007>
- Schermelleh-Engel, K., Moosbrugger, H., y Müller, H. (2003). "Evaluating the fit of structural equation models: Tests of significance and descriptive goodness-of-fit measures". *Methods of psychological research online*, 8(2), 23-74.
- Schreiber, J. B., Nora, A., Stage, F. K., Barlow, E. A. and King, J. (2006). "Reporting Structural Equation Modeling and Confirmatory Factor Analysis Results: A Review". *The Journal of Educational Research*, 99(6), 323-338. <https://doi.org/10.3200/JOER.99.6.323-338>
- Schumacker, R. E., y Lomax, R. G. (2004). *A beginner's guide to structural equation modeling*. Psychology press.
- Sellin, N. (1995). "Partial Least Square modeling in research on educational achievement". En Bos, W. y R.H. Lehmann (Eds.) *Reflections on Educational Achievement*, (pp. 256-267). Waxmann Munster.
- Shmueli, G., Ray, S., Velasquez, J. y Chatla, S. (2016). "The elephant in the room: Evaluating the predictive performance of PLS models". *Journal of Business Research*, 69(10), 4552-4564. <https://doi.org/10.2139/ssrn.2659233>.
- Shmueli, G., Sarstedt, M., Hair, J., Cheah, J., Ting, H., Vaithilingam, S. y Ringle, C. (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>.

- Simon, H. (1988). "Causal ordering, comparative statics and near decomposability". *Journal of Econometrics*, 39, 149-173.
- Simon, H. A. (1954). "Spurious correlation: A causal interpretation". *Journal of the American statistical Association*, 49(267), 467-479.
- Solís, S., Zerón, M. y Sánchez, Y. (2019). "Effects of the Absorption Capacity in the Innovation of the Industrial Sector in the North of Mexico". *Nova Scientia*, 11(23), 447 - 472. <https://doi.org/10.21640/ns.v11i23.2039>
- Sosik, J., Kahai, S. y Piovoso, M. (2009). "Silver bullet or voodoo statistics? A primer for using the Partial Least Squares data analytic technique in group and organization research". *Group y Organization Management*, 34(1), 5-36. <https://doi.org/10.1177/1059601108329198>.
- Triola, M. (2004). *Probabilidad y estadística*. Pearson Educación.
- Tristán-López, A. (2008). "Modificación al modelo de Lawshe para el dictamen cuantitativo de la validez de contenido de un instrumento objetivo". *Avances en Medición*, 6(1), 37-48.
- Urbach, N. y Ahlemann, F. (2010). "Structural Equation Modeling in information systems research using Partial Least Squares". *Journal of Information Technology Theory and Application*, 11(2), 5-40.
- Véliz, C. (2016). *Análisis multivariante. Métodos estadísticos multivariantes para la investigación*. CENGAGE Learning Editores SA.
- Wall, M.M., Guo, J. y Amemiya, Y. (2012). "Mixture factor analysis for approximating a nonnormally distributed continuous latent factor with continuous and dichotomous observed variables". *Multivariate Behavioral Research*, 47(2), 276- 313.
- Williams, L., Vandenberg, R. y Edwards, J. (2009). Structural equation modeling in management research: A guide for improved analysis. *Academy of Management Annals*, 3(1), 543-604. <https://doi.org/10.5465/19416520903065683>
- Wold, H. (1980). "Model Construction and Evaluation when Theoretical Knowledge is Scarce: An Example of the Use of Partial Least Squares". En: Kmenta, J. y J.B. Ramsey (Eds.), *Evaluation of Econometric Models* (pp. 47-74). Elsevier, Academic Press. <https://doi.org/10.1016/b978-0-12-416550-2.50007-8>.
- Wold, H. (1982). "Soft modeling: The basic design and some extensions". En Jöreskog, K.G. and H. Wold (Eds.), *Systems under indirect observations: Part II*. (pp. 1-54) North-Holland.
- Wold, H. (1985). "Partial Least Squares". En: Kotz, S. y N.L. Johnson (Eds.), *Encyclopedia of Statistical Sciences*. Vol. 6. (pp. 581-591). Wiley Editorial.
- Wooldridge, J. M. (2016). *Introductory econometrics: A modern approach*. Nelson Education.
- Zikmund, W., Babin, B., Carr, J. y Griffin, M. (2003). *Business Research Methods* (9a Ed.). CENAGE Learning.

Técnicas de análisis multivariante aplicadas a las ciencias sociales Volumen I, de Demian Ábrego Almazán, José Melchor Medina Quintero, Yesenia Sánchez Tovar y Manuel H. de la Garza Cárdenas, publicado por la Universidad Autónoma de Tamaulipas y Colofón, se terminó de imprimir en julio de 2021 en los talleres de Ultradigital Press S.A. de C.V. Centeno 195, Col. Valle del Sur, C.P. 09819, Ciudad de México. El tiraje consta de 400 ejemplares impresos de forma digital en papel Cultural de 75 gramos. El cuidado estuvo a cargo del Consejo de Publicaciones UAT.

Las técnicas estadísticas han evolucionado en la investigación científica; en el ayer quedaron las herramientas como análisis de regresión, ANOVA, chi-cuadrada y otras. Hoy las necesidades son distintas, se requiere de análisis multivariante que permita crear una red entramada de variables de las que son codependientes. Surge con ello la segunda generación y la aparición del Modelado de Ecuaciones Estructurales (SEM, por sus siglas en inglés de Structural Equation Modeling), que es una herramienta que identifica más de una relación entre variables y su análisis inferencial se realiza en forma simultánea. AMOS y PLS son dos técnicas progresistas de SEM que han sobresalido en la investigación de alto impacto, ya que sus algoritmos, desde una sola corrida de datos y aunque no exista una conexión directa, obtienen los resultados de las conexiones establecidas por el investigador.

Esta obra se enfoca esencialmente a temas de las ciencias sociales, del comportamiento y administrativas, y tomando en cuenta que los trabajos empíricos llevados a cabo con AMOS y PLS destacan por su alto valor académico, sus aportaciones al conocimiento y un sustento robusto que combina la parte teórica con la inferencia estadística; precisamente por ello su estudio y análisis.

Sin duda, este documento es una guía para investigadores y practicantes noveles y expertos que les proporcionará la oportunidad de evaluar sus trabajos de investigación e incursionar en las publicaciones de alta calidad, ya que a través de su lectura se conoce la teoría, para aplicarlo en un ejemplo ilustrativo que lleva al lector desde el diseño de un modelo gráfico con AMOS y PLS hasta el desarrollo robusto de un análisis estadístico.

Publicación financiada con recurso PROFEXCE 2020

ISBN UAT: 978-607-8750-60-3

ISBN Obra Completa: 978-607-8750-59-7

ISBN Colofón: 978-607-635-236-6

ISBN Obra Completa: 978-607-635-235-9

