

Modelo de Ecuaciones Estructurales: Una guía para ciencias médicas y ciencias de la salud

Structural Equation modeling: A guide for Medical and Health sciences

Manuel S. Ortiz

Departamento de Psicología.
Laboratorio de Estrés y Salud.
Universidad de La Frontera, Chile

Montserrat Fernández-Pera

Programa de Doctorado en Psicología.
Universidad de La Frontera, Chile.
Vice Gran Cancillería.
Universidad Católica de Temuco.

Recibido (09/01/2017) Aceptado (25/05/2017)

Correspondencia: La correspondencia relativa a este artículo debe ser enviada a Manuel S. Ortiz. Universidad de La Frontera, Departamento de Psicología. Montevideo 0831, Temuco, Chile. +56 45 2592440 manuel.ortiz@ufrontera.cl

Esta investigación fue financiada por CONICYT (Comisión Nacional de Investigación Científica y Tecnológica, Gobierno de Chile), Proyecto FONDECYT de INICIACION N°11140454 cuyo investigador principal fue el Dr. Manuel S. Ortiz.

Resumen

El modelo de ecuaciones estructurales (SEM) es una técnica de análisis estadística multivariada, que permite analizar patrones complejos de relaciones entre variables, realizar comparaciones entre e intragrupos, y validar modelos teóricos y empíricos. SEM puede ser utilizado para responder una amplia variedad de preguntas de investigación tanto en diseños experimentales como no experimentales. Pese a sus ventajas sobre técnicas tradicionales como la regresión múltiple o ANOVA, su uso en ciencias médicas y de la salud es poco frecuente. Por tanto, el objetivo de este artículo es introducir esta técnica de análisis a investigadores de las ciencias médicas y de la salud, explicando su aplicación con ejemplos del estudio chileno de predictores psicológicos de obesidad y síndrome metabólico (PPOMS). Se espera contribuir a la comprensión de esta técnica de análisis entre lectores de manuscritos científicos y estimular su uso entre investigadores de las ciencias médicas y de la salud.

Palabras clave: Estadísticas, análisis multivariado, modelo de ecuaciones estructurales.

Abstract

Structural equation modeling (SEM) is a multivariate statistical analysis technique, utilized to analyze complex patterns of relationships among a set of variables, conduct between-groups and within-groups comparisons, and validate theoretical and empirical models. SEM can be used to answer several research questions including those formulated in the context of experimental and non-experimental designs. Despite the several advantages that SEM has over traditional procedures, such as multiple regression or ANOVA, it has not been applied frequently in the medical and health science domains. Therefore, the purpose of this article is to present SEM as a robust and comprehensive analytical technique capable of strengthening and increasing the accuracy of the analyses in medical and health research. The functioning and applications of SEM are illustrated through one research example: a Chilean study of psychological predictors of obesity and metabolic syndrome. This article is aimed at contributing to a better understanding of this technique among readers of scientific publications and facilitating its implementation in future research.

Keywords: Statistics; Multivariate analysis; Structural equation modeling.

Introducción

El modelo de ecuaciones estructurales (SEM por su nombre en inglés “Structural Equation Modeling”) es una técnica de análisis estadístico multivariada, ampliamente utilizada en ciencias sociales y de la conducta (Bentler y Weeks, 1980; Beran y Violato, 2010; Ullman y Bentler, 2013), que permite analizar el patrón de relaciones entre una o más variables independientes (VIs) y una o más variables dependientes (VDs). SEM provee un marco flexible para el desarrollo y análisis de relaciones múltiples, útil para validar teorías usando modelos empíricos (Salinas-Oñate, Ortiz, Baeza-Rivera, y Betancourt, 2017) y extensamente utilizado en procesos de validación de instrumentos psicométricos (Fonseca-Pedrero y Ortu, 2015; Ortiz, Gómez-Pérez, Cancino, y Barrera-Herrera, 2016). Permite

modelar el error de medición, siendo ésta una de sus mayores ventajas por sobre otros métodos tradicionales de análisis de datos. SEM puede ser aplicado a una amplia variedad de preguntas de investigación, tipos de variables (continuas y discretas), y diseños de investigación (experimentales y no experimentales).

Aunque el uso de esta técnica es común en ciencias sociales y de la conducta, su uso en medicina y epidemiología es menos frecuente (Fonseca-Pedrero y Ortu, 2015). Por tanto, el objetivo de este artículo es introducir las principales características de SEM y guiar al lector en su implementación, desarrollando ejemplos y resultados con datos obtenidos del proyecto FONDECYT de INICIACIÓN 11140454 “predictores psicológicos de obesidad y síndrome metabólico en una muestra de adultos chilenos”.

SEM y otras técnicas estadísticas.

SEM integra múltiples técnicas estadísticas, tales como el análisis de la varianza (ANOVA), análisis de regresión múltiple y análisis factorial. Por ejemplo, comparaciones de dos grupos o más, o comparaciones intragrupos pueden ser analizadas con ANOVA o SEM, llegando indistintamente al mismo resultado. SEM permite también responder preguntas que implican análisis de regresión múltiple. En un nivel muy simple, un investigador podría plantear la relación entre una variable observada, por ejemplo, estrés agudo, estrés crónico o sexo y otra variable medida directamente, por ejemplo, perímetro de cintura. Tal como se observa en la Figura 1, el análisis más simple en este caso, implica la regresión múltiple. Las cuatro variables medidas son representadas con rectángulos y conectadas por líneas con flechas que indican que las variables independientes (VIs) “estrés agudo”, “estrés crónico” y “sexo” predicen la variable dependiente (VD) “perímetro de cintura”. Líneas con doble flecha indican covarianza entre las variables. La presencia de error o residuo (e y d) indica que la predicción no es perfecta.

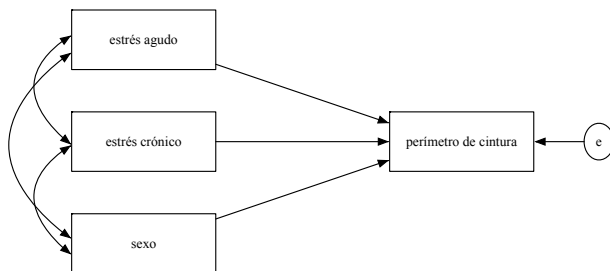


Figura 1. Modelo de sendero (path).

Fuente: Elaboración propia.

La Figura 2 representa un modelo más complejo, en el cual el Síndrome Metabólico es un factor latente (representado por círculos) que no es medido directamente sino más bien evaluado indirectamente usando seis indicadores (perímetro de cintura, triglicéridos, colesterol HDL, glucosa, presión sistólica y presión diastólica). El Síndrome Metabólico es predicho por otro factor latente de Estrés Psicológico, el cual es medido indirectamente por dos indicadores que sí son medidos directamente (estrés agudo y estrés crónico). El sexo (variable observada) predice simultáneamente al Síndrome Metabólico y al Estrés Psicológico. La nomenclatura de SEM indica que los factores latentes sean identificados con una mayúscula inicial y las variables observadas o indicadores con minúsculas.

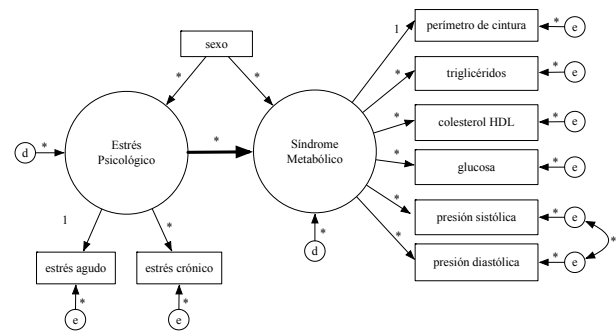


Figura 2. Modelo de ecuaciones estructurales

Fuente: Elaboración propia.

Las Figuras 1 y 2 son ejemplos de diagramas “path” o SEM, por su denominación en inglés. Estos diagramas permiten al investigador representar gráficamente sus hipótesis, facilitando que éstas sean convertidas en las ecuaciones necesarias para el análisis.

Para desarrollar un diagrama se debe cumplir con diversas convenciones. Las variables medidas, también llamadas variables observadas, indicadores o variables manifiestas, son representadas por cuadrados o rectángulos. Los factores tienen dos o más indicadores y pueden ser llamados factores latentes, variables latentes, constructos o variables no observadas. Los factores son representados por círculos u óvalos. Las relaciones entre variables son indicadas por líneas. La ausencia de líneas conectando variables indica que no se ha hipotetizado una relación entre ellas. Las líneas tienen una o dos flechas. Una línea con una flecha indica una relación directa entre las dos variables, siendo aquella a la cual la flecha apunta la VD. Una línea con doble flecha indica una covarianza entre las variables, que no implica dirección o efecto.

En la Figura 2, el Síndrome Metabólico es un factor latente predicho por el Estrés Psicológico, el cual es otro factor latente. Nótese que el sexo predice tanto al Síndrome Metabólico como al Estrés Psicológico. La línea con una flecha indica la dirección de la predicción. Las líneas con doble flecha entre presión sistólica y diastólica, indican asociación (covarianza) entre las variables observadas, pero sin hacer una predicción sobre la dirección o el efecto. La dirección de las flechas que conectan al Síndrome Metabólico (factor latente) con sus indicadores, indican que el Síndrome Metabólico predice a las variables medidas. En otras palabras, el Síndrome Metabólico como constructo, resulta imposible de medir directamente, por tanto, se mide a sus indicadores (perímetro de cintura, triglicéridos,

colesterol HDL, glucosa, presión sistólica y presión diastólica). Este es el mismo caso para el factor latente de Estrés Psicológico. Tal como en la regresión lineal, nada es medido perfectamente, habiendo, por tanto, error o residuos. En SEM el error de las variables observadas y los factores latentes es representado con esferas con las letras “e” y “d” respectivamente. El error del indicador (e) es la varianza que no es explicada por el factor y que suele considerarse error de medición. El error del factor en la regresión (d) representa cualquier variable no considerada en la predicción y, por tanto, permite estimar el porcentaje de varianza explicado por la regresión (Kline, 2013; Tu, 2009).

SEM: Modelo de medición y modelo estructural.

Un modelo de medición es aquel que evalúa la relación entre variables observadas y factores latentes. En la Figura 2 puede observarse que los indicadores perímetro de cintura, triglicéridos, colesterol HDL, glucosa, presión arterial sistólica y presión arterial diastólica comparten una causa común subyacente denominada Síndrome Metabólico. Similarmente, un puntaje de estrés agudo y uno de estrés crónico correlacionan entre sí y son explicados por un factor latente que ha sido designado Estrés Psicológico. El objetivo de un factor latente en un modelo de medición es identificar el número y la naturaleza de los factores que dan cuenta de la variación común entre un conjunto de indicadores (Babin y Svensson, 2012; Brown y Moore, 2012; Kaplan, 2001). Un factor latente, por tanto, da cuenta de la correlación entre los indicadores. En otras palabras, las variables observadas están inter-correlacionadas porque comparten una causa común (están influidas por el mismo concepto subyacente) (Babin y Svensson, 2012).

En un modelo estructural se evalúa la relación que existe entre una o varias variables observadas con un factor latente, o la relación existente entre factores latentes. Tal como se observa en la Figura 2, se hipotetiza que el Estrés Psicológico se asocia directamente con un Síndrome Metabólico. Asimismo, se espera que el sexo se asocie con ambos factores. La inclusión de la variable sexo al modelo representado en la Figura 2, permite establecer diferencias de medias entre hombres y mujeres tanto en el factor latente de Estrés Psicológico como en el de Síndrome Metabólico.

SEM y parámetros estimados

Tanto en un modelo de medición como en un modelo estructural, diferentes parámetros han de ser especificados y estimados. Los parámetros a ser especificados son: libres, fijos y restringidos. Un parámetro libre es aquel que tiene un valor desconocido y que se estima a partir del análisis. En la Figura 2, éstos son representados con asteriscos. Un parámetro fijo es pre-especificado por el investigador para adoptar un valor determinado, siendo 10 en la mayoría de los casos. Un parámetro restringido también tiene un valor desconocido, pero el investigador le impone un valor determinado. Por ejemplo, es común imponer una restricción de igualdad en el cual uno de los parámetros asume el mismo valor en diferentes grupos que se comparan (Jackson et al., 2009; Schmitt, 2011).

Etapas en SEM

SEM implica una serie de etapas, las cuales se describen a continuación: Especificación del modelo, estimación del modelo, evaluación de ajuste global del modelo y re-especificación del modelo. Cada una de estas etapas será ilustrada con ejemplos y datos obtenidos del Proyecto Fondecyt de Iniciación 11140454 “predictores psicológicos de obesidad y síndrome metabólico en una muestra de adultos chilenos”.

Especificación del modelo

En esta etapa inicial se desarrollan las hipótesis acerca de la relación entre las variables. El investigador debe decidir si las relaciones son unidireccionales o bidireccionales, basándose en la teoría, en la evidencia previa o en ambas (Bentler y Weeks, 1980). Estas relaciones pueden ser directas o indirectas, en la medida que en el modelo se incluyan variables que medien el efecto de una variable sobre la otra (Bou y Satorra, 2010; Byrne, 2008). SEM es, por tanto, una técnica confirmatoria más que exploratoria.

En la Figura 2 se observa que el investigador hipotetiza la existencia de dos factores latentes, cada uno de los cuales tiene sus propios indicadores. En el caso del Estrés Psicológico, sus indicadores son: estrés agudo y estrés crónico. En tanto, para el Síndrome Metabólico sus indicadores son: perímetro de cintura, triglicéridos, colesterol HDL, glucosa, presión sistólica y presión diastólica.

De esta manera, valores elevados en las variables (indicadores) psicológicas implican alto estrés psicológico. Más aún, alto Estrés Psicológico predice elevado estrés agudo y crónico (obsérvese la dirección de las flechas). Lo anterior es similar para el caso del factor latente Síndrome Metabólico.

En este modelo también se hipotetiza que el Estrés Psicológico se asocia directamente con el Síndrome Metabólico y que el sexo predice tanto al Estrés Psicológico como al Síndrome Metabólico (véase la dirección de las flechas).

Estas relaciones son traducidas a ecuaciones, las cuales permiten estimar el modelo. Un método de estimación ampliamente utilizado es el denominado método Bentler-Weeks (Bentler y Weeks, 1980), en el cual los parámetros que se estiman son: cargas o pesos factoriales, coeficientes de regresión, varianzas y covarianzas, errores y residuos.

Identificación

En SEM sólo modelos que están identificados pueden ser estimados. Para que un modelo esté identificado, el número de parámetro que será estimado (asteriscos en la Figura 2) no puede exceder al número de datos observados. Aplicando la siguiente fórmula se obtiene el total de “data points”, donde p equivale al número de variables observadas:

$$\text{número de “datapoints”} = \frac{p(p+1)}{2}$$

En la Figura 2, p equivale a 9, por tanto, el número de “data points” es 45. Los parámetros a ser estimados son 20, con lo cual el modelo está sobre-identificado; condición necesaria para proceder con el análisis. En el caso que exista el mismo número de “data points” y parámetros, el modelo estará solo-identificado y, por tanto, el análisis será irrelevante pues la hipótesis del ajuste del modelo a los datos no podrá ser comprobada. Pese a esto, las hipótesis sobre coeficientes de regresión pueden ser testeados. Si existieran menos “data points” que parámetros, el modelo estaría sub-identificado y sus parámetros no podrían ser estimados. Para mayores detalles sobre identificación refiérase a Kline (2013).

Estimación del modelo

Una vez que el modelo ha sido especificado e identificado, sus parámetros son estimados con el objetivo de minimizar

la diferencia entre la matriz de covarianza observada y la matriz de covarianza poblacional (Hayes, 2013).

Aunque hay muchos procedimientos de estimación disponibles, tres son presentados en este artículo. El primero de ellos es la estimación de máxima-verosimilitud (ML por su nombre en inglés “maximum-likelihood”). Este es el método que por defecto utilizan la mayoría de los softwares que ejecutan SEM. Esta estimación es un proceso iterativo que determina en qué medida el modelo predice los valores de la matriz de covarianza, con valores cercanos a cero indicando un mejor ajuste. ML requiere de tamaños muestrales grandes (generalmente $n > 300$) y ha demostrado ser robusta con datos que no distribuyen en forma normal (Hayes y Preacher, 2014; Hoyle, 2012). Otro método común de estimación es el de mínimos cuadrados (LS por su nombre en inglés “least squares”). LS es similar a ML por cuanto estima los patrones de relaciones entre las variables, pero lo hace minimizando la suma de cuadrados de la desviación entre el modelo hipotetizado y el observado. LS funciona mejor que ML con muestras pequeñas y provee una mejor estimación cuando los supuestos de distribución normal e independencia son violados (Bollen y Long, 1993; Saris et al., 2009).

Finalmente, el tercer método de estimación es conocido como distribución asimptótica libre (ADF por su nombre en inglés “asymptotically distribution free”). Aunque esta estimación es menos utilizada que las anteriores, funciona bien si la distribución de datos es muy asimétrica. Este método requiere tamaños muestrales grandes (200 a 500) para obtener estimadores confiables. Para mayores detalles sobre estimación revise a Hu, Bentler y Kano (Hu, Bentler, y Kano, 1992).

Evaluación del ajuste global del modelo

Típicamente el ajuste del modelo es evaluado con el chi-cuadrado (χ^2), el cual es muy sensible a tamaños muestrales grandes. Por tal motivo, la evaluación global del modelo suele realizarse considerando otros indicadores de bondad de ajuste, los cuales tienen un rango que va de 0 a 1, con valores altos sugiriendo mayor varianza explicada por el modelo. El “Comparative Fit Index” (CFI) y el “Tucker Lewis Index” (TLI) son ampliamente empleados y comparan el modelo existente con uno nulo. Valores de CFI superiores a 0.95 y TLI mayores a 0.90, son indicadores de un buen

ajuste (Hooper et al., 2008; Hu y Bentler, 1999). Un buen calce de los datos al modelo es también logrado con valores de residuos cercanos a 0, lo cual representa el porcentaje de varianza no explicado por el modelo. De éstos, el “Root Mean Square Error of Approximation (RMSEA)” y el “Standardized Root Mean Residual” (SRMR) son frecuentemente reportados, con valores de RMSEA menores

a 0.6 y SRMR inferiores a 0,08 considerados aceptables (Fan y Sivo, 2007). En el caso del modelo desarrollado en este artículo, los indicadores de bondad de ajuste son los siguientes $\chi^2(24) = 68.6$; $p < 0.001$; CFI = 0.96; TLI = 0.93; RMSEA = 0.069; SRMR = 0.049. En conjuntos estos indicadores señalan un muy buen ajuste global del modelo a los datos (ver Figura 3).

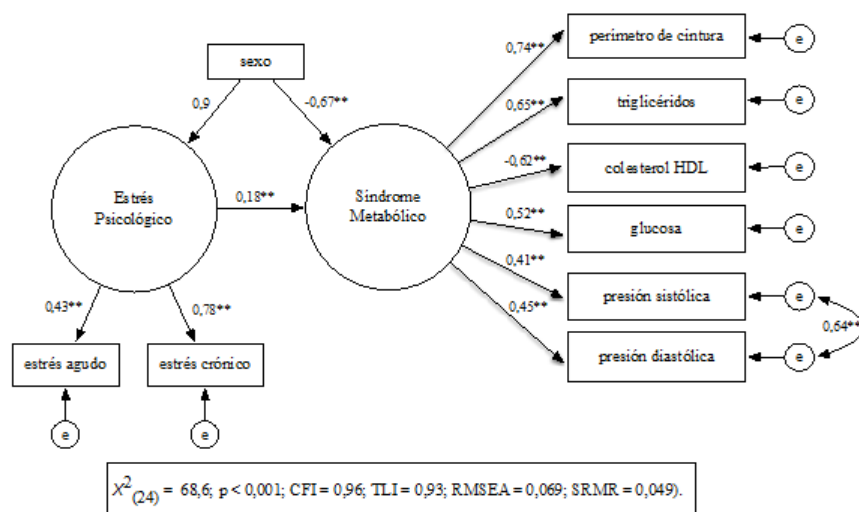


Figura 3. Modelo de ecuaciones estructurales con resultados

Nota. Los valores de los errores y las varianzas de los factores latentes han sido omitidos para claridad del modelo.
Fuente: Elaboración propia.

Re-especificación del modelo

Una vez que se ha evaluado el ajuste global del modelo a los datos, se pueden tomar dos decisiones. La primera de ellas es la interpretación y reporte de los resultados y la segunda la re-especificación del modelo. Aunque la primera opción es la deseada, la evaluación del ajuste global del modelo a los datos no es siempre óptima, necesiándose, por tanto, re-especificar el modelo. Nótese que la re-especificación requiere la reconsideración de la “identificación” del modelo, luego volver a estimar el modelo y evaluar su bondad de ajuste. Razones para re-especificar el modelo pueden ser la búsqueda de un patrón de correlación omitido inicialmente, la detección de un residuo de gran valor o el uso de algún algoritmo estadístico, como por ejemplo el Test de Lagrange

(“Lagrange Multiplier Test”) el cual evalúa en forma incremental la mejora del ajuste global del modelo por cada parámetro que es restringido o liberado del modelo (Buse, 1982).

Interpretación y reporte

La principal preocupación que debe tenerse al interpretar SEM es el significado de un determinado parámetro en el modelo y el grado en el cual éste da cuenta de los datos observados. No debe olvidarse que el modelo a priori refleja una teoría o un conjunto de hipótesis interrelacionadas que pueden derivarse de diversas teorías, por tanto, la interpretación de un parámetro es siempre teórica, sin importar cuál sea su resultado. Aunque la causalidad en la interpretación de un parámetro es algo que escapa el foco de este artículo, usando SEM es posible hacer una interpretación causal siempre y cuando ésta sea adecuada y correctamente fundamentada, incluso en diseños correlacionales y transversales (Hayes, 2013; Pearl, 2012).

Ejemplo de reporte de resultados

En el caso del modelo representado en la Figura 3, es posible identificar un efecto directo del factor latente de Estrés Psicológico en el Síndrome Metabólico ($\beta = 0.18$; $p < 0.01$). El sexo de los participantes tiene un efecto directo en el Síndrome Metabólico. Dado que esta variable ha sido consignada con los valores 0 para hombres y 1 para mujeres, se observa que las mujeres presentan menos síndrome metabólico que los hombres ($\beta = -0.67$; $p < 0.01$). Finalmente, no existe efecto del sexo en el estrés ($\beta = 0.09$; $p = 0.31$). Los indicadores de bondad de ajuste del modelo global son muy buenos ($\chi^2(24) = 68.6$; $p < 0.001$; CFI = 0.96; TLI = 0.93; RMSEA = 0.069; SRMR = 0.049). De esta manera, es posible determinar que existe un efecto estadísticamente significativo del estrés psicológico sobre el síndrome metabólico, y que las mujeres presentan menos síndrome que los hombres, en esta muestra de adultos chilenos.

Conclusiones

SEM es una poderosa técnica de análisis estadístico superior a la regresión múltiple pues permite analizar múltiples relaciones teóricas simultáneamente y además hacer comparaciones entre grupos e intragrupos siendo, por tanto, similar a ANOVA. Con SEM es factible modelar el error y el uso de factores latentes reduce el error de medida. Aunque SEM ha sido mayoritariamente utilizado en ciencias sociales y del comportamiento, es menos utilizada en ciencias médicas y epidemiología. Su uso en estas disciplinas puede contribuir a generar modelos complejos que incorporen múltiples variables o indicadores, mejorando la comprensión de cómo diversas variables o factores se interrelacionan entre sí y explican desenlaces o comportamientos en salud.

Pese a las virtudes aquí presentadas, el investigador debe tener un conocimiento teórico profundo del fenómeno de estudio, pues SEM es una técnica confirmatoria. En su implementación, se debe usar el diseño metodológico correcto, considerar los tamaños muestrales adecuados y tener medidas robustas para las variables observadas.

Referencias

- Babin, B. J., y Svensson, G. (2012). Structural equation modeling in social science research. *European Business Review*, 24, 320–330. <http://doi.org/10.1108/09555341211242132>
- Bentler, P. M., y Weeks, D. G. (1980). Linear structural equations with latent variables. *Psychometrika*, 45, 289–308. <http://doi.org/10.1007/BF02293905>
- Beran, T. N., y Violato, C. (2010). Structural equation modeling in medical research: A primer. *BMC Research Notes*, 3, 267. <http://doi.org/10.1186/1756-0500-3-267>
- Bollen, K., y Long, J. S. (1993). *Testing structural equation models* (Vol. v. 154). Newbury Park: Sage Publications. Retrieved from http://www.worldcat.org/title/testing-structural-equation-models/oclc/26856260?referer=brief_results
- Bou, J. C., y Satorra, A. (2010). A multigroup structural equation approach: A demonstration by testing variation of firm profitability across EU samples. *Organizational Research Methods*, 13, 738–766. <http://doi.org/10.1177/1094428109340433>
- Brown, T. A., y Moore, M. T. (2012). Confirmatory factor analysis. In R. Hoyle (Ed.), *Handbook of Structural Equation Modeling* (pp. 361–379). New York: The Guildfor Press.
- Buse, A. (1982). The likelihood ratio, wald, and lagrange multiplier tests: An expository note. *The American Statistician*, 36, 153–157. <http://doi.org/10.1080/00031305.1982.10482817>
- Byrne, B. M. (2008). Testing for multigroup equivalence of a measuring instrument: A walk through the process. *Psicothema*, 20, 872–882.
- Fan, X., y Sivo, S. A. (2007). Sensitivity of fit indices to model misspecification and model types. *Multivariate Behavioral Research*, 42, 509–529. <http://doi.org/10.1080/00273170701382864>
- Fonseca-Pedrero, E., y Ortu, J. (2015). Detección del riesgo para los trastornos del espectro bipolar: Evidencias de validez del Mood Disorder Questionnaire en adolescentes y adultos jóvenes. *Revista de Psiquiatría y Salud Mental*, 9, 4–12. <http://doi.org/10.1016/j.rpsm.2015.04.003>
- Hayes, A. (2013). Introduction to mediation, moderation, and conditional process analysis. *New York, NY: Guilford*, 3–4. <http://doi.org/978-1-60918-230-4>
- Hayes, A., y Preacher, K. (2014). Statistical mediation analysis with a multicategorical independent variable. *British Journal of Mathematical and Statistical Psychology*, 67, 451–470. <http://doi.org/10.1111/bmsp.12028>
- Hooper, D., Coughlan, J., y Mullen, M. R. (2008). Structural equation modelling: Guidelines for determining model fit. *Electronic Journal of Business Research Methods*, 6, 53–60.
- Hoyle, R. H. (2012). *Handbook of structural equation modeling*. New York: The Guildford Press.
- Hu, L., Bentler, P., y Kano, Y. (1992). Can test statistics in covariance structure analysis be trusted? *Psychological Bulletin*, 112, 351–362.
- Hu, L., y Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling: A Multidisciplinary Journal*, 6, 1–55. <http://doi.org/10.1080/10705519909540118>
- Jackson, D. L., Gillasp, J. A., y Purc-Stephenson, R. (2009). Reporting practices in confirmatory factor analysis: An overview and some recommendations. *Psychological Methods*, 14, 6–23. <http://doi.org/10.1037/a0014694>

- Kaplan, D. (2001). Structural Equation Modeling. *International Encyclopedia of the Social y Behavioral Sciences*, 15215–15222. <http://doi.org/10.1080/10705510802154356>
- Kline, R. B. (2013). Principales and practice of structural equation modeling. *Journal of Chemical Information and Modeling*, 53, 1689–1699. <http://doi.org/10.1017/CBO9781107415324.004>
- Ortiz, M. S., Gómez-Pérez, D., Cancino, M., y Barrera-Herrera, A. (2016). Validación de la versión en Español de la Escala de Optimismo Disposicional (LOT-R) en una muestra Chilena de estudiantes universitarios. *Terapia Psicológica*, 34, 53–58. <http://doi.org/10.4067/S0718-48082016000100006>
- Pearl, J. (2012). The causal foundations of structural equation modeling. In R. H. Hoyle (Ed.), *Handbook of structural equation modeling* (pp. 68–91). New York: The Guildford Press.
- Salinas-Oñate, N., Ortiz, M. S., Baeza-Rivera, M. J., y Betancourt, H. (2017). Desarrollo de un instrumento culturalmente pertinente para medir creencias en psicoterapia. *Terapia Psicológica*, 35, 15–22.
- Saris, W. E., Satorra, A., y van der Veld, W. M. (2009). Testing structural equation models or detection of misspecifications? *Structural Equation Modeling: A Multidisciplinary Journal*, 16, 561–582. <http://doi.org/10.1080/10705510903203433>
- Schmitt, T. (2011). Current Methodological Considerations in Exploratory and Confirmatory Factor Analysis. *Journal of Psychoeducational Assessment*, 29, 304–321. <http://doi.org/10.1177/0734282911406653>
- Tu, Y. K. (2009). Commentary: Is structural equation modeling a step forward for epidemiologist. *International Journal of Epidemiology*, 38, 1–3.
- Ullman, J. B., y Bentler, P. M. (2013). *Structural equation modeling*. In S. J. A. y W. F. Velicer (Eds.), *Handbook of Psychology* (pp. 661–690). New York: Hoboken. <http://doi.org/10.1016/B978-0-444-53737-9.50010-4>