



International Journal of Quantitative Research in Education

ISSN online: 2049-5994 - ISSN print: 2049-5986

<https://www.inderscience.com/ijqre>

Analysing observed categorical data in SPSS AMOS: a Bayesian approach

Hongwei Yang, Lihua Xu, Mark Malisa, Menglin Xu, Qintong Hu, Xing Liu, Hyungsoo Kim, Jing Yuan

DOI: [10.1504/IJORE.2023.10051888](https://doi.org/10.1504/IJORE.2023.10051888)

Article History:

Received:	29 February 2020
Accepted:	24 March 2022
Published online:	28 March 2023

Analysing observed categorical data in SPSS AMOS: a Bayesian approach

Hongwei Yang*

University of West Florida,
11000 University Pkwy,
Pensacola, FL 32514, USA
Email: pyang@uwf.edu
*Corresponding author

Lihua Xu

Orange County Public Schools,
445 W. Amelia St.,
Orlando FL 32801, USA
Email: lihua.xu@ocps.net

Mark Malisa

University of West Florida,
11000 University Pkwy,
Pensacola, FL 32514, USA
Email: mmalisa@uwf.edu

Menglin Xu

Ohio State University,
29 W Woodruff Ave.,
Columbus, OH, 43210, USA
Email: xumenglin920@gmail.com

Qintong Hu

Shandong University of Science and Technology,
No. 579 QianWanGang Rd., Qingdao,
Shandong, 266590, China
Email: qhuvols@163.com

Xing Liu

Eastern Connecticut State University,
83 Windham Street,
Willimantic, CT 06226, USA
Email: liux@easternct.edu

Hyungsoo Kim

University of Kentucky,
316 FB Family Sciences University of Kentucky,
Lexington, KY 40506, USA
Email: hkim3@uky.edu

Jing Yuan

Guangxi Normal University,
Guilin, Guangxi 541003, China
Email: 734216184@qq.com

Abstract: This study has a didactic purpose to help applied investigators and practitioners to understand the roles of observed categorical data (OCD) in structural equation modelling (SEM) and the appropriate ways of analysing such data under SPSS AMOS. To that end, the study reviews types of OCD (nominal, ordinal, dichotomous and polytomous) and their incorporation into SEM under AMOS to play different roles. The study presents two applications from the health and retirement study where Bayesian statistical inference is used to analyse one set of OCD variables serving as endogenous variables with/without groups created by another OCD variable. Besides, the study demonstrates the typical ways of summarising, reporting and interpreting the results from Bayesian statistics, and compares AMOS with several other SEM programmes (Mplus, R lavaan, Stata and SAS PROC CALIS) on handling OCD. The study concludes with summaries of the findings for its intended audience.

Keywords: structural equation modelling; SEM; multi-group analysis; Bayesian statistics; SPSS AMOS; categorical data analysis; Mplus; R lavaan; Stata; SAS PROC CALIS; health and retirement study; HRS.

Reference to this paper should be made as follows: Yang, H., Xu, L., Malisa, M., Xu, M., Hu, Q., Liu, X., Kim, H. and Yuan, J. (2022) 'Analysing observed categorical data in SPSS AMOS: a Bayesian approach', *Int. J. Quantitative Research in Education*, Vol. 5, No. 4, pp.399–430.

Biographical notes: Hongwei Yang is an Assistant Professor in the School of Education at the University of West Florida. His scholarship is found in online teaching and learning, family science, human and animal medical care, survey research, quantitative methodology, and psychometrics.

Lihua Xu works as a Senior Administrator in the Department of Research and Evaluation in Orange County Public Schools in Orlando, FL.

Mark Malisa is an Associate Professor in the School of Education at the University of West Florida. His research focuses on educational research methodology and its applications to various contemporary issues in education. He mentors and advises EdD students who are mostly school teachers and administrators to earn doctoral degrees.

Menglin Xu is a researcher from the Department of Internal Medicine at the Ohio State University. Her research focuses on psychometrics, causal inference, longitudinal study, and the application of quantitative methods to education, psychology, and health fields. She also provides methodology support to faculty and graduate students for grant proposals and research projects.

Qintong Hu is a faculty member at the College of Mathematics and System Sciences, Shandong University of Science and Technology and was on the faculty of mathematics education in Columbia College in the USA. Her research interests lie primarily in the fields of mathematics education, comparative studies and international education where she has published in top-tier peer-reviewed journals.

Xing Liu is a Professor of Educational Research and Assessment in the Education Department at Eastern Connecticut State University. He is the author of *Categorical Data Analysis and Multilevel Modeling Using R* and *Applied Ordinal Logistic Regression Using Stata: From Single-Level to Multilevel Modeling*. His research interests include categorical data analysis, multilevel modelling, longitudinal data analysis, educational assessment, data science, and Bayesian methods.

Hyungsoo Kim is an Associate Professor of Family Finance and Consumer Economics at the University of Kentucky. His research centres on financial security and health among middle-aged and older Americans. His studies on non-cognitive motivational retirement savings using psychological self-regulation models and cognitive contributors of wealth accumulation such as financial literacy within limited financial resources have greatly contributed to the current field.

Jing Yuan is currently working at Guangxi Normal University. She has strong and extensive methodological and content expertise. She focuses on statistics and analytics and is also applying educational psychology to STEAM education. Her education and experience have helped her develop technical capabilities in data analysis and scientific exploration.

1 Introduction

Structural equation modelling (SEM) provides a powerful framework for modelling various complex relationships using multivariate observed data. The observed data analysed in SEM could be categorical as well as continuous. In many fields of studies, observed categorical data (OCD) are prevalent:

- 1 1 = yes or 2 = no about the existence of a symptom
- 2 1 = agree, 2 = neutral or 3 = disagree about a policy, among other things.

The values assigned to the options of a categorical variable are arbitrarily selected (e.g., 0 and 1, or 1 and 2), letters or strings. Therefore, such data are usually considered to be qualitative, or non-numeric even if they may be labelled in the form of discrete, numeric values. Given that OCD violates many assumptions of statistical models optimised for continuous data, the analysis of such data usually requires special considerations and could be challenging to some (Bhardwaj, 2015; Hadi, 2015). Besides, another challenge with OCD is that SEM software programmes vary significantly in terms of how to incorporate such data into the model for analysis.

SPSS AMOS is a popular SEM programme among applied investigators and practitioners, such as those in social and behavioural sciences (Boateng, 2020; Perera, 2013; Peterson et al., 2013). AMOS is capable of analysing OCD in multiple ways to

serve various research needs. However, a review of the literature indicates the categorical analysis capability of AMOS is not as well known as many of its other capabilities to the extent that some even mistakenly believe AMOS can only be used to analyse continuous data (Gupta, 2016; Yorgason, n.d.). Among those who have some knowledge of the categorical analysis features in AMOS, some still feel confused about how to properly interpret the analysis results (Cristofaro, 2016a, 2016b).

Several factors may account for the lack of understanding in and inadequate use of AMOS for analysing OCD. For one, the only platform for handling OCD provided in AMOS is Bayesian SEM (Arbuckle, 2019). Bayesian statistical inference is relatively new and is much more technically challenging, therefore making it difficult for many AMOS users to even understand the Bayesian method, let alone implementing it. In fact, users who choose to use AMOS which is best known for its point-and-click environment are probably themselves not as well trained in statistics as those who choose to use another SEM programme (e.g., Mplus) which primarily relies on a syntax-based environment for model specification. For another, the literature is scant on the use and performance of the categorical data analysis features of AMOS, particularly little, if any, discussion on its Bayesian SEM platform. In the absence of such information (preferably having been thoroughly evaluated by peers), AMOS users mainly count on online discussions or informal talks for empirical evidences which may or may not be adequately compelling and convincing in research projects requiring scientific rigor.

This study focuses on the handling of OCD in AMOS to familiarise applied investigators and practitioners with the appropriate strategies of analysing such data under this popular SEM programme. To that end, the study presents several routine SEM analyses involving different types of OCD (nominal, ordinal, dichotomous, polytomous) playing various roles in the model (endogenous, exogenous, creating groups). During the demonstrations, the Bayesian platform in AMOS is used to estimate each model and draw substantive conclusions through Bayesian statistical inference.

In sum, the organisation of the study is as follows. The section that follows immediately discusses OCD and their common uses in SEM. Following that, the next section covers briefly Bayesian statistics and the Bayesian platform in AMOS, which precedes another section where four other SEM programmes are discussed regarding their OCD handling capabilities, as a comparison with the capabilities in AMOS. Then, using a dataset from the health and retirement study (HRS), two related numeric examples where different types of OCD variables play different roles are presented in AMOS and three of the four comparison software programmes. The study concludes with summaries of findings for its intended audience.

2 Categorical data and their applications in SEM

According to Agresti (2013), Lee and Song (2003), Lee et al. (2010) and Olsson (1979), an OCD variable can be dichotomous or polytomous. With both dichotomous and polytomous data, a distinction is made between two types of categorical scales:

- 1 nominal categorical
- 2 ordinal/ordered categorical.

Variables having categories without an ordering are measured on a nominal scale whereas those with an ordering on an ordinal/ordered scale. A variable's measurement scale determines which statistical methods are appropriate. It is usually best to apply methods appropriate for the actual scale: application of the right estimator capable of handling data specified/declared to be measured on a certain scale (e.g., continuous, categorical, etc.) (Li, 2021). Many SEM software programmes, like Mplus, R lavaan, Stata and AMOS, allow the specification of individual variables as categorical and also provide the appropriate estimators (and model structures/formulations) for handling the specified categorical data. When there is a mismatch between the statistical methods and the distribution of the data (e.g., statistical methods designed for nominal categorical data applied to ordinal categorical data, and vice versa), misleading results may occur. For example, Lee et al. (2010, pp.294–296) identified inflated standard errors when a statistical method designed for ordinal dichotomous data was applied to nominal dichotomous data, and vice versa. Finally, when used in SEM, an OCD variable can be an endogenous variable, but it can also serve a different purpose than being endogenous (i.e., non-endogenous).

2.1 *OCD endogenous variable*

When an OCD variable is endogenous in a model specified in AMOS, it should be (declared as) ordinal (including ordinal dichotomous data) in order for the software to properly deal with it (at least, not treating it the same way as a continuous variable). To declare a categorical variable as ordinal/ordered categorical, the *allow non-numeric data* box in the *data files* window (found under *file/data files*) should be checked first. Next, the categorical variable should be declared as *ordered-categorical* within the *data recode* window (found under *tools*). With the declaration completed, AMOS will use an ordinal probit model formulation to analyse the OCD endogenous variable in an ordinal probit regression. This declaration process should be repeated for each OCD endogenous variable in the model.

For each OCD endogenous variable which has been declared as ordinal, a continuous latent distribution is assumed to be underlying the ordered response categories. There are response boundaries or thresholds associated with the underlying continuous distribution, with which to progress from one response category to the next (e.g., from agree to strongly agree in a five-point Likert scale item). For an OCD variable with k categories, a total of $(k - 1)$ response boundaries are needed. In AMOS, after a variable has been declared as ordered categorical, the boundaries can be either user-specified or automatically determined by the software.

Finally, as of this writing, AMOS is still not able to process nominal (including nominal dichotomous) endogenous data where the categories are not ordered. Although, within the *ordered-categorical* window, such a variable can nevertheless be declared as ordinal categorical so that AMOS will not treat it as a continuous variable, the literature indicates this practice is likely to lead to misleading results [Lee et al., (2010), pp.294–296].

2.2 *OCD non-endogenous variable*

When an OCD variable is not endogenous, it can play one of multiple roles. First, it can be exogenous. In this instance, AMOS should be set up to treat it as a continuous

variable. Therefore, regardless of whether the OCD variable (e.g., with k categories) is ordinal or nominal, replacing it with $(k - 1)$ dummy indicator variables is always a viable option, which is similar to how a categorical predictor is handled in a multiple linear regression analysis. The dummy variables for replacing the OCD exogenous variable have to be created manually outside AMOS (e.g., in SPSS Statistics) which has little, if any, data management and manipulation capability. Besides, if the OCD exogenous variable is ordinal, allocating to the categories a set of numeric codes (e.g., 1, 2, 3, ...), whose ordering is consistent with the ordering of the categories, and analysing these ordered numeric codes as if they were the values of a continuous variable is a second viable, but also sub-optimal option due to the arbitrary metric for the categories of the ordinal OCD exogenous variable created by the arbitrarily allocated numeric codes [Kutner et al., (2005), pp.321–322].

Next, besides being exogenous, an OCD non-endogenous variable can also serve to create groups for comparison based on its categories (i.e., multi-group analysis or MGA (e.g., for testing moderated mediation) versus single group analysis or SGA) (IBM SPSS, 2020). In particular, when the model is formulated as a factor analytic model (i.e., measurement model), MGA usually serves as an analysis of measurement invariance for evaluating if the same construct(s) is/are being measured similarly across specified groups. Finally, when the structural equation model goes beyond a measurement model to also include the structural relationships, MGA can be implemented to find out if a structural coefficient is significantly different across the groups specified by the categories of the OCD non-endogenous variable.

3 Bayesian statistics and AMOS Bayesian SEM

Over the past decades, Bayesian estimation has become increasingly widely used in social and behavioural sciences, including structural equation modelling (Depaoli, 2021; Lee, 2007; Song and Lee, 2012), item response theory (Chang and Sheng, 2017; Fox, 2010; Kuo and Sheng, 2015, 2016, 2017; Sheng, 2015, 2017a, 2017b), statistics and psychometrics in general (Gill and Walker, 2005; Gill, 2007; Gill and Witko, 2013; Jackman, 2009; Kaplan, 2014; Kruschke, 2014; Levy and Mislevy, 2016; Lynch, 2007; van de Schoot et al., 2014; Wagner and Gill, 2005; Wang and Preacher, 2015; Yuan and MacKinnon, 2009), among others. Bayesian statistics is based on the Bayes theorem and treats parameters as random quantities represented using probability distributions. Bayesian statistics aims to obtain the posterior distributions of parameters, given:

- 1 the data (combined with the model specification to derive the likelihood function)
- 2 the prior knowledge.

It is the summaries (posterior means, posterior standard deviations (SD), etc.) of the random parameter values obtained from the posterior distributions that are usually reported and interpreted at the conclusion of a Bayesian analysis.

A central task in Bayesian statistics is to conduct statistical inference using the posterior distributions of model parameters. In most practical applications, the posterior distribution cannot be derived by analytical means, and therefore simulation-based Markov chain Monte Carlo (MCMC) methods are frequently used to approximate and sample the posterior distribution. Usually, an MCMC process begins with a burn-in

period when the samples are drawn/simulated through one or more constructed Markov chains and subsequently discarded (i.e., not used in approximating the posterior distribution) to minimise the effects of initial values of the model parameter vector on the posterior inference. Notably, AMOS further defines and implements a pre-burn-in period with which to get past those values of the parameter vector which are of very low probability (IBM SPSS, 2018). After the burn-in period ends, the samples continue to be simulated through the chains to approximate the posterior distribution. Under strict conditions of ergodicity and reversibility, the chains will gradually converge to an equilibrium distribution as the target distribution.

In AMOS, Bayesian statistics is the only built-in platform capable of handling OCD. According to AMOS Development Corporation (2021a), the Bayesian platform offers two MCMC algorithms:

- 1 random walk metropolis (RWM)
- 2 Hamiltonian Monte Carlo (HMC).

RWM is one of the earliest MCMC methods developed by Metropolis et al. (1953). Because RWM is conceptually simple to understand and is easy to implement, it is still popular in numerous applications. Basically, RWM selects a candidate point (for the parameter vector) θ^* by taking a random perturbation around the current point (for the parameter vector) $\theta^{(t)}$ (i.e., $\theta^* = \theta^{(t)} + a\epsilon$) with a being a tuning parameter. The possibilities for the density for ϵ are endless, and they include the uniform distribution, the normal distribution (implemented in AMOS), etc.

On the other hand, the enticing simplicity of RWM is associated with poor performance of the algorithm with increasing dimensionality and complexity of the target distribution. RWM tends to explore the target distribution slowly in high dimensional problems, which are typical in SEM applications, and the MCMC estimates could be highly biased [Betancourt, (2018), p.16].

Over the years, substantial research has been devoted to improving the sampling efficiency of MCMC algorithms. Among such research is that based on the Hamiltonian dynamics known as the HMC method (Neal, 2011). HMC achieves better sampling efficiency than RWM by using the Hamiltonian dynamics in order for the chains to more efficiently explore the target distribution. Specifically, HMC makes use of the gradient information of the posterior distribution to reduce the random walk behaviour typical in RWM and allow the chains to move in the direction of high probabilities. Therefore, HMC is capable of substantially enhancing the efficiency of MCMC simulations even for highly parameterised models with complex multivariate dependencies among parameters.

Both RWM and HMC algorithms have their own parameters which influence whether the MCMC simulation can start (e.g., getting past the pre-burn-in period in AMOS) and the speed at which the simulation converges to a stationary distribution. In AMOS, RWM has one tuning parameter which represents the random perturbation around the current point of the parameter vector; HMC has two parameters representing, respectively, the number of leapfrog steps and the leapfrog step size.

Unfortunately, as of this writing, there is little literature on the proper selection of these MCMC parameter values in AMOS. The study recommends the default settings of each MCMC algorithm be used. In the case of any difficulty in getting the algorithm to start or the algorithm taking too long to reach convergence, one solution is to click the wrench-shaped *adapt* button in AMOS Bayesian SEM to have the software automatically

adjust the parameter(s) of the MCMC algorithm. The adjustment is made by using the information contained in the MCMC samples already simulated. Besides, an alternative solution is to take a trial-and-error approach to experiment with a variety of MCMC parameter values to see which one/ones performs/performs the best. This approach may be easier with the RWM algorithm with only one parameter to adjust.

4 Comparisons with other SEM software programmes

For the purpose of comparison, besides AMOS, also discussed and demonstrated here are several other SEM programmes:

- 1 Mplus
- 2 R lavaan
- 3 Stata
- 4 SAS PROC CALIS (discussion only).

These programmes have stable releases and, as of this writing, are well-maintained. More importantly, they are among those most widely used in SEM applications from various fields of studies. Out of these comparison SEM software programmes, R lavaan is the only non-commercial programme which does not require a paid license; SAS PROC CALIS is the only programme which only offers estimators for handling non-normal (not necessarily categorical, though) data but without also providing the capability of specifying individual variables as categorical or the appropriate model structure/formulation for handling OCD.

Mplus provides a *CATEGORICAL* (for ordinal categorical data) option and a *NOMINAL* (for nominal categorical data) option to deal with OCD endogenous variables including dichotomous data (Muthén and Muthén, 2017). These options should not be applied to OCD exogenous variables, though. When an exogenous variable in a model is declared to be categorical (nominal or ordinal), Mplus will stop executing the code and issue an error. In other words, Mplus treats the values of exogenous variables, including those of dummy indicators of categories, as continuous data. Next, in Mplus, the estimation of a model with OCD endogenous data could, but does not have to, rely on the Bayesian estimator. The software provides two frequentist estimators for that purpose:

- 1 a weighted least squares (WLS) estimator (probit regression)
- 2 a maximum likelihood (ML) estimator (logistic regression).

When the OCD endogenous data (including dichotomous data) are ordinal, all three estimators (Bayesian, WLS and ML) mentioned above can be used; when the OCD endogenous data (including the dichotomous data) are nominal, only the ML estimator can be used.

The lavaan package in R is capable of handling ordinal OCD endogenous data including dichotomous data (Rosseel, 2012). However, as of version 0.6–5, R lavaan provides no support for nominal OCD endogenous variables. Regarding the estimation of a model containing one or more ordinal OCD endogenous variables, the lavaan package provides both a three-stage WLS method and a pairwise likelihood method (Katsikatsou, n.d.; Rosseel, 2020). To invoke either estimator, the data should be declared as

categorical in the data frame containing the data (i.e., `base::ordered()` function) or, alternatively, in one of the model estimation functions (e.g., `lavaan::cfa()`, `lavaan::sem()` or `lavaan::lavaan()` function). Finally, when an exogenous variable in a model is declared as ordinal categorical, R lavaan will continue to execute the code but also issue a warning.

Stata offers a *gsem* command which allows the incorporation of OCD endogenous variables into the model through the generalised linear modelling framework (StataCorp, 2021). To properly use the command, the user should be familiar with statistical concepts like link function and distribution of the random component, which may be challenging to some. Even though Mplus, R lavaan and AMOS also use the generalised linear modelling framework when it comes to OCD endogenous variables, their choice of and reliance on easier-to-understand syntax key words or an intuitive point-and-click interface make them more user-friendly than Stata requiring statistical jargons as the input. Finally, Stata does not allow the specification of an exogenous variable as categorical (nominal or ordinal), so any such variable should be first dummy-coded to create indicators of categories which are subsequently treated as continuous data.

Finally, compared with all the other programmes previously discussed, PROC CALIS in SAS offers only limited features in handling OCD variables in SEM. The instructions on using these features are found only in a usage note on the SAS company website but not in the official documentation of PROC CALIS (SAS Institute, 2017, 2018). In the usage note, two solutions are provided. First, PROC CALIS provides two estimators, *METHOD=WLS* and *METHOD=MLM*, which are capable of producing asymptotically unbiased parameter estimates and adjusting the standard errors for non-normal data (again, non-normal data are not necessarily categorical). The two estimators can work with a dataset including dichotomous or ordinal data regardless of whether there are also continuous variables. However, the note does not indicate PROC CALIS allows an OCD endogenous variable to be analysed as a generalised linear model (e.g., probit regression, already a built-in feature in the other SEM programmes). Second, PROC CALIS could also accept the input of a covariance matrix that is based on the polychoric and polyserial correlations from OCD variables. This covariance matrix is treated as if it were the covariance matrix from continuous variables and thus could be analysed under any estimation method. However, the note does not specifically indicate how to compute this covariance matrix from the polychoric or the polyserial correlation matrix. A further conversation with the SAS Technical Support team suggested the covariance matrix was just the correlation matrix (i.e., correlations treated as covariances in the input covariance matrix for PROC CALIS to analyse) (SAS Technical Support Statistics, personal communication, October 5, 2021).

5 Numeric examples

5.1 Research context

This dataset used in the two demonstrations comes from the 2010 HRS. The HRS is a panel study of middle-aged and older Americans and provides both health-related outcomes (e.g., lifestyle, chronic conditions) and financial status variables (e.g., income, wealth). In 2010, the HRS included an additional module of 16 items measuring health literacy (*HealthLiteracy*). With a composite dichotomous health literacy measure created

from the 16 items (Brockett et al., 2002), it was possible to investigate the effect of health literacy on retirement worth. Out of a total of 1,308 participants asked to respond to the special module, 1,069 of them completed the assessment and their responses are analysed in this study. Table 1 outlines the variables used in the demonstrations, particularly whether an OCD is actually coded as *ordered-categorical* in AMOS. Notably, the OCD exogenous variable *HealthLiteracy* is dichotomous consisting of 0's and 1's. This variable is used as a continuous variable. For one, to stay consistent with the other SEM programmes. As has been discussed above, other software programmes (e.g., Mplus) could run into issues when an exogenous variable is declared as categorical. For another, when this variable is specified as *ordered-categorical* in AMOS, the Bayesian algorithm sometimes has difficulty getting past the pre-burn-in period. Finally, AMOS is set up to automatically decide the boundaries of each ordered categorical variable.

Table 1 Summary of variables in the conceptual model

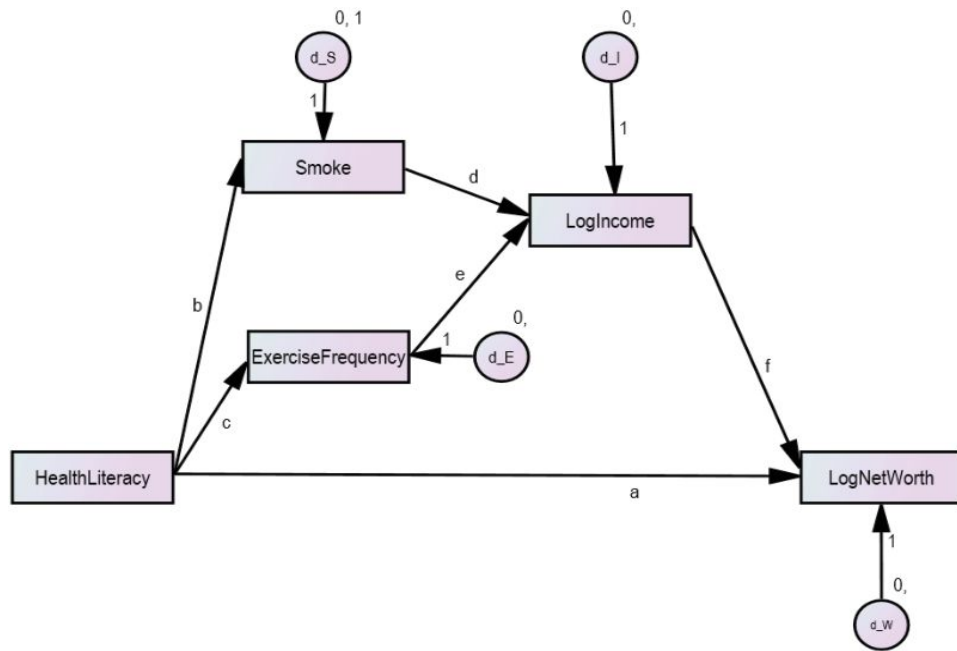
<i>Variable type</i>	<i>Variable name</i>	<i>Values</i>	<i>Number of categories, if applicable</i>	<i>Role in the model</i>	<i>Coded as ordered-categorical</i>
Categorical	HealthLiteracy	1 = High; 0 = Low	Two	Exogenous	No
Categorical	Smoke	1 = No; 0 = Yes	Two	Endogenous	Yes
Categorical	ExerciseFrequency	3 = Frequent; 2 = Occasional; 1 = Never	Three	Endogenous	Yes
Continuous	LogIncome	Household income on the logarithm scale	N/A	Endogenous	NA
Continuous	LogNetWorth	Retirement net worth (wealth) on the logarithm scale	N/A	Endogenous	NA
Categorical	Female	1 = Female; 0 = Male	Two	Grouping (MGA)	NA

Using the HRS data, two demonstrations are provided on the same conceptual model analysed, respectively, in a single- and multi-group analysis. The MGA consists of two sub-analyses, respectively, for females (MGA-F) and males (MGA-M). Figure 1 presents the conceptual model specified by subject matter experts. In the model, the two OCD endogenous variables (*ExerciseFrequency* and *Smoke*) are coded as ordinal categorical in AMOS. For the dichotomous OCD endogenous variable *Smoke*, an additional step is taken to fix its error variance at one to prevent this variable from causing the entire model to become under-identified (AMOS Development Corporation, 2021b; Grace, 2009). Both ordinal OCD endogenous variables are handled through the probit regression formulation. Next, for MGA only, the conceptual model is specified to have the same structure (i.e., the same set of paths) across the two groups created by the OCD variable representing gender: *Female*.

In both SGA and MGA, a default, diffuse prior distribution is specified for all model parameters. This default prior is a uniform distribution on the interval $[-3.4 \times 10^{-38}, 3.4 \times 10^{38}]$ which evidently spreads its probability over a very wide range of parameter

values. Because the prior distribution introduces very little information and the sample size is large relative to the number of model parameters (the ratio of sample size to number of model parameters being over 28 to 1 for MGA, and over 70 to 1 for SGA), the Bayesian analysis settings actually allow the data to speak for themselves.

Figure 1 Conceptual model in SGA and MGA (see online version for colours)



Notably, AMOS provides a total of three prior options. The first two options are specific distributions:

- 1 uniform distribution (default diffuse prior as described above)
- 2 normal distribution.

The parameters of both prior distributions are adjustable in the software. The third option is a customisable prior which allows the user to create/draw a ‘freehand’, user-specified prior distribution (AMOS Development Corporation, n.d.). Given multiple priors, a decision on which prior to use and how to specify its parameters should be made as a part of invoking the Bayesian estimator in AMOS. As a general discussion on analysing categorical data in SEM, this study does not assume any prior knowledge that can serve to make the prior non-diffuse (i.e., informative). Therefore, the use of a diffuse prior as described above is appropriate here in that it allows the Bayesian SEM platform to be invoked so that the model containing the specified OCD variables can be estimated in AMOS without bringing in the complications from incorporating specific prior information into the model. Certainly, under a different research context, there could well be prior knowledge that should be properly addressed through the selection of the right prior distribution and the appropriate specification of its parameters, a critically important topic in Bayesian statistics. Because the topic is beyond the scope of the study, the readers are referred to the following studies for additional, up-to-date instructions:

Arbuckle (2019, p.405), Depaoli (2021, pp.34–42), Gelman et al. (2014, pp.102–104), Gill (2015, pp.97–143), Jackman (2009, pp.15–18, pp.80–97), Kaplan (2014, pp.17–22), and Song and Lee (2012, pp.35–40).

Next, to run Bayesian SEM in both SGA and MGA, the RWM method (random seed = 103,828,397, tuning parameter $\alpha = 0.70$) is used to approximate the posterior distribution. Following Skrondal and Rabe-Hesketh (2004, p.213), the upper limit of 25,000 burn-in observations is taken to maximise the chance of stabilisation of the (only) chain and convergence on the posterior distribution. After the burn-in process ends, the sampling continues and is monitored for the convergence of MCMC. AMOS provides a convergence statistic (CS) for the overall model and for each individual parameter. The CS is a modified version of the potential scale reduction factor (PSRF) statistic by Gelman et al. (2014, p.285). Arbuckle (2019) provides a threshold of 1.002 for the CS statistic for assessing convergence. Using the CS and several other measures, the study assesses the convergence of each model in the following way:

- 1 the overall and individual convergence statistics each drop to below 1.002 and become stabilised there
- 2 the posterior probability density plots for individual parameters are each approximating a normal density
- 3 the trace plots for individual parameters each exhibit a tight, horizontal band. No long term trend is identified in each trace plot
- 4 the autocorrelation plots for individual parameters each drop to close to zero (no higher than 0.15).

The graphics outlined above are not included here in the interest of space, but are available upon request. To learn more about assessing MCMC convergence and other Bayesian diagnostics, please also refer to the previously recommended readings on Bayesian statistics.

In the end, the two models are also analysed under three of the four comparison software programmes to assess the consistency of parameter estimates across different software programmes:

- 1 Mplus
- 2 R lavaan
- 3 Stata.

However, this study does not use SAS PROC CALIS to estimate the models because of its limitations discussed above.

5.2 *Bayesian estimation results and interpretation*

Tables 2 and 3 present the SGA and MGA posterior summaries of unstandardised parameter estimates. Following Kaplan (2014), interpreted as the Bayesian point estimate is each posterior mean (i.e., expected a posteriori or EAP). Also included in both tables are posterior standard deviation, 95% posterior probability interval (PPI) signifying a 95% probability that the effect falls into its lower and upper limits and the probability

estimate for Bayesian hypothesis testing (Bayesian p-value) that the effect is less than or equal to zero [Arbuckle, (2019), p.404; Gill, (2007), p.237].

Table 2 Posterior summary of paths from single group analysis

<i>Paths</i>	<i>EAP¹</i>	<i>S.D.</i>	<i>95% PPI lower</i>	<i>95% PPI upper</i>	<i>Bayesian p-value</i>
HealthLiteracy → LogNetWorth (a)	1.486	0.381	0.735	2.235	0.000
HealthLiteracy → Smoke (b)	0.240	0.124	−0.004	0.479	0.027
HealthLiteracy → ExerciseFrequency (c)	0.577	0.184	0.220	0.943	0.001
Smoke → LogIncome (d)	0.167	0.059	0.047	0.279	0.004
ExerciseFrequency → LogIncome (e)	0.143	0.028	0.089	0.200	0.000
LogIncome → LogNetWorth (f)	0.969	0.098	0.776	1.161	0.000

Notes: ¹EAP = expected A posteriori.

Results obtained using the random walk MCMC provided in SPSS AMOS 25 with a tuning parameter of 0.70, a random seed of 103,828,397, a burn-in sample of 25,000 observations, a posterior sample of 27,920 observations and a thinning factor of 4.

Table 3 Posterior summary of paths from multi-group analysis (MGA-F/MGA-M)

<i>Paths</i>	<i>EAP¹</i>	<i>S.D.</i>	<i>95% PPI lower</i>	<i>95% PPI upper</i>	<i>Bayesian p-value</i>
HealthLiteracy → LogNetWorth (a)	0.885/2.067	0.562/0.538	−0.217/1.014	1.996/3.120	0.058/0.000
HealthLiteracy → Smoke (b)	0.253/0.207	0.172/0.185	−0.087/−0.164	0.590/0.562	0.071/0.131
HealthLiteracy → ExerciseFrequency (c)	0.672/0.636	0.238/0.300	0.217/0.059	1.155/1.244	0.002/0.016
Smoke → LogIncome (d)	0.158/0.178	0.077/0.087	0.001/0.001	0.305/0.342	0.024/0.024
ExerciseFrequency → LogIncome (e)	0.137/0.134	0.039/0.039	0.063/0.058	0.216/0.214	0.000/0.000
LogIncome → LogNetWorth (f)	1.079/0.909	0.180/0.116	0.723/0.679	1.433/1.136	0.000/0.000

Notes: ¹EAP = expected A posteriori.

Results obtained using the random walk MCMC provided in SPSS AMOS 25 with a tuning parameter of 0.70, a random seed of 103,828,397, a burn-in sample of 25,000 observations, a posterior sample of 29,837 observations and a thinning factor of 4.

In Table 2 outlining the SGA results, it is observed that the unstandardised EAP estimate for the direct effect of health literacy on net worth (path a) is 1.486, with a posterior standard deviation of 0.381. The EAP is positive, indicating that a higher value on health literacy is associated with a higher net worth. The 95% PPI indicates there is a 95% probability that the direct effect of health literacy on net worth is between 0.735 and 2.235. Next, based on an EAP of 0.240 and 0.577, people in the high health literacy group are more likely not to smoke but are more likely to exercise more frequently. The path for not smoking is associated with a 95% PPI, [−0.004, 0.479], that contains zero,

indicating zero is a credible value for this effect, but the probability is only 0.027 of this effect being equal to or less than zero. Next, for paths d and e, the EAP estimates of 0.167 and 0.143 suggest no smoking and more intensive exercise are expected to be associated with higher household income. Finally, the last path connecting income with retirement wealth indicates, as income increases, so does retirement wealth. Since the 95% PPI, [0.776, 1.161], for this path does not contain zero, zero is not a credible value for this effect. Based on the Bayesian p-value, the probability is (extremely close to) zero of this path being less than or equal to zero.

Table 4 Posterior summary of differences in paths from multi-group analysis (difference = MGA-F – MGA-M)

<i>Paths</i>	<i>EAP¹</i>	<i>S.D.</i>	<i>95% PPI lower</i>	<i>95% PPI upper</i>	<i>Bayesian p-value</i>
HealthLiteracy → LogNetWorth (a)	–1.182	0.776	–2.717	0.330	0.936
HealthLiteracy → Smoke (b)	0.046	0.252	–0.448	0.540	0.429
HealthLiteracy → ExerciseFrequency (c)	0.036	0.384	–0.720	0.784	0.460
Smoke → LogIncome (d)	–0.020	0.116	–0.247	0.210	0.570
ExerciseFrequency → LogIncome (e)	0.004	0.056	–0.106	0.113	0.474
LogIncome → LogNetWorth (f)	0.170	0.213	–0.246	0.588	0.212

Notes: ¹EAP = expected A posteriori.
Results obtained using the random walk MCMC provided in SPSS AMOS 25 with a tuning parameter of 0.70, a random seed of 103,828,397, a burn-in sample of 25,000 observations, a posterior sample of 29,837 observations and a thinning factor of 4.

In Tables 3 and 4 outlining the MGA results, they focus on assessing if the paths are different across groups (difference = MGA-F – MGA-M), which is essentially a test of the moderation effects of gender on direct and indirect effects in the conceptual model (Wang and Preacher, 2015). Notably, the posterior distributions for between-group parameter differences are not directly available from AMOS. So, user-defined estimands (*Custom Estimands*) are coded in AMOS to obtain those posterior distributions for Bayesian statistical inference. Table 4 provides a posterior summary of these parameter differences. Because all 95% PPIs contain zero as a credible value, the relationships from these paths largely remain invariant (i.e., absence of the moderation effects of gender) across the two groups. For path a only, its Bayesian p-value indicates the probability is as high as 0.936 that this path is less than or equal to zero, suggesting the strength of this unstandardised path for the female group is very likely to be less than that for the male group [Kline, (2016), p.395].

Together with unstandardised estimates, also computed are standardised estimates and several functions of unstandardised and standardised estimates. Many times, these computations have to be based on user-defined, custom estimands in AMOS. Sample AMOS script in the form of Visual Basic code for performing these computations is presented in two appendices:

- 1 Appendix A for SGA
- 2 Appendix B for MGA comparing groups of males and females.

5.3 Comparison with Mplus, R lavaan and Stata

The two models are also analysed using Mplus, R lavaan, and Stata to assess how consistent the estimates are across AMOS and these comparison programmes. Examined here are the unstandardised estimates of nine parameters under three different types from the four software programmes (for AMOS, it is only the EAP estimates from the Bayesian results):

- 1 five paths (a, b, c, d, and e)
- 2 two intercepts of the continuous endogenous variables (int_W for LogNetWorth, and int_I for LogIncome)
- 3 two variances of the disturbances of the two continuous endogenous variables (d_W for LogNetWorth, and d_I for LogIncome).

First, the pairwise Pearson correlations between the four sets of parameter estimates from the four software programmes demonstrate the parameter estimates exhibit a highly similar pattern across the programmes. Under SGA, the correlations are all as high as 0.999. Under MGA-F, the correlations range from 0.996 to 0.999, and under MGA-M, the correlations range from 0.991 to 0.999.

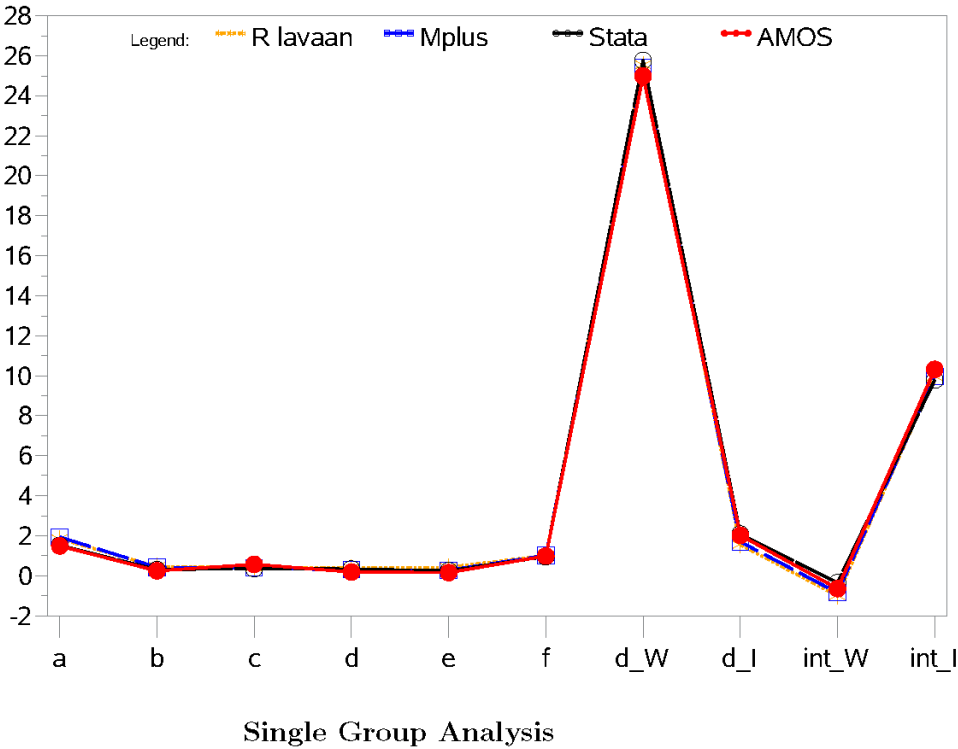
Second, Figure 2 presents an overlay of the four line charts from the SGA documenting the parameter estimates. Figure 3 presents the same set of estimates under MGA: MGA-F [Figure 3(a)] and MGA-M [Figure 3(b)]. Based on Figure 2, the four software programmes are pretty close in terms of these parameter estimates, as is evidenced by the substantial overlap of the four line charts across the nine parameters. Based on Figure 3, a similar conclusion may be drawn on most of the nine parameters. However, evident differences are found between AMOS and the other three comparison programmes on three of the nine parameters:

- 1 regression path a
- 2 disturbance variance d_W
- 3 intercept int_W.

Third, separately under SGA, MGA-F and MGA-M, the study computes the absolute value differences in parameter estimates between each pair of software programmes, leading to six variables representing these absolute value differences. Sequentially in two steps, the study further summarises the computed differences using four descriptive statistics (means, standard deviations, minimums and maximums) when taking into account parameter type (i.e., paths, intercepts and variances) and software comparison in pairs.

In step one, within each of the six absolute value difference variables, the four descriptive statistics are computed within each parameter type, therefore resulting in four descriptive statistics summarising each of the three types of parameters under each of the six difference variables. These four descriptive statistics within each combination of parameter type and difference variable are next further summarised.

Figure 2 Single group analysis results under R, Mplus, Stata and AMOS (see online version for colours)



In step two, two separate approaches are taken on the descriptive statistics from step one. In approach one, within each parameter type, each of the four descriptive statistics is averaged across all six difference variables to arrive at four means of the descriptive statistics (i.e., four means for each of the three parameter types: a total of $4 * 3 = 12$ means). In approach two, within each difference variable, each of the four descriptive statistics is averaged across all three parameter types to arrive at another four means of the descriptive statistics (i.e., four means for each of the six difference variables: a total of $4 * 6 = 24$ means).

The outlined process is conducted separately for the three analyses of SGA, MGA-F, and MGA-M. The $4 * 3 = 12$ means from each analysis under approach one are presented using line charts in Figure 4 and the $4 * 6 = 24$ means from each analysis under approach two presented in Figure 5. In both figures, the means are in the vertical axis. The horizontal axis in Figure 4 represents parameter type, and that in Figure 5 software pair under comparison.

Figure 4 demonstrates similar patterns in three of the four subfigures:

- 1 means (top left)
- 2 standard deviations (top right)
- 3 maximums (bottom right).

Figure 3 Multi-group analysis results under R, Mplus, Stata and AMOS, (a) female group (MGA-F) (b) male group (MGA-M) (see online version for colours)

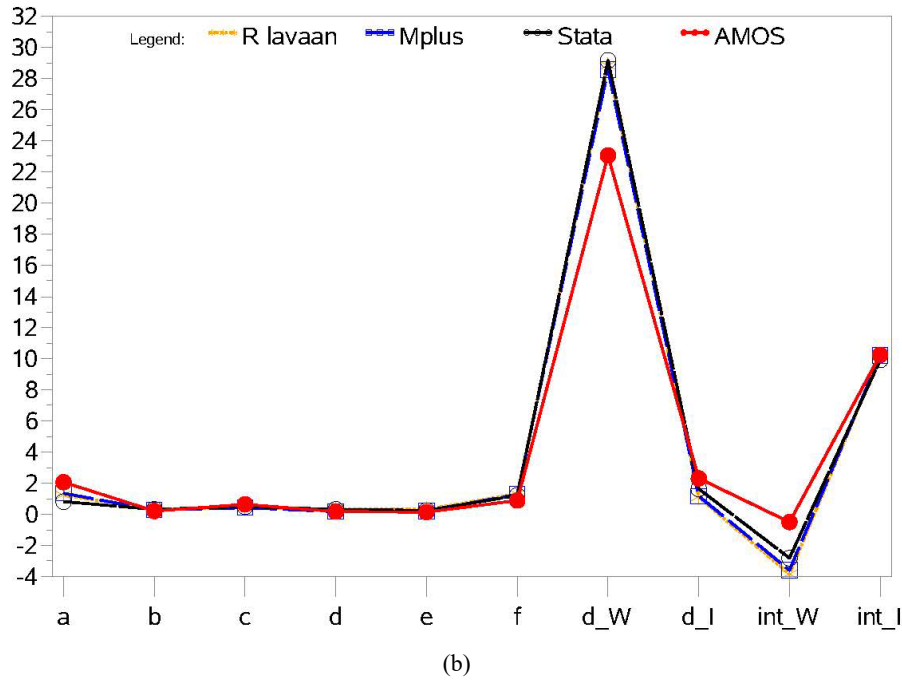
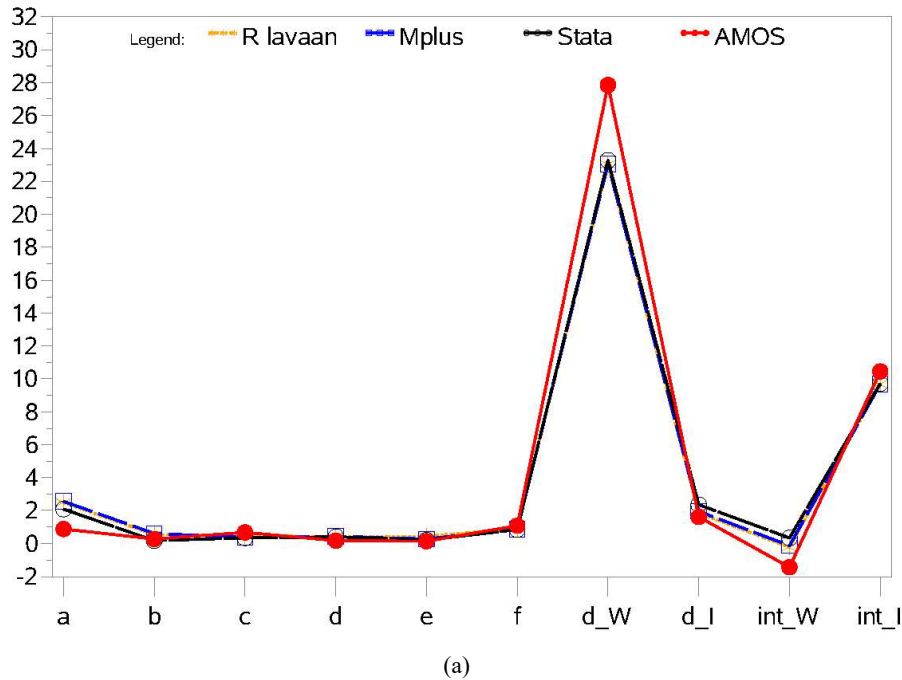
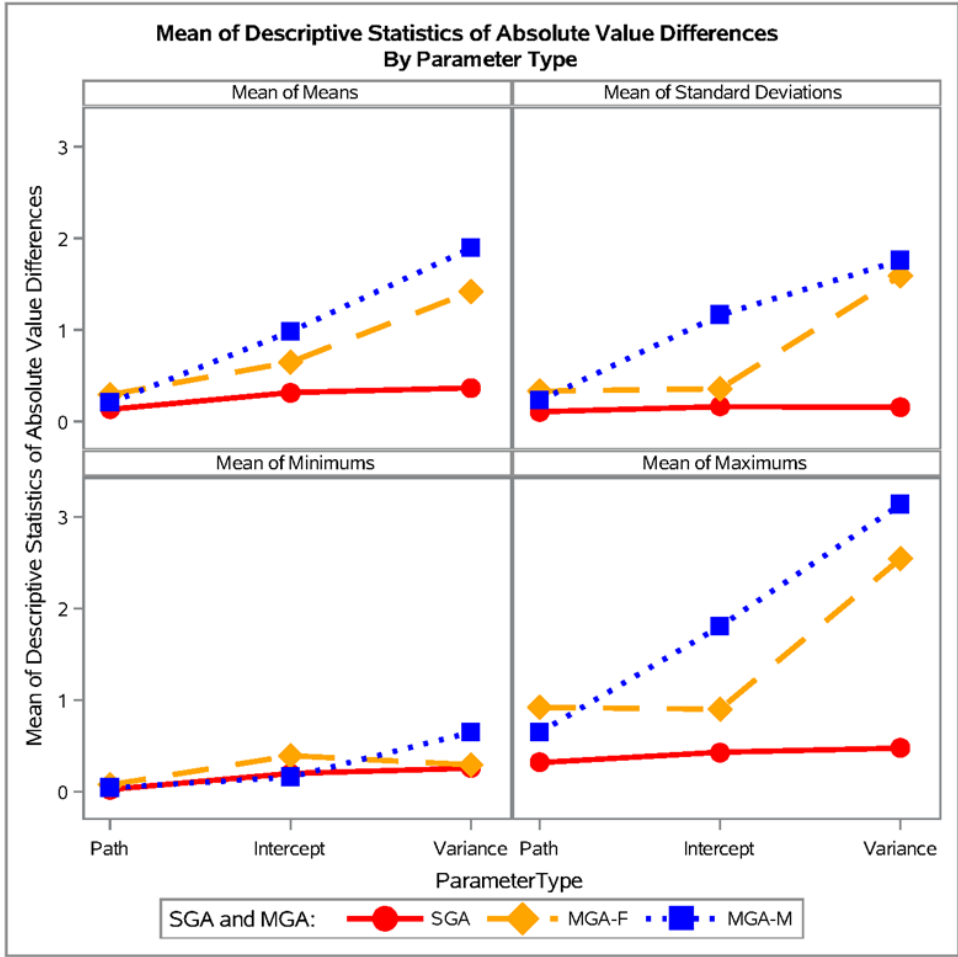


Figure 4 Mean of descriptive statistics of absolute value differences by parameter type (see online version for colours)



Taking as an example the subfigure for means (i.e., mean discrepancies in the absolute value differences between the software programmes), the SGA line chart is almost completely horizontal running through the zero of the vertical axis, suggesting that the mean discrepancies in the absolute value differences between the software programmes tend to be zero and this pattern is true of all three parameter types. Next, the two MGA line charts deviate noticeably away from being horizontal, the MGA-M chart in particular. Therefore, under MGA, the mean discrepancies in the absolute value differences between the software programmes do not tend to be consistent across the three parameter types: on average, minimum, intermediate and maximum mean discrepancies exhibited, respectively, in paths, intercepts, and variances. Next, similar patterns are found in the subfigures for standard deviations and maximums. Finally, the line charts for the minimum discrepancies in the absolute value differences between the software programmes demonstrate that the minimum discrepancies tend to be zero, and

this pattern is largely consistent across all three parameter types and across all three analyses.

Figure 5 Mean of descriptive statistics of absolute value differences by software pair (see online version for colours)

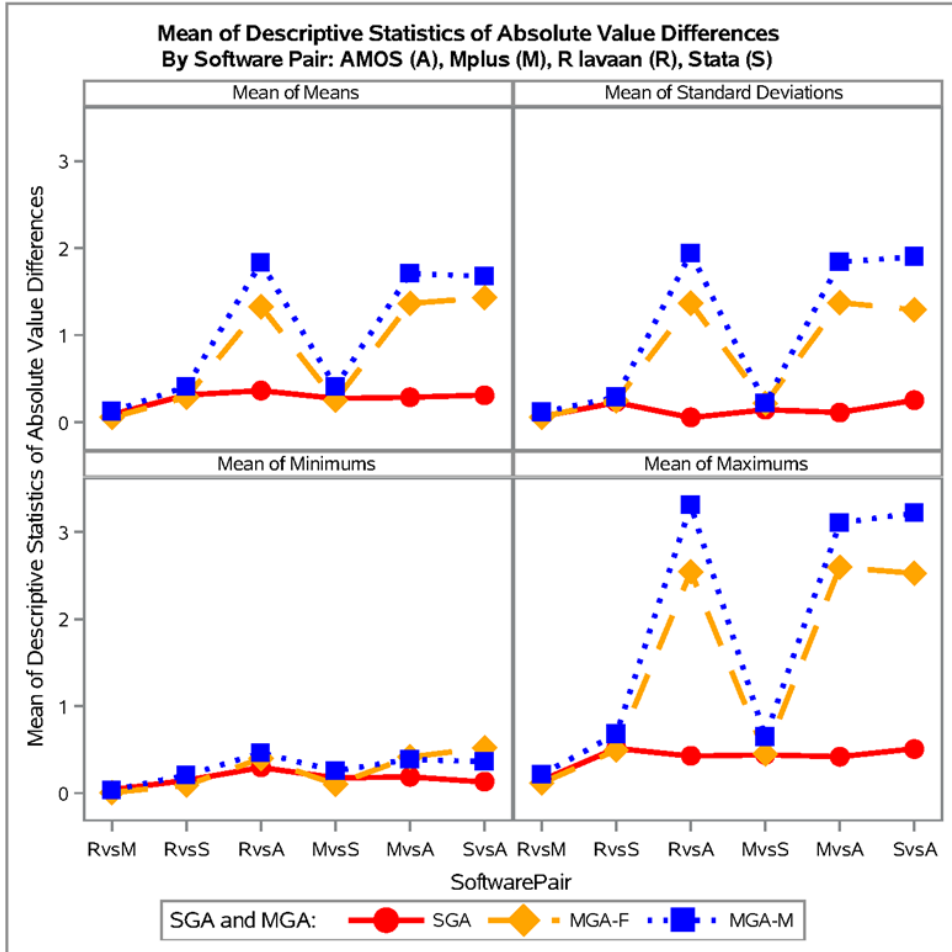


Figure 5 can be interpreted in a way similar to Figure 4 because the patterns in its four subfigures are similar to their counterparts in Figure 4. Taking as an example the subfigure for means (i.e., mean discrepancies in the absolute value differences between the software programmes), the SGA line chart is almost completely horizontal running through the zero of the vertical axis, suggesting that the mean discrepancies in the absolute value differences between the software programmes tend to be zero and this pattern is true of all six software comparisons. Next, the two MGA line charts overlap with the SGA chart on three of the six software comparisons not involving AMOS, but deviate noticeably away from the SGA chart on the other three comparisons involving AMOS, the MGA-M chart in particular. Therefore, under MGA, the mean discrepancies in the absolute value differences between the software programmes do not tend to be consistent across the six software comparisons: on average, smaller mean discrepancies

in comparisons not involving AMOS than in comparisons involving AMOS. Next, similar patterns are found in the subfigures for standard deviations and maximums. Finally, the line charts for the minimum discrepancies in the absolute value differences between the software programmes demonstrate that the minimum discrepancies tend to be zero, and this pattern is largely consistent across all six software comparisons and across all three analyses.

6 Discussion

This study provides a didactic demonstration of the features of AMOS for handling OCD variables under both SGA and MGA, compares AMOS with several other SEM software programmes (Mplus, R lavaan, Stata and SAS PROC CALIS) on these features, and also assesses the consistency of parameter estimates across AMOS and three of the four comparison programmes. An OCD variable features non-continuous data representing qualitative categories. Among other things, the treatment of such data depends primarily on whether the categorical variable is endogenous or non-endogenous (exogenous or creating groups) in the model and varies from one SEM software programme to another.

When an OCD variable is endogenous, AMOS allows it to be first coded as ordinal categorical to prevent the software from treating it as continuous. In the programme, whenever a variable is specified as ordinal categorical, Bayesian SEM is the only estimation option available for estimating the model. If the categories of the OCD endogenous variable are nominal without any meaningful ordering, AMOS unfortunately cannot handle such data as of this writing. Besides, incorrectly treating nominal data as ordinal is likely to lead to inflated standard errors and misleading results [Lee et al., (2010), pp.294–296].

When an OCD variable is non-endogenous but exogenous, AMOS allows it to be handled in at least three different ways. First, given an OCD exogenous variable consisting of ($k \geq 2$) categories, a usual practice is to replace it with a set of ($k - 1$) dummy variables, which is the same as the handling of a categorical predictor in a classic multiple linear regression model. For example, if the exogenous variable is dichotomous consisting of ($k = 2$) categories, it is to be replaced with one ($= (k - 1)$) dummy variable in the form of 0's and 1's. Second, if the OCD exogenous variable is ordinal, a coding scheme reflecting the ordering of categories may be used (e.g., 1, 2, 3, ...) before next treating the ordinal variable in the same way as any other continuous exogenous variables. Third, AMOS also allows the specification of an OCD exogenous variable as ordinal categorical. However, in this study, doing so led to difficulty in getting the MCMC algorithm to move beyond the pre-burn-in period and also an extended amount of time to reach convergence.

When an OCD variable is non-endogenous but represents groups of comparison, AMOS provides a multiple-group analysis (e.g., MGA for assessing measurement invariance) option. In this case, unless there are other conditions necessitating the use of Bayesian SEM (e.g., OCD variable specified as ordinal categorical), both Bayesian and frequentist (e.g., ML) estimation methods can be used.

When comparing AMOS with Mplus, R lavaan and Stata, the study focuses on the consistency of parameter estimates as measured by the absolute value differences in parameter estimates from each pair of software programmes under SGA and MGA. Overall, the consistency between software programmes tends to be better in SGA than in

either of MGA-F and MGA-M, which might have been due to the fact that the sample sizes in MGA are each much smaller than that in SGA. On the other hand, the consistency of parameter estimates among the three comparison programmes tends to be better than that between AMOS and each comparison programme. However, the study is in no position to argue one software programme produces more accurate estimates than another because it is limited to just one real world dataset with no more than a handful of applications where the true parameter values are always unknown. For an in-depth investigation of the accuracy of parameter estimates, the study calls for large scale simulation studies under varying research contexts (e.g., varying model and/or data structures).

Besides, for full transparency, the estimation methods used in each of the three demonstrated comparison programmes (Mplus, R lavaan and Stata) are shared here. In Mplus, with OCD endogenous data in the model, it is a robust WLS estimator under a diagonal weight matrix that is used in deriving the parameter estimates. This estimator is invoked by using the *ESTIMATOR* option of the *ANALYSIS* command: *ESTIMATOR = WLSMV* (Muthén and Muthén, 2017). Next, the model estimation in R lavaan is based on a WLS estimator as well. Specifically, R lavaan utilises a diagonally weighted least squares (DWLS) estimator to obtain the estimates of model parameters. After (all or part of) the endogenous data are declared as categorical in R lavaan, this estimator is the default setting and thus automatically applied, but it can also be specified using *estimator = "DWLS"* when calling one of the model estimation functions (e.g., *lavaan::cfa()*, *lavaan::sem()* or *lavaan::lavaan()* function) (Rosseel, 2012, 2020). In the end, in Stata, after incorporating OCD endogenous variables into the model through the generalised linear modelling framework, a ML estimator is used in estimating the model. Since the estimator is the default and also the only estimator provided in the *gsem* command, it is automatically invoked, but can also be specified using *method(ml)* in the command (StataCorp, 2021).

Finally, when presenting the AMOS features for handling OCD, the study takes the opportunity to familiarise the intended readership of the study with the typical ways of summarising, reporting and interpreting the results of Bayesian statistical analyses including the appropriate language to use and references to cite from the recent literature of Bayesian statistics primarily in the social and behavioural sciences. It is hoped that the knowledge will contribute to demystifying Bayesian statistical inference and encourage its understanding and use among applied investigators and practitioners to solve real world problems.

References

- Agresti, A. (2013) *Categorical Data Analysis*, 3rd ed., John Wiley & Sons, Inc., Hoboken, NJ, USA.
- AMOS Development Corporation (2021a) *How the MCMC Algorithm Works [Computer Software Manual]* [online] <http://amosdevelopment.com/webhelp/7653.html> (accessed 17 September 2022).
- AMOS Development Corporation (2021b) *Parameter Identification with Dichotomous Variables [Computer Software Manual]* [online] <http://amosdevelopment.com/webhelp/7994.html> (accessed 17 September 2022).

- AMOS Development Corporation (n.d.) *An Introduction to Bayesian Estimation with AMOS*, Video, AMOS Development [online] <http://amosdevelopment.com/video/bayesian/intro/flash/intro.html> (accessed 17 September 2022).
- Arbuckle, J.L. (2019) *IBM SPSS AMOS 26 User's Guide*, AMOS Development Corporation, Chicago, IL, USA.
- Betancourt, M. (2018) *A Conceptual Introduction to Hamiltonian Monte Carlo* [online] <https://arxiv.org/pdf/1701.02434.pdf> (accessed 17 September 2022).
- Bhardwaj, R. (2015) *Re: SEM with Categorical Variable/Parceling/How to Enter in AMOS? [Online Discussion Group]* [online] <http://www.talkstats.com/threads/sem-with-categorical-variable-parceling-how-to-enter-in-amos.55030/page-2> (accessed 17 September 2022).
- Boateng, S.L. (2020) *Structural Equation Modelling (SEM) Made Easy for Business and Social Science Research using SPSS and AMOS*, 2nd ed., Kindle Direct Publishing, Seattle, WA, USA.
- Brockett, P.L., Derrig, R.A., Golden, L.L., Levine, A. and Alpert, M. (2002) 'Fraud classification using principal component analysis of RIDITs', *The Journal of Risk and Insurance*, Vol. 69, No. 3, pp.341–371, <https://doi.org/10.1111/1539-6975.00027>.
- Chang, M-I. and Sheng, Y. (2017) 'A comparison of two MCMC algorithms for the 2PL IRT model', in van der Ark, L.A., Wiberg, M., Culppepper, S.A., Douglas, J.A. and Wang, W.C. (Eds.): *Quantitative Psychology. IMPS 2016. Springer Proceedings in Mathematics & Statistics*, Vol. 196, Springer, https://doi.org/10.1007/978-3-319-56294-0_7.
- Cristofaro, M. (2016a) *Categorical Data in SEM Using AMOS. How to Code and Interpret Results?* [online] <https://stats.stackexchange.com/questions/192908/categorical-data-in-sem-using-amos-how-to-code-and-interpret-results> (accessed 17 September 2022).
- Cristofaro, M. (2016b) *Categorical Data in SEM Using AMOS. Is it Reliable the Bayesian Estimation? How Can I Interpret Results?* [online] <https://www.researchgate.net/post/Categorical-data-in-SEM-using-AMOS-Is-it-reliable-the-Bayesian-estimation-How-can-I-interpret-results> (accessed 17 September 2022).
- Depaoli, S. (2021) *Bayesian Structural Equation Modeling*, The Guilford Press, NYC, NY, USA.
- Fox, J-P. (2010) *Bayesian Item Response Modeling: Theory and Applications*, Springer, NYC, NY, USA.
- Gelman, A., Carlin, J.B., Stern, H.S., Dunson, D.B., Vehtari, A. and Rubin, D.B. (2014) *Bayesian Data Analysis*, 3rd ed., CRC Press, Boca Raton, FL, USA.
- Gill, J. (2007) *Bayesian Methods: A Social and Behavioral Sciences Approach*, 2nd ed., CRC Press, Boca Raton, FL, USA.
- Gill, J. (2015) *Bayesian Methods: A Social and Behavioral Sciences Approach*, 3rd ed., CRC Press, Boca Raton, FL, USA.
- Gill, J. and Walker, L.D. (2005) 'Elicited priors for Bayesian model specifications in political science research', *The Journal of Politics*, Vol. 67, No. 3, pp.841–872, <https://doi.org/10.1111/j.1468-2508.2005.00342.x>.
- Gill, J. and Witko, C. (2013) 'Bayesian analytical methods: a methodological prescription for public administration', *Journal of Public Administration Research and Theory*, Vol. 23, No. 2, pp.457–494, <https://doi.org/10.1093/jopart/mus091>.
- Grace, J.B. (2009) *SE Modeling when Some Response Variables are Categorical: The Special Case of Binary (Dichotomous) Variables* [online] http://www.structuralequations.com/resources/BinaryResponseModeling.Mar30_2009.pps (accessed 17 September 2022).
- Gupta, S. (2016) *Can Categorical Variable be Used as Dependent Variable in SEM? If Yes Which Software AMOS, MPlus, Liserl Can Be Used for it* [online] https://www.researchgate.net/post/Can_Categorical_Variable_be_used_as_Dependent_Variable_in_SEM.If_yes_which_software_AMOS_MPlus_Liserl_can_be_used_for_it#:~:text=Yes%2C%20Mplus%20can%20do%20that.&text=Mplus%20is%20a%20extensive%20structural,your%20analysis%20with%20this%20software (accessed 17 September 2022).

- Hadi, N.U. (2015) *Any Suggestion for Categorical Variables (Nominal, Ordinal, and Dichotomous) Analysis in CB-SEM or VB-SEM?* [online] <https://www.researchgate.net/post/Any-suggestion-for-categorical-variables-Nominal-ordinal-and-dichotomous-analysis-in-CB-SEM-or-VB-SEM> (accessed 17 September 2022).
- IBM SPSS (2018) *AMOS Bayesian Estimation Triggers "Waiting to Accept a Transition..." Warning* [online] <https://www.ibm.com/support/pages/amos-bayesian-estimation-triggers-waiting-accept-transition-warning> (accessed 17 September 2022).
- IBM SPSS (2020) *Binary Variables in AMOS* [online] <https://www.ibm.com/support/pages/binary-variables-amos> (accessed 17 September 2022).
- Jackman, S. (2009) *Bayesian Analysis for the Social Sciences*, Wiley, Chichester, West Sussex, UK.
- Kaplan, D. (2014) *Bayesian Statistics for the Social Sciences*, The Guilford Press, NYC, NY, USA.
- Katsikatsou, M. (n.d.) *The Pairwise Likelihood Method for Structural Equation Modelling with Ordinal Variables and Data with Missing Values using the R Package Lavaan* [online] https://users.ugent.be/~yrosseel/lavaan/pml/PL_Tutorial.pdf (accessed 17 September 2022).
- Kline, R.B. (2016) *Principles and Practice of Structural Equation Modeling*, 4th ed., The Guilford Press, NYC, NY, USA.
- Kruschke, J.K. (2014) *Doing Bayesian Data Analysis: A Tutorial with R, JAGS, and Stan*, 2nd ed., Academic Press, Boston, MA, USA.
- Kuo, T-C. and Sheng, Y. (2015) 'Bayesian estimation of a multi-unidimensional graded response IRT model', *Behaviormetrika*, Vol. 42, No. 2, pp.79–94, <https://doi.org/10.2333/bhmk.42.79>.
- Kuo, T-C. and Sheng, Y. (2016) 'A comparison of estimation methods for a multi-unidimensional graded response IRT model', *Frontiers in Psychology*, Vol. 7, p.880, <https://doi.org/10.3389/fpsyg.2016.00880>.
- Kuo, T-C. and Sheng, Y. (2017) 'Fitting graded response models to data with non-normal latent traits', In van der Ark, L.A., Wiberg, M., Culppepper, S.A., Douglas, J.A. and Wang, W.C. (Eds.): *Quantitative Psychology. IMPS 2016. Springer Proceedings in Mathematics & Statistics*, Vol. 196, Springer, https://doi.org/10.1007/978-3-319-56294-0_4.
- Kutner, M.H., Nachtsheim, C.J., Neter, J. and Li, W. (2005) *Applied Linear Statistical Models*, 5th ed., McGraw-Hill, NYC, NY, USA.
- Lee, S-Y. (2007) *Structural Equation Modeling: A Bayesian Approach*, John Wiley & Sons, Chichester, West Sussex, UK.
- Lee, S-Y. and Song, X-Y. (2003) 'Bayesian analysis of structural equation models with dichotomous variables', *Statistics in Medicine*, Vol. 22, pp.3073–3088, <https://doi.org/10.1002/sim.1544>.
- Lee, S-Y., Song, X-Y. and Cai, J-H. (2010) 'A Bayesian approach for nonlinear structural equation models with dichotomous variables using logit and probit links', *Structural Equation Modeling: A Multidisciplinary Journal*, Vol. 17, pp.280–302, <https://doi.org/10.1080/10705511003659425>.
- Levy, R. and Mislevy, R.J. (2016) *Bayesian Psychometric Modeling*, CRC Press, Boca Raton, FL, USA.
- Li, C-H. (2021) 'Statistical estimation of structural equation models with a mixture of continuous and categorical observed variables', *Behavior Research Methods*, Vol. 53, pp.2191–2213, <https://doi.org/10.3758/s13428-021-01547-z>.
- Lynch, S.M. (2007) *Introduction to Applied Bayesian Statistics and Estimation for Social Scientists*, Springer, NYC, NY, USA.
- Metropolis, N., Rosenbluth, A.W., Rosenbluth, M.N., Teller, A.H. and Teller, E. (1953) 'Equations of state calculations by fast computing machines', *The Journal of Chemical Physics*, Vol. 21, No. 6, pp.1087–1092, <https://doi.org/10.1063/1.1699114>.
- Muthén, L.K. and Muthén, B.O. (2017) *Mplus User's Guide*, 8th ed. [online] https://www.statmodel.com/download/usersguide/MplusUserGuideVer_8.pdf (accessed 17 September 2022).

- Neal, R.M. (2011) 'MCMC using Hamiltonian dynamics', in Brooks, S., Gelman, A., Jones, G. and Meng, X.-L. (Eds.): *Handbook of Markov Chain Monte Carlo*, pp.113–162, Hall/CRC, Chapman.
- Olsson, U. (1979) 'Maximum likelihood estimation of the polychoric correlation coefficient', *Psychometrika*, Vol. 44, No. 4, pp.443–460, <https://doi.org/10.1007/BF02296207>.
- Perera, H.N. (2013) 'A novel approach to estimating and testing specific mediation effects in educational research: explication and application of Macho and Ledermann's (2011) phantom model approach', *International Journal of Quantitative Research in Education*, Vol. 1, No. 1, pp.39–60, <https://doi.org/10.1504/IJQRE.2013.055640>.
- Peterson, E.R., Brown, G.T.L. and Hamilton, R.J. (2013) 'Cultural differences in tertiary students' conceptions of learning as a duty and student achievement', *International Journal of Quantitative Research in Education*, Vol. 1, No. 2, pp.167–181, <https://doi.org/10.1504/IJQRE.2013.056462>.
- Rosseel, Y. (2012) Lavaan: an R package for structural equation modeling', *Journal of Statistical Software*, Vol. 48, No. 2, pp.1–36 [online] <http://www.jstatsoft.org/v48/i02/> (accessed 17 September 2022).
- Rosseel, Y. (2020) *The Lavaan Tutorial* [online] <http://lavaan.ugent.be/tutorial/tutorial.pdf> (accessed 17 September 2022).
- SAS Institute (2017) *SAS Usage Note 22529: Can PROC CALIS Analyze Categorical Data?* [online] <http://support.sas.com/kb/22/529.html> (accessed 17 September 2022).
- SAS Institute (2018) *SAS/STAT 15.1 User's Guide*, SAS Institute Inc. [online] <http://support.sas.com/documentation/onlinedoc/stat/151/introcalis.pdf> (accessed 17 September 2022).
- Sheng, Y. (2015) 'Bayesian estimation of the four-parameter IRT model using Gibbs sampling', *International Journal of Quantitative Research in Education*, Vol. 2, Nos. 3–4, pp.194–212, <https://doi.org/10.1504/IJQRE.2015.071736>.
- Sheng, Y. (2017a) 'Editorial: Fitting psychometric models: issues and new developments', *Frontiers in Psychology*, Vol. 8, p.856, <https://doi.org/10.3389/fpsyg.2017.00856>.
- Sheng, Y. (2017b) 'Investigating a weakly informative prior for item scale hyperparameters in hierarchical 3PNO IRT models', *Frontiers in Psychology*, Vol. 8, p.123, <https://doi.org/10.3389/fpsyg.2017.00123>.
- Skrondal, A. and Rabe-Hesketh, S. (2004) *Generalized Latent Variable Modeling: Multilevel, Longitudinal, and Structural Equation Models*, Chapman & Hall/CRC, Boca Raton, FL, USA.
- Song, X.Y. and Lee, S.Y. (2012) *Basic and Advanced Bayesian Structural Equation Modeling: With Applications in the Medical and Behavioral Sciences*, John Wiley & Sons, Chichester, West Sussex, UK.
- StataCorp (2021) *Stata Structural Equation Modeling Reference Manual Release 17*, StataCorp LLC [online] <https://www.stata.com/manuals/sem.pdf> (accessed 17 September 2022).
- van de Schoot, R., Kaplan, D., Denissen, J., Asendorpf, J.B., Neyer, F.J. and van Aken, M.A.G. (2014) 'A gentle introduction to Bayesian analysis: applications to developmental research', *Child Development*, Vol. 85, No. 3, pp.842–860, <https://doi.org/10.1111/cdev.12169>.
- Wagner, K. and Gill, J. (2005) 'Bayesian inference in public administration research: substantive differences from somewhat different assumptions', *Journal of Public Administration*, Vol. 28, pp.5–35, <https://doi.org/10.1081/PAD-200044556>.
- Wang, L. and Preacher, K.J. (2015) 'Moderated mediation analysis using Bayesian methods', *Structural Equation Modeling: A Multidisciplinary Journal*, Vol. 22, pp.249–263, <https://doi.org/10.1080/10705511.2014.935256>.
- Yorgason, J. (n.d.) *Troubleshooting Problems with SEM Models that have "Heywood" Cases such as Negative Variance Parameters and Non-Positive Definite Covariance Matrices* [online] <https://brightspotcdn.byu.edu/ba/fd/65505714418ebf6c0ad09f579b60/working-with-difficult-errors-in-sem.pptx> (accessed 17 September 2022).

Yuan, Y. and MacKinnon, D.P. (2009) 'Bayesian mediation analysis', *Psychological Methods*, Vol. 14, No. 4, pp.301–322, <https://doi.org/10.1037/a0016972>.

Appendix 1

AMOS Visual Basic Code for Single Group Analysis

#Region "Header"

Imports System

Imports Microsoft.VisualBasic

Imports AmosEngineLib

Imports AmosEngineLib.AmosEngine.TMatrixID

Imports AmosExtensions.CustomEstimand

Imports MiscAmosTypes

Imports MiscAmosTypes.cDatabaseFormat

#End Region

Public Class CEstimand

Implements IEstimand

Public Sub DeclareEstimands() Implements IEstimand.DeclareEstimands

 'Custom estimands for unstandardized estimates

 newestimand("a_unstd")

 newestimand("b_unstd")

 newestimand("c_unstd")

 newestimand("d_unstd")

 newestimand("e_unstd")

 newestimand("f_unstd")

 newestimand("path1direct_unstd")

 newestimand("path2indirect_Smoke_unstd")

 newestimand("path3indirect_ExerciseFrequency_unstd")

 newestimand("totalindirect_unstd")

 newestimand("total_unstd")

 newestimand("path1direct_lessthanorequaltozero_unstd")

 newestimand("path2indirect_Smoke_lessthanorequaltozero_unstd")

 newestimand("path3indirect_ExerciseFrequency_lessthanorequaltozero_unstd")

```
newestimand("totalindirect_lessthanorequaltozero_unstd")
newestimand("total_lessthanorequaltozero_unstd")
```

End Sub

Public Function CalculateEstimands(sem As AmosEngine) As String Implements
IEstimand.CalculateEstimands

‘Get Unstandardized Estimates.

```
estimand("a_unstd").value = sem.GetEstimate(DirectEffects, "LogNetWorth",  
"HealthLiteracy")
```

```
estimand("b_unstd").value = sem.GetEstimate(DirectEffects, "Smoke",  
"HealthLiteracy")
```

```
estimand("c_unstd").value = sem.GetEstimate(DirectEffects,  
"ExerciseFrequency", "HealthLiteracy")
```

```
estimand("d_unstd").value = sem.GetEstimate(DirectEffects, "LogIncome",  
"Smoke")
```

```
estimand("e_unstd").value = sem.GetEstimate(DirectEffects, "LogIncome",  
"ExerciseFrequency")
```

```
estimand("f_unstd").value = sem.GetEstimate(DirectEffects, "LogNetWorth",  
"LogIncome")
```

```
estimand("path1direct_unstd").value=estimand("a_unstd").value
```

```
estimand("path2indirect_Smoke_unstd").value=
```

```
estimand("b_unstd").value * estimand("d_unstd").value *  
estimand("f_unstd").value
```

```
estimand("path3indirect_ExerciseFrequency_unstd").value=
```

```
estimand("c_unstd").value * estimand("e_unstd").value *  
estimand("f_unstd").value
```

```
estimand("totalindirect_unstd").value=sem.GetEstimate(IndirectEffects,  
"LogNetWorth", "HealthLiteracy")
```

```
estimand("total_unstd").value=sem.GetEstimate(TotalEffects, "LogNetWorth",  
"HealthLiteracy")
```

‘Consistent with the null hypothesis (parameter ≤ 0) leading to Bayesian p
value (Gill, 2008, p. 237)

```
estimand("path1direct_lessthanorequaltozero_unstd").value=(estimand("path1direct  
_unstd").value<=0)
```

```
estimand("path2indirect_Smoke_lessthanorequaltozero_unstd").value=(estimand("p  
ath2indirect_Smoke_unstd").value<=0)
```

```

estimand("path3indirect_ExerciseFrequency_lessthanorequaltozero_unstd").value=(
estimand("path3indirect_ExerciseFrequency_unstd").value<=0)
estimand("totalindirect_lessthanorequaltozero_unstd").value=(estimand("totalindirect
_unstd").value<=0)
estimand("total_lessthanorequaltozero_unstd").value=(estimand("total_unstd").value
<=0)

```

Return ""Return an empty string if no error occurred

End Function

End Class

Appendix 2

AMOS visual basic code for multi-group analysis

#Region "Header"

Imports System

Imports Microsoft.VisualBasic

Imports AmosEngineLib

Imports AmosEngineLib.AmosEngine.TMatrixID

Imports AmosExtensions.CustomEstimand

Imports PBayes

Imports MiscAmosTypes

Imports MiscAmosTypes.cDatabaseFormat

#End Region

Public Class CEstimand

Implements IEstimand

Public Sub DeclareEstimands() Implements IEstimand.DeclareEstimands

 'Your code goes here.

 'Unstandardized effects (one direct effect plus three indirect effects) plus
four dichotomous estimands (see the AMOS manual regarding dichotomous
estimands)

 'Female group

 newestimand("a_unstd_female_1")

 newestimand("b_unstd_female_1")

```

newestimand("c_unstd_female_1")
newestimand("d_unstd_female_1")
newestimand("e_unstd_female_1")
newestimand("f_unstd_female_1")
newestimand("path1direct_unstd_female_1")
newestimand("path2indirect_Smoke_unstd_female_1")
newestimand("path3indirect_ExerciseFrequency_unstd_female_1")
newestimand("totalindirect_unstd_female_1")
newestimand("total_unstd_female_1")
newestimand("path1direct_lessthanorequaltozero_unstd_female_1")
newestimand("path2indirect_Smoke_lessthanorequaltozero_unstd_female_1")
newestimand("path3indirect_ExerciseFrequency_lessthanorequaltozero_unstd_female_1")
newestimand("totalindirect_lessthanorequaltozero_unstd_female_1")
newestimand("total_lessthanorequaltozero_unstd_female_1")

```

'Male group

```

newestimand("a_unstd_female_0")
newestimand("b_unstd_female_0")
newestimand("c_unstd_female_0")
newestimand("d_unstd_female_0")
newestimand("e_unstd_female_0")
newestimand("f_unstd_female_0")
newestimand("path1direct_unstd_female_0")
newestimand("path2indirect_Smoke_unstd_female_0")
newestimand("path3indirect_ExerciseFrequency_unstd_female_0")
newestimand("totalindirect_unstd_female_0")
newestimand("total_unstd_female_0")
newestimand("path1direct_lessthanorequaltozero_unstd_female_0")
newestimand("path2indirect_Smoke_lessthanorequaltozero_unstd_female_0")

```

```
newestimand(("path3indirect_ExerciseFrequency_lessthanorequaltozero_unstd_f
emale_0"))
```

```
newestimand(("totalindirect_lessthanorequaltozero_unstd_female_0"))
```

```
newestimand(("total_lessthanorequaltozero_unstd_female_0"))
```

```
'Unstandardized differences:Female = 1 MINUS Female = 0
```

```
newestimand(("path1_diff_unstd_female_1_0"))
```

```
newestimand(("path2_diff_unstd_female_1_0"))
```

```
newestimand(("path3_diff_unstd_female_1_0"))
```

```
newestimand(("path1_diff_lessthanorequaltozero_unstd_female_1_0"))
```

```
newestimand(("path2_diff_lessthanorequaltozero_unstd_female_1_0"))
```

```
newestimand(("path3_diff_lessthanorequaltozero_unstd_female_1_0"))
```

```
End Sub
```

```
Public Function CalculateEstimands(ByVal sem As AmosEngine) As String
Implements IEstimand.CalculateEstimands
```

```
    'Your code goes here.
```

```
    'UNStandardized estimates: Parameters and functions of parameters
```

```
    'Female = 1 (Group 1 model for Females)
```

```
    estimand("a_unstd_female_1").value = sem.GetEstimate(DirectEffects,
    "LogNetWorth", "HealthLiteracy", 1)
```

```
    estimand("b_unstd_female_1").value = sem.GetEstimate(DirectEffects,
    "Smoke", "HealthLiteracy", 1)
```

```
    estimand("c_unstd_female_1").value = sem.GetEstimate(DirectEffects,
    "ExerciseFrequency", "HealthLiteracy", 1)
```

```
    estimand("d_unstd_female_1").value = sem.GetEstimate(DirectEffects,
    "LogIncome", "Smoke", 1)
```

```
    estimand("e_unstd_female_1").value = sem.GetEstimate(DirectEffects,
    "LogIncome", "ExerciseFrequency", 1)
```

```
    estimand("f_unstd_female_1").value = sem.GetEstimate(DirectEffects,
    "LogNetWorth", "LogIncome", 1)
```

```
    estimand("path1direct_unstd_female_1").value=estimand("a_unstd_female_1").valu
e
```

```
    estimand("path2indirect_Smoke_unstd_female_1").value=
```

```
    estimand("b_unstd_female_1").value *
```

```
    estimand("d_unstd_female_1").value *
```

```
    estimand("f_unstd_female_1").value
```

```
    estimand("path3indirect_ExerciseFrequency_unstd_female_1").value=
```

```

    estimand("c_unstd_female_1").value *
    estimand("e_unstd_female_1").value *
    estimand("f_unstd_female_1").value

    estimand("totalindirect_unstd_female_1").value=sem.GetEstimate(IndirectEffects,
    "LogNetWorth", "HealthLiteracy", 1)

    estimand("total_unstd_female_1").value=sem.GetEstimate(TotalEffects,
    "LogNetWorth", "HealthLiteracy", 1)

    'Consistent with the null hypothesis (parameter <= 0) leading to Bayesian p
    value (Gill, 2008, p. 237)

    estimand("path1direct_lessthanorequaltozero_unstd_female_1").value=(estimand("p
    ath1direct_unstd_female_1").value<=0)

    estimand("path2indirect_Smoke_lessthanorequaltozero_unstd_female_1").value=(es
    timand("path2indirect_Smoke_unstd_female_1").value<=0)

    estimand("path3indirect_ExerciseFrequency_lessthanorequaltozero_unstd_female_1
    ").value=(estimand("path3indirect_ExerciseFrequency_unstd_female_1").value<=0)

    estimand("totalindirect_lessthanorequaltozero_unstd_female_1").value=(estimand("t
    otalindirect_unstd_female_1").value<=0)

    estimand("total_lessthanorequaltozero_unstd_female_1").value=(estimand("total_un
    std_female_1").value<=0)

    'Female = 0 (Group 2 Model for Males)

    estimand("a_unstd_female_0").value = sem.GetEstimate(DirectEffects,
    "LogNetWorth", "HealthLiteracy", 2)

    estimand("b_unstd_female_0").value = sem.GetEstimate(DirectEffects,
    "Smoke", "HealthLiteracy", 2)

    estimand("c_unstd_female_0").value = sem.GetEstimate(DirectEffects,
    "ExerciseFrequency", "HealthLiteracy", 2)

    estimand("d_unstd_female_0").value = sem.GetEstimate(DirectEffects,
    "LogIncome", "Smoke", 2)

    estimand("e_unstd_female_0").value = sem.GetEstimate(DirectEffects,
    "LogIncome", "ExerciseFrequency", 2)

    estimand("f_unstd_female_0").value = sem.GetEstimate(DirectEffects,
    "LogNetWorth", "LogIncome", 2)

    estimand("path1direct_unstd_female_0").value=estimand("a_unstd_female_0").valu
    e

    estimand("path2indirect_Smoke_unstd_female_0").value=

```

```

estimand("b_unstd_female_0").value *
estimand("d_unstd_female_0").value *
estimand("f_unstd_female_0").value

estimand("path3indirect_ExerciseFrequency_unstd_female_0").value=
estimand("c_unstd_female_0").value *
estimand("e_unstd_female_0").value *
estimand("f_unstd_female_0").value

estimand("totalindirect_unstd_female_0").value=sem.GetEstimate(IndirectEffects,
"LogNetWorth", "HealthLiteracy", 2)

estimand("total_unstd_female_0").value=sem.GetEstimate(TotalEffects,
"LogNetWorth", "HealthLiteracy", 2)

'Consistent with the null hypothesis (parameter <= 0) leading to Bayesian p
value (Gill, 2008, p. 237)

estimand("path1direct_lessthanorequaltozero_unstd_female_0").value=(estimand("p
ath1direct_unstd_female_0").value<=0)

estimand("path2indirect_Smoke_lessthanorequaltozero_unstd_female_0").value=(es
timand("path2indirect_Smoke_unstd_female_0").value<=0)

estimand("path3indirect_ExerciseFrequency_lessthanorequaltozero_unstd_female_0
").value=(estimand("path3indirect_ExerciseFrequency_unstd_female_0").value<=0)

estimand("totalindirect_lessthanorequaltozero_unstd_female_0").value=(estimand("t
otalindirect_unstd_female_0").value<=0)

estimand("total_lessthanorequaltozero_unstd_female_0").value=(estimand("total_un
std_female_0").value<=0)

'Differences: unstandardized

'Female = 1 MINUS Female = 0

estimand("path1_diff_unstd_female_1_0").value=estimand("path1direct_unstd_fem
ale_1").value - estimand("path1direct_unstd_female_0").value

estimand("path2_diff_unstd_female_1_0").value=estimand("path2indirect_Smoke_u
nstd_female_1").value - estimand("path2indirect_Smoke_unstd_female_0").value

estimand("path3_diff_unstd_female_1_0").value=estimand("path3indirect_Exercise
Frequency_unstd_female_1").value -
estimand("path3indirect_ExerciseFrequency_unstd_female_0").value

Differences: Unstandardized less than or equal to zero

estimand("path1_diff_lessthanorequaltozero_unstd_female_1_0").value=(estimand("
path1_diff_unstd_female_1_0").value<=0)

estimand("path2_diff_lessthanorequaltozero_unstd_female_1_0").value=(estimand("
path2_diff_unstd_female_1_0").value<=0)

```

```
estimand("path3_diff_lessthanorequaltozero_unstd_female_1_0").value=(estimand("path3_diff_unstd_female_1_0").value<=0)
```

```
Return ""Return an empty string if no error occurred
```

```
End Function
```

```
End Class
```