

Psychology 454: Latent Variable Modeling

How do you know if a model works?

William Revelle

Department of Psychology
Northwestern University
Evanston, Illinois USA



NORTHWESTERN
UNIVERSITY

November, 2012

Outline

- 1 Goodness of fit measures
 - Absolute fit indices
 - Incremental or relative fit indices
 - Distribution free fit functions – after Loehlin and Browne
- 2 Measures of fit
- 3 Fits and sample size
- 4 Advice

A number of tests of fit taken from Marsh et al. (2005)

- ① Marsh, Hau & Grayson (2005) lists 40 different proposed measures of goodness of fit
- ② Measures of absolute fit
 - I_o = index of fit for original or baseline model
 - I_t = index of fit for target or “true” model
- ③ Measures of incremental fit Type I
 - $\frac{|I_t - I_o|}{\text{Max}(I_o, I_t)}$ which is either
 - $\frac{I_o - I_t}{I_o}$
 - or $\frac{I_t - I_o}{I_t}$
- ④ Measures of incremental fit Type II
 - $\frac{|I_t - I_o|}{E(I_t - I_o)}$ which is either
 - $\frac{I_o - I_t}{I_o - E(I_t)}$
 - or $\frac{I_t - I_o}{E(I_t) - I_o}$

Fit functions from Jöreskog

- ① Ordinary least squares $F = \frac{1}{2}tr(S - \Sigma)^2$
 - The squared difference between the observed (S) and model (Σ) covariance matrices
 - tr means trace of the sum of the diagonal values of the matrix of squared deviations
- ② Generalized least squares $F = \frac{1}{2}tr(I - S^{-1}\Sigma)^2$
 - I is the identity matrix
 - if the model = data, then $S^{-1}\Sigma$ should be I
 - weight the fit by the inverse of the observed covariances
- ③ Maximum Likelihood $F = \log|\Sigma| + tr(S\Sigma^{-1}) - \log|S| - p$
 - weight the fit by the inverse of the modeled covariance
 - p is the number of variables
 - $tr(I) = p$, and thus the ML should be 0 if the model fits the data

Fit-function based indices

① Fit Function Minimum fit function (FF)

- $FF = \frac{\chi^2}{(N-1)}$

② Likelihood ratio $LHR = e^{-1/2FF}$

③ χ^2 (minimum fit function chi square)

- $\chi^2 = tr(\Sigma^{-1}S - I) - \log|\Sigma^{-1}S| = (N-1)FF$

④ $p(\chi^2)$ probability of observing a χ^2 this larger or larger given that the model fits

⑤ $\frac{\chi^2}{df}$ has expected value of 1

Non-centrality based indices

- ① Dk: Rescaled noncentrality parameter (McDonald & Marsh, 1990)
 - $Dk = FF - df / (N - 1) = \frac{\chi^2 - df}{N - 1}$
- ② PDF (population discrepancy function = DK normed to be non-negative)
 - $PDF = \max(\frac{\chi^2 - df}{N - 1}, 0)$
- ③ Mc: Measure of centrality (CENTRA, MacDonald Fit Index (MFI))
 - $Mc = e^{\frac{-(\chi^2 - df)}{2(N - 1)}}$
- ④ Non-centrality parameter
 - $NCP = \chi^2 - df$

Error of approximation indices

How large are the residuals, estimated several different ways

① RMSEA (root mean square error of approximation)

- $RMSEA = \sqrt{PDF/df} = \sqrt{\frac{\max(\frac{\chi^2 - df}{N-1}, 0)}{df}}$
- based upon the non-central χ^2 distribution to find the error of fit

② MSEA (mean square error of approximation – unnormed version of RMSEA)

- $MSEA = \frac{Dk}{df} = \frac{\chi^2 - df}{(N-1)df}$

③ RMSEAP (root mean square error of approximation of close fit)

- $RMSEA < .05$

④ RMR Root mean square residual (or, if S and Σ are standardized, the SRMR). Just

- square root of the average squared residual
- $\sqrt{\frac{2 \sum (S - \Sigma)^2}{p*(p+1)}}$

Information indices

Compare the information of a model with the number of parameters used for the model. These allow for comparisons of different models with different degrees of freedom.

- 1 AIC (Akaike Information Criterion for a model penalizes for using up df)

- $AIC = \chi^2 + p * (p + 1) - 2df = \chi^2 + 2K$
- where $K = \frac{p*(p+1)}{2} - df$

- 2 Bayesian Information Criterion

- $-2\text{Log}(L) + p\log(N) = \chi^2 - K\log(N(.5(p(p + 1)))$

Goodness of fit indices

- ① GFI from LISREL
 - $GFI = 1 - \frac{tr(\Sigma^{-1}S - I)^2}{tr(\Sigma^{-1}S)^2}$
- ② Adjusted Goodness of Fit (Lisrel)
 - $AGFI = 1 - \frac{p(p+1)}{2df}(1 - GFI)$
- ③ Unbiased GFI (from Steiger)
 - $GFI = \frac{p}{2 \frac{(\chi^2 - df)}{(N-1)} + p}$

Comparing solutions to solutions

- ① Incremental fit indices without correction for model complexity
 - RNI (relative non-centrality) McDonald and Marsh
 - CFI Comparative fit index (normed version of RNI) Bentler
 - Normed Fit index (Bentler and Bonett)
- ② Incremental fit indices correcting for model complexity
 - Tucker - Lewis Index
 - Normed Tucker Lewis
 - Incremental Fit index
 - Relative Fit Index
- ③ Parsimony indices

Incremental fit indices without correction for model complexity

- ① RNI (relative non-centrality) McDonald and Marsh
 - $RNI = 1 - \frac{DK_t}{DK_n}$
 - where $DK = \frac{\chi^2 - df}{N - 1}$ for either the null or the tested model
- ② CFI Comparative fit index (normed version of RNI) Bentler
 - Just norm the RNI to be greater than 0.
 - $CFI = 1 - \frac{MAX(NCP_t, 0)}{MAX(NCP_n, 0)}$
- ③ Normed Fit index (Bentler and Bonett)

Fitting functions from Loehlin

- ① Let S be the “strung out” data matrix
- ② Let Σ be the “strung out” model matrix
- ③ $Fit = (S - \Sigma)'W^{-1}(S - \Sigma)$
- ④ Where $W =$
 - Ordinary Least Squares $W = I$
 - Generalized Least Squares $W = SS'$
 - Maximum likelihood $W = \Sigma\Sigma'$

Practical advice

- ① Taken from Kenny
 - <http://davidkenny.net/cf/fit.htm>
- ② Bentler and Bonnet Normed Fit Index
 - $\frac{\chi^2_{Null} - \chi^2_{Model}}{\chi^2_{Null}}$
 - Between .90 and .95 is acceptable
 - > .95 is “good”
- ③ RMSEA
 - if $\chi^2 < df$, then $RMSEA = 0$
 - “good” models have $RMSEA < .05$
 - “poor” models have $RMSEA > .10$
- ④ p of close fit
 - Null hypothesis is that $RMSEA$ is .05
 - test if $RMSEA > .05$
 - Claim good fit if $p(RMSEA > .05) > .05$

Fits and sample size

- 1 See associated simulation results
 -
- 2
 -

Considering rules of thumb and fit

- ① Fit functions have distributions and thus are susceptible to problems of type I and type II error.
 - Compare the fits for correct model as well as those for a simple incorrect
- ② Should we just use chi square and reject models that don't fit, or should we reason about why they don't fit

What does it mean if the model does not fit

- 1 Model is wrong
- 2 Measurement is wrong
- 3 Structure is wrong
- 4 Assumptions are wrong
- 5 at least one of above, but which one?

Specification & Respecification

- ① Is the measurement model consistent
 - revise it
 - evaluate loadings
 - evaluate error variances
 - more or fewer factors
 - correlated errors?
- ② Structural model:
 - adjust paths
 - drop paths
 - add paths
- ③ Equivalent models
 - What models are equivalent
 - Do they make equally good sense

Conclusion

- ❶ Latent variable models are a powerful theoretical aid but do not replace theory
- ❷ Nor do latent modeling algorithms replace the need for good scale development
- ❸ Latent variable models are a supplement to the conventional regression models of observed scores.
- ❹ Other latent models (not considered) include
 - item Response Theory
 - Latent Class Analysis
 - Latent Growth Curve analysis

Marsh, H. W., Hau, K.-T., & Grayson, D. (2005). Goodness of Fit in Structural Equation Models. In A. Maydeu-Olivares & J. J. McArdle (Eds.), *Contemporary Psychometrics* chapter 10, (pp. 275–340). New York: Routledge.

McDonald, R. & Marsh, H. (1990). Choosing a multivariate model: Noncentrality and goodness of fit. *Psychological Bulletin*, 107(2), 247–255.