

Small Sample Bayesian Factor Analysis

Author: Dirk Heerwegh, Business & Decision Life Sciences, Brussels, Belgium

ABSTRACT

Factor analysis (FA) is a statistical technique to explain the correlation structure among observed variables. It can e.g. be useful in the analysis of Patient Reported Outcomes, which typically use multiple questions to measure a single concept. Maximum likelihood factor analysis (ML-FA) relies on large sample theory, and it is consequently often recommended to use it only in large samples (e.g. $N=200$ or more). In smaller samples, ML-FA can run into problems (e.g. model non-convergence, negative residual variances). Bayesian statistics typically perform better in small samples, and may therefore be useful in (clinical) studies that rely on smaller sample sizes.

This paper provides a gentle introduction to Bayesian FA and shows how it can be implemented in Mplus 7.0. The paper then shows the relative performance of Bayesian FA as compared to ML-FA for two CFA models: a one-factor model with four factor indicators, and a two-factor model with 3 indicators for each factor.

Several criteria are taken into account to evaluate the FA results in samples of sizes 200, 100, 50, and 25: problems with model fitting (model convergence, negative residual variances), and statistical power. The influence of priors in Bayesian FA on the parameter estimates is also investigated in terms of bias on factor loadings and factor correlations.

The paper explains the results from the Monte Carlo studies in practical terms. It concludes by pointing to some novel capabilities in Bayesian FA such as estimation of near-zero cross-loadings or quasi zero residual covariances, which represent a more flexible approach than can typically be done in ML-FA.

1. INTRODUCTION

2. FACTOR ANALYSIS

Exploratory factor analysis (EFA) is a statistical technique to identify the number and nature of the underlying factors that are responsible for the correlation structure in the data (Hatcher, 1994, p. 59). For example, we expect that the 12 items (Q1-Q6c) in the Veterans VR-12 survey (Spira et al., 2004), are correlated with each other because the responses to these items are assumed to be influenced by two underlying factors (broad concepts referring to Physical and Mental functioning, see Figure 1). As a consequence, we'd expect an EFA to show a solution with two factors, and we would expect to be able to characterize the nature of these factors (i.e. label them appropriately) by investigating the content of the items which load on each of the two factors.

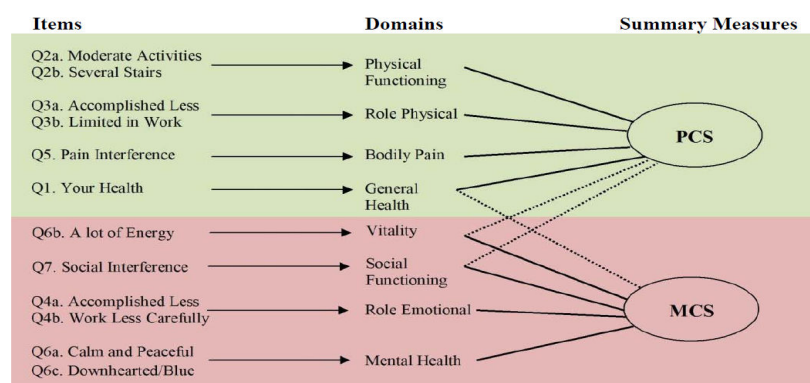


Figure 1. Theoretical model mapping the 12 items from the Veterans Rand survey to 2 summary measures (PCS-Physical Component Summary, and MCS-Mental Component Summary). Source: Centers for Medicare (2012), p. 6.

EFA is appropriate when little is known about the number and nature of the factors that may be responsible for the correlation structure in the data. If a validated measurement instrument is used, the number and nature of the underlying factors is known prior to the data collection and analysis. In such circumstances, a more appropriate factor analysis method is Confirmatory Factor Analysis (CFA). CFA differs from EFA in that it imposes an a priori model on the data, and tests the degree to which it is plausible that the data were generated by the proposed model. Again taking the Veterans VR-12 survey as an example, one can not only ask CFA to extract 2 factors from the data, but also to do so in accordance with a theoretical model which specifies which items can load on which factors. Figure 1 shows the theoretical model imposed on the VR-12 survey. Question 2a ("moderate activities") for instance, should load on the physical component summary (PCS), but not on the mental component summary (MCS). Technically, this means that the factor loading of Q2a on MCS should be restricted to zero. Such restrictions give rise to an imposed factor structure (graphically represented by the absence and presence of arrows in the model). Obviously, each restriction can bring with it a certain degree of model misfit, and the purpose of the CFA analysis is to assess whether the degree of misfit is within the range of what is statistically acceptable. For a worked example on the VR-12 items, see Heerwegh (2013).

CFA can be conducted with Sas® PROC CALIS, and other software packages are available that were specifically designed to perform CFA (and related analyses that are more generally termed Structural Equation Models, abbreviated as SEM), e.g. Mplus (www.statmodel.com), LISREL (www.ssicentral.com/lisrel/index.html), AMOS (<http://www-03.ibm.com/software/products/us/en/spss-amos/>), EQS (<http://www.mvsoft.com/>), and several packages within R (<http://www.r-project.org/>), such as lavaan (<http://cran.r-project.org/web/packages/lavaan/index.html>), and sem (<http://cran.r-project.org/web/packages/sem/index.html>). In this paper, we will use Mplus. A demo version of Mplus can be downloaded from <http://www.statmodel.com>.

It is routinely recommended to utilize Maximum Likelihood factor analysis with large samples. De Winter et al. (2009) provide an overview of recommended sample sizes for ML-EFA from the early literature (1950s to late 1970s). Many different sample sizes were recommended, but N=200 or 250 and N=500 seem to be the most often recommended sample sizes (Guilford, 1954; Gorsuch, 1973; Comrey, 1973; Cattell, 1978 as cited in de Winter et al. 2009). Other researchers focused on the ratio of the number of cases per variable (N/p), with recommendations ranging from 3:1-6:1 to 20:1 (Cattell 1978; Hair et al 1979 cited in de Winter et al. 2009). Later studies acknowledged that other elements are of importance, including the communalities, the ratio of the number of variables per factor (p/f), factor loadings, and number of factors. De Winter et al. (2009) showed through simulation studies that samples of well below N=50 are sufficient when the factor loadings are high, the number of factors in the model is low, and the number of observed variables is high. For example, 1-factor models with factor loadings of 0.60 and 6 observed variables, required a sample of only N=18 for satisfactory factor recovery.

A number of studies have focused on sample size for ML-CFA. Recommended sample sizes ranged from 200 (Boomsma, 1982 cited in de Winter et al. 2009) to 100 (Marsh and Hau, 1999 cited in de Winter et al. 2009). Given the results from the ML-EFA studies, it is possible that ML-CFA also performs well at lower sample sizes when the data is well conditioned (i.e. high factor loadings, low number of factors, high number of observed variables).

3. BAYESIAN STATISTICS

At least two properties of Bayesian statistics are appealing. The first is that Bayesian approaches allow to incorporate background (prior) knowledge into the analyses instead of testing essentially the same null hypothesis over and over, ignoring the lessons from previous studies (van de Schoot et.al., 2014). The second is that Bayes does not rely on large-sample theory, and as a consequence can perform better in small samples (Muthén & Asparouhov, 2012).

There are three main ingredients to a Bayesian approach (van de Schoot et.al., 2014):

1. The background knowledge on the parameters of the model being tested. This refers to all knowledge available *before* seeing the data, and is captured in the prior distribution. The variance of this distribution reflects the degree of uncertainty about the population value of the parameter of interest (the larger the variance, the more uncertain we are).
2. The information in the data themselves.
3. A combination of the first two ingredients into the posterior distribution. It is a compromise between the prior and the data, and reflects one's updated knowledge about the model parameters of interest.

Prior distributions can be "non-informative" or "informative". A uniform prior distribution is said to be non-informative because it essentially states that all values are equally likely. Normal distributions (with variance < infinity) for instance, are informative, because they are characterized by a mean and a variance. By stipulating a Normal distributed prior, we are effectively expressing a belief regarding the likely value of the mean. The variance of the prior reflects our uncertainty (with a larger variance reflecting more uncertainty). Other prior distributions are also possible, but software packages may impose limitations on the available prior distributions (van de Schoot et.al., 2014). Mplus offers a wide variety of priors: the Normal, the Inverse Gamma, the Uniform, the Gamma, the Inverse Wishart and the Dirichlet distribution (Asparouhov & Muthén, 2010a, pp. 34-36).

As sample sizes get smaller, and/or the prior becomes more informative, the posterior becomes more influenced by the prior (van de Schoot et.al., 2014).

4. BAYESIAN FACTOR ANALYSIS

Mplus offers Bayesian analysis capability as of version 6 (<http://www.statmodel.com/verhistory.shtml>). This capability is offered for a variety of statistical models, including but not limited to CFA. The hallmarks of Bayesian analysis methods carry over to Bayesian CFA in that it is possible to incorporate priors in the model, and that the use of small samples should – at least theoretically – be less problematic than MLE.

Below, we show the syntax in Mplus to fit a two-factor CFA model on the VR-12 survey. In the left panel, the syntax is given for ML-CFA and in the right panel the syntax is given for Bayesian CFA. Note that the example does not use all variables available in the dataset (consult Figure 1 to verify that in fact there are 12 variables available to measure MCS and PCS). The reason is that the demo version on Mplus is limited to 6 dependent variables, so we restricted the example syntax to 6 variables so that readers who only have access to the demo version of Mplus can also run the programs.

Table 1. Mplus version 7 syntax to perform ML-CFA (left panel) and Bayesian CFA (right panel)

<pre>TITLE: ML 2 factor CFA for VR-12 data DATA: FILE IS "sample1.csv"; VARIABLE: NAMES ARE age race gender educ q1 q2a q2b q3a q3b q4a q4b q5 q6a q6b q6c q7 lang region; USEVAR ARE q4a q4b q6a q3a q3b q5; MISSING ARE ALL(99999); ANALYSIS: TYPE=GENERAL; ESTIMATOR=ML; MODEL: mcs BY q4a* q4b q6a; pcs by q3a* q3b q5; mcs@1; pcs@1; pcs WITH mcs; OUTPUT: STAND MOD;</pre>	<pre>TITLE: BAYES 2 factor CFA for VR-12 data DATA: FILE IS "sample1.csv"; VARIABLE: NAMES ARE age race gender educ q1 q2a q2b q3a q3b q4a q4b q5 q6a q6b q6c q7 lang region; USEVAR ARE q4a q4b q6a q3a q3b q5; MISSING ARE ALL(99999); ANALYSIS: TYPE = GENERAL ESTIMATOR = BAYES; PROC = 2; FBITER = 10000; MODEL: mcs BY q4a* q4b q6a (I1-I3); pcs by q3a* q3b q5 (I4-I6); mcs@1; pcs@1; pcs WITH mcs (psi1); MODEL PRIORS: I1-I6 ~ N(0.6,0.05); psi1 ~ N(0.5,0.05); OUTPUT: STAND;</pre>
--	---

An Mplus syntax file consists of several parts. The TITLE: statement is optional and serves to give the program (and the corresponding output) a title. The DATA: statement is required and serves to define the data to be used. In the syntax above, the data is read in from a comma separated values file, "sample1.csv". The VARIABLE: statement is also required and allows one to supply information about the variables in the source file (NAMES ARE), which variables will be used in the current analysis (USEVAR ARE), and which numeric values represent missing values (MISSING ARE). In this file, missing values are represented by value 99999.

The ANALYSIS: statement allows one to specify the analysis settings. In the ML-CFA syntax, TYPE is set to GENERAL (the default analysis type in Mplus), and ESTIMATOR is set to ML (Maximum Likelihood). (In fact, TYPE=GENERAL and ESTIMATOR=ML are the default Mplus settings, so in principle these should not have been explicitly declared in the program.) In the BAYES-CFA syntax, ESTIMATOR is set to BAYES. The PROC = option allows the user to define how many processors will be used to run the program. Increasing the number of processors increases the speed of the model fitting. The FBITER= option sets a fixed number of iterations for the MCMC chains.

In the MODEL: statement, the actual CFA model is specified. It can be seen that two factors are defined, "MCS" and "PCS". The MCS factor is measured by (BY) q4a, q4b, and q6a. The asterisk (q4a*) means that the factor loading for q4a will be estimated freely. This requires fixing the factor variance to 1 (done through mcs@1). The default parameterization in Mplus is to fix the factor loading of the first variable of each factor to 1, and to freely estimate the factor variance. Another parameterization is to freely estimate all factor loadings and fix the factor variance to 1, as is done in the current example. PCS is measured by q3a q3b and q5. Again, we use the alternative parameterization and freely estimate all factor loadings and fix the factor variance to 1. Also note that we explicitly allow PCS and MCS to correlate with each other (pcs WITH mcs). (This is also the default behaviour of Mplus, so the statement could have been omitted). Note that the MODEL statements are identical in the BAYES-CFA syntax, except for the addition of labels which are added at the end of each line between parentheses. The factor loading of q4a on MCS is labelled "I1" (short for lambda 1), the factor loading of q4b on MCS is labelled "I2" (short for lambda 2), ... and the factor loading of q5 on PCS is labelled "I6" (lambda 6). The covariance between MCS and PCS is labelled "psi1". These labels could also have been added in the ML-CFA, but they would not have served any meaningful purpose in the current example. The reason they were included in the BAYESIAN-CFA is that they are needed to define the priors. This is done in the MODEL PRIORS: statement. From the statements, it can be seen that all lambdas are given the same prior: a Normal distribution with mean 0.6 and variance 0.05. The prior for the factor covariance is also a Normal distribution, but with mean 0.5 and factor variance 0.05. Note that in practice, the priors would be chosen more carefully, based on previous research, and that the priors on the factor loadings would in reality probably not all be the same.

The OUTPUT: statement finally, allows one to request additional output (on top of the default Mplus output). In the ML-CFA syntax, standardized parameter estimates are requested through the specification of STAND, and modification indices are also requested through MOD (These can be helpful to modify the model in order to improve model fit). In Bayesian CFA, the MOD option is not supported, so it does not appear in the syntax.

STUDY AIMS

Some simulation studies have already been undertaken to evaluate the performance of Bayesian CFA (e.g. Asparhouhov & Muthén, 2010b). The current paper seeks to build on these studies by evaluating Bayesian CFA in comparison to ML-CFA for a set of scenarios which are more or less typical for the pharmaceutical setting (i.e. small samples). More specifically, this paper seeks to answer the following general and specific research questions:

- How does Bayesian CFA compare to ML-CFA at moderate to very small sample sizes?
- What is the influence of priors in Bayesian CFA?
 - Non-informative versus informative priors
 - Correctly versus incorrectly specified priors

These research questions are tackled using simulation (Monte Carlo) studies conducted in Mplus version 7.

In general, it can be expected that Bayesian CFA will perform better in small to very small samples as compared to ML-CFA. However, with smaller samples, the influence of the prior should increase. Therefore, it becomes important to assess the effect of miss specified priors on the model results.

This paper will assess the quality of the analysis results on 4 criteria:

- Percentage parameter bias:
 - The degree (in percentage) to which parameter values are under- or over-estimated can be calculated as:
 - $100 \times (\text{mean estimated value} - \text{population value}) / \text{population value}$
- Mean Squared Error (MSE)
 - MSE is a measure that takes into account both bias and variance of the estimates across the replications in the Monte Carlo simulation. It is calculated as:
 - Variance of the estimates across the replications plus the square of the bias.
 - Where bias = mean estimated value – population value
- Statistical power
 - In general it will be desirable to have sufficient statistical power to reject false null hypotheses.
- Non-convergence of the models or errors in the model estimation
 - It can be expected that ML-CFA will have problems estimating the models when sample sizes are (very) small. These problems can be manifested in a number of ways: outright model non-convergence, a residual covariance matrix that is not positive definite, and/or untrustworthy standard errors as identified by Mplus. We will take all these potential problems into account and use a measure of “error free runs” to refer to runs that are not affected by any of these problems.

5. MONTE CARLO STUDIES

5.1. STUDY 1: 1-FACTOR MODEL

In study 1, a 1-factor model with 4 observed variables will be used. The model is graphically depicted in Figure 2.

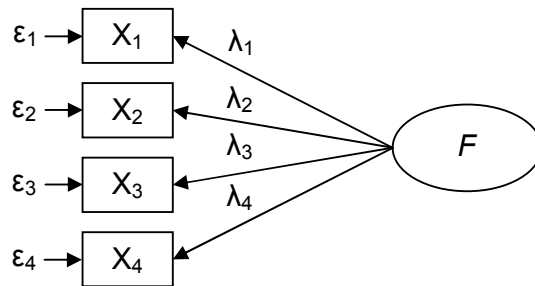


Figure 2. One factor model with 4 observed variables.

The simulation study includes 48 conditions:

- 4 levels for sample size ($N=200$, $N=100$, $N=50$, and $N=25$)
- 3 levels for the standardized factor loadings (.80, .60, and .40)
 - These represent, respectively, high, moderate and low factor loadings. Note that these are standardized factor loadings (i.e. range $[-1, +1]$). Also note that all 4 factor loadings in the model will be set at the same value (i.e. we will not test models where for instance 2 factor loadings are set at 0.80, one factor loading at 0.60 and one at 0.40).
- 4 levels for the type of estimation:
 - Maximum likelihood
 - Bayes, with non-informative prior: Normal distribution with mean 0 and variance *infinity*.
 - Bayes, with informative prior: Normal distribution with mean 0.40 and variance 0.05.
 - Bayes, with informative prior: Normal distribution with mean 0.40 and variance 0.01.

As can be seen in Table 2, using a prior $N(0.40, 0.05)$ leads to 90% of the factor loadings to fall within the range of $[0.03, 0.77]$. With a prior $N(0.40, 0.01)$, 90% fall within $[0.24, 0.56]$. So the prior with variance 0.05 is a relatively weak prior, stipulating in essence that the factor loading is expected to be positive, while the prior with variance 0.01 is more informative in that it not only expects the factor loading to be a positive value, but also that it is of a relatively weak to moderate size.

Table 2. 90% and 95% limits of factor loadings by selected variance

Variance	90% Limits	95% Limits
0.05	[0.03;0.77]	[-0.04;0.84]
0.01	[0.24;0.56]	[0.20;0.60]

The Mplus default priors for CFA model parameters are shown in Table 3 (Asparhouhov & Muthén, 2010b, p.58). In simulation study 1, we will only vary the prior distribution for the factor loadings as explained above.

Table 3. Default priors in Mplus

Model parameter	Default prior	Comments
Loadings	$N(0, \infty)$	
Intercepts	$N(0, \infty)$	
Residual variances	$IG(0, -1)$	
Factor covariances	$IW(0, -p-1)$	Where p is size of the matrix

5.2. STUDY 1: RESULTS

The results from the study are summarized in Table 4 (see section 8).

Parameter bias

ML-CFA leads to factor loading estimates which are close to the population values when the population factor loadings are moderate (0.60) or high (0.80). With weak factor loadings (0.40), the estimates are biased upwards, especially at (very) small sample sizes (N=50 or less).

In contrast to ML-CFA, Bayesian CFA performs well at weak factor loadings (0.40). Especially the informative priors perform well, and this at all sample sizes. Not surprisingly, Bayesian CFA performs less well when the informative priors are miss specified (true factor loadings 0.60 or 0.80 but the prior mean set at 0.40). In this case, the factor loadings are under-estimated (and this to a larger degree when the prior was more informative, i.e. with a smaller variance). Interestingly, a non-informative prior seems to lead to over-estimated factor loadings when the true factor loadings are 0.60 or 0.80. This upward bias is reduced with an increasing sample size.

So from these results it appears ML-CFA performs well when the factor loadings are high, but the Bayesian CFA performs better when the factor loadings are rather weak.

Mean Squared Error (MSE)

When using ML-CFA, the MSE increases with decreasing factor loadings. The MSE also decreases with increasing sample sizes, so that at N=200, the MSE approaches zero with all factor loadings.

Bayesian CFA performs better than ML-CFA when the factor loadings are weak (0.40), even with a non-informative prior. With informative, correctly specified priors, Bayesian CFA performs markedly better than ML-CFA. Interestingly, even a miss specified prior performs slightly better than ML-CFA at moderate factor loadings of 0.60. In this case, it seems better to use a miss specified weak informative prior than a non-informative prior. This is probably due to the way the MSE is calculated, since it incorporates the variance of the posterior – with an informative prior the variance is smaller than with a non-informative prior. At sample size of 200 and a non-informative prior, the average estimated factor loading in the simulations was 0.6081 for the first observed variable, with a standard deviation of 0.0830. The MSE therefore equals $0.0830^2 + (0.6081 - 0.6000)^2 = 0.0070$. With an informative prior (Normal distribution with mean = 0.40 and variance = 0.05), the average estimated factor loading in the simulations was 0.5796 for the first observed variable, with a standard deviation of 0.0719. The MSE therefore equals $0.0719^2 + (0.5796 - 0.6000)^2 = 0.0056$. This numerical example shows that the point estimate is not far off the population value, while the variance is reduced. The consequence is a smaller MSE when using a weak informative prior as compared to a non-informative prior.

Statistical power

ML-CFA returns near 100% power at high factor loadings (0.80). With sample sizes of N=50 or larger, statistical power is also (nearly) 100% with moderate factor loadings (0.60). Power drops to 75% at N=25. With weak factor loadings (0.40), statistical power ranges from 25% to 90% as the sample size moves from N=25 to N=200.

Bayesian CFA also returns close to 100% power at high factor loadings (0.80), even when using miss-specified priors for the factor loadings (with a mean set at 0.40). Informative priors (although miss-specified at 0.40), return close to 100% power at all sample sizes. A non-informative prior returns a lower power than the ML-CFA at sample size N=25. With weak factor loadings (0.40), the informative prior with a low variance (0.01) consistently gives 100% power. As the prior become less informative (variance = 0.05 and infinity), the power drops with the sample size. The lowest power is achieved at N=25 and a non-informative prior (power = 12%).

Error-free runs

The results are straightforward. While ML-CFA runs into problems at low sample sizes and weak to moderate factor loadings, Bayesian CFA consistently runs without errors.

Overall conclusions simulation study 1

The results from study 1 clearly show that Bayesian CFA has advantages over ML-CFA when the sample size is (very) small and when the factor loadings are weak. ML-CFA is at least as good as Bayesian CFA when the sample size is larger or when the factor loadings are high.

When using small samples and when the factor loadings are expected to be low (around 0.40), Bayesian CFA can bring advantages. In such scenarios, it pays to use a weak informative prior if one is not very confident about the factor loadings to expect. If one is fairly confident, then a stronger prior should be used. The potential drawback of using a stronger prior is that the bias and MSE on the factor loadings can increase. Using a non-informative prior in small sample size and low expected factor loadings situations does not seem the best strategy because the MSE is expected to be larger and the power much lower. A non-informative prior does not seem a good choice in many cases since at least the sign of the factor loading should be quite clear (e.g. a negatively worded item should load negatively on a positively scaled factor, so one could at least use a prior that would restrict the standardized factor loading to be concentrated in the negative region from 0 to -1).

5.2. STUDY 2: 2-FACTOR MODEL

In study 2, a 2-factor model with 3 observed variables each will be used. The model is graphically depicted in Figure 3.

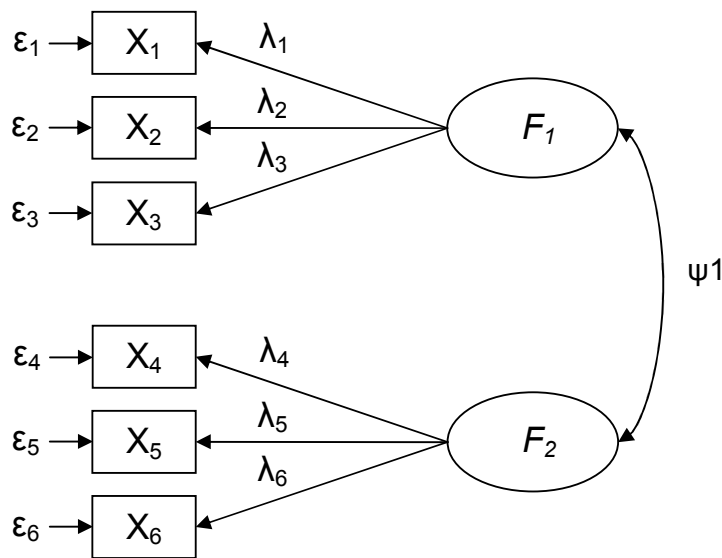


Figure 3. Two factor model with 6 observed variables (3 per factor) and a factor covariance.

The simulation study includes 96 conditions:

- 4 levels for sample size (N=200, N=100, N=50, and N=25)
- 3 levels for the standardized factor loadings (.80, .60, and .40)
 - These represent, respectively, high, moderate and low factor loadings. Note that these are standardized factor loadings (i.e. range [-1, +1]). Also note that all 6 factor loadings in the model will be set at the same value (i.e. we will not test models where for instance 2 factor loadings are set at 0.8, 2 factor loadings at 0.6 and two at 0.4).
- 4 levels for the type of estimation:
 - Maximum likelihood
 - Bayes, with non-informative prior: Normal distribution with mean 0 and variance *infinity*.
 - Bayes, with informative prior: Normal distribution with mean 0.4 and variance 0.05.
 - Bayes, with informative prior: Normal distribution with mean 0.4 and variance 0.01.
- 2 levels for the factor correlation (0.25 and 0.40)

- Note that when Bayesian CFA is used, a non-informative prior will be used with regard to the factor correlation.

5.2. STUDY 2: RESULTS

The results from study 2 are summarized in Table 5 (see section 8).

Parameter bias on the factor loadings (λ bdas)

ML-CFA leads to factor loading estimates which are close to the population values when the population factor loadings are moderate (0.60) or high (0.80). With weak factor loadings (0.40), the estimates are biased upwards, especially at (very) small sample sizes ($N=50$ or less). This finding replicates the findings based on the 1-factor model.

In contrast to ML-CFA, Bayesian CFA performs well at weak factor loadings (0.40). Especially the informative priors perform well, and this at all sample sizes. Not surprisingly, Bayesian CFA performs less well when the informative priors are miss specified (true factor loadings 0.60 or 0.80 but the prior mean set at 0.40). In this case, the factor loadings are under-estimated (and this to a larger degree when the prior was specified with a smaller variance). Interestingly, a non-informative prior seems to lead to over-estimated factor loadings when the true factor loadings are 0.60 or 0.80. This upward bias is reduced with an increasing sample size. Again, these findings replicate those from the 1-factor model.

So from these results it appears ML-CFA performs well when the factor loadings are high, but the Bayesian CFA performs better when the factor loadings are rather weak. This is in line with the simulations on the 1-factor model.

Mean Squared Error (MSE) on the factor loadings (λ bdas)

When using ML-CFA, the MSE increases with decreasing factor loadings. The MSE also decreases with increasing sample sizes. These findings correspond to what was observed in the 1-factor model. However, the MSE appears to be larger in the 2-factor model than in the 1-factor model. While in the 1-factor model, the MSE approached zero at sample size $N=200$ for factor loadings of 0.40, this is not the case in the 2-factor model (with an MSE still close to 0.10). One interpretation is that ML-CFA is having more difficulties fitting more complicated models.

Bayesian CFA performs better than ML-CFA when the factor loadings are weak (0.40), even with a non-informative prior. With informative, correctly specified priors, Bayesian CFA performs markedly better than ML-CFA. This echoes the findings from the 1-factor model. At moderate factor loadings (0.60), Bayesian CFA performs better than ML-CFA, except when using a miss specified prior with low variance (mean = 0.40 and variance = 0.10) when the sample is of size $N=100$ or $N=200$. At high factor loadings, ML-CFA and non-informative Bayes have similar MSE levels. The miss specified Bayes CFA analyses fare worse than ML-CFA and Bayes CFA with a non-informative prior, although a larger sample size provides some protection when using a weakly informative prior (mean = 0.40 and variance = 0.05).

Parameter bias on the factor correlation (ψ)

ML-CFA leads to factor correlation estimates which are close to the population values when the population factor loadings are moderate (0.60) or high (0.80). With weak factor loadings (0.40), the estimates are biased upwards, especially at (very) small sample sizes ($N=50$ or less).

Bayesian CFA performs better than ML-CFA at low factor loadings (0.40), especially at smaller sample sizes. A non-informative prior leads to an under estimation of the factor correlation, while an informative prior leads to an over estimate. At moderate and high factor loadings, and miss specified priors on the factor loadings, the factor correlations are over estimated by as much as up to 200%. The implied covariance between two observed variables (e.g. x_1 and x_3) that are indicators of two different factors (say, factor F_1 and factor F_2), is given by equation¹:

¹ The implied covariance between x_1 and x_3 is derived as follows. From the CFA model (cf. the diagram in Figure 2), we know that:

$$x_1 = \lambda_1 \cdot F_1 + \varepsilon_1$$

$$x_3 = \lambda_3 \cdot F_2 + \varepsilon_3$$

Using the formula

$$\text{cov}(aX + bY, cZ + dU) = ac \cdot \text{cov}(X, Z) + ad \cdot \text{cov}(X, U) + bc \cdot \text{cov}(Y, Z) + bd \cdot \text{cov}(Y, U), \text{ we derive:}$$

$$\text{cov}(x_1, x_3) = \lambda_1 \cdot \lambda_3 \cdot \text{cov}(F_1, F_2)$$

In order to reproduce the observed covariance as closely as possible, it can be seen from this equation that when the lambdas are under-estimated, the covariance between the two factors is necessarily over-estimated.

Mean Squared Error (MSE) on the factor correlation (psi)

ML-CFA leads to low MSE when the factor loadings are moderate to high (0.60 and 0.80), but to a higher MSE when the factor loadings are weak (0.40).

Bayesian CFA leads to low MSE on the factor correlation when the prior is non-informative. Rather unexpectedly, the MSE on the factor correlation is higher when a correctly specified informative prior is used on the factor loadings (not on the factor correlation).

Statistical power for the factor correlation

Statistical power remains of course a function of sample size and the population value of the parameter to estimate. When investigating the results of the ML-CFA, higher sample sizes and higher factor correlation yield higher statistical power.

Bayesian CFA does better than ML-CFA at lower sample sizes, if the priors on the factor loadings are informative. Non-informative priors on the factor loadings consistently return lower power on the factor loadings than ML-CFA or informative Bayes. It should be noted that miss specified priors on the factor loadings, which led to under-estimated factor loadings and over-estimated factor correlations, may be partly responsible for the higher statistical power obtained for informative Bayes at true factor loadings of 0.60 and 0.80 at lower sample sizes.

Error-free runs

Just like for the 1-factor model, the results are straightforward. While ML-CFA runs into problems at low sample sizes and weak to moderate factor loadings, Bayesian CFA consistently runs without errors.

Overall conclusions Monte Carlo study 2

The conclusions from Monte Carlo study 2 are similar to those from Monte Carlo study 1. ML-CFA performs well at large sample sizes and/or when the factor loadings are moderate to high (0.60 or 0.80). When either the sample size or the factor loadings decline, Bayesian CFA starts to bring clear advantages. The most obvious advantage is that the model fitting process can be executed without errors (model non-convergence or problematic parameter estimates). Well specified, informative priors can aid in keeping bias down and statistical power up. One has to be cautious about using informative priors though. From the results, it was clear that informative priors which under estimate the true factor loadings may lead to large bias in other model parameters, in this case, the factor correlation. So while Bayesian CFA may bring advantages, it also requires a careful consideration of which prior to choose for the model parameters. The current study did not investigate the use of an informative prior on the factor correlation, which should of course also influence the results.

$$\text{cov}(x_1, x_3) = \text{cov}(\lambda_1 \cdot F_1 + \varepsilon_1, \lambda_3 \cdot F_2 + \varepsilon_3)$$

$$\text{cov}(x_1, x_3) = \lambda_1 \cdot \lambda_3 \cdot \text{cov}(F_1, F_2) + \lambda_1 \cdot \text{cov}(F_1, \varepsilon_3) + \lambda_3 \cdot \text{cov}(\varepsilon_1, F_2) + \text{cov}(\varepsilon_1, \varepsilon_3)$$

$$\text{cov}(x_1, x_3) = \lambda_1 \cdot \lambda_3 \cdot \text{cov}(F_1, F_2)$$

(the simplification arises because the covariance between factors and errors are assumed to be zero, as well as covariances between errors).

6. DISCUSSION AND CONCLUSIONS

The results from both Monte Carlo studies generally confirm two anticipated key points. The first is that Bayesian CFA models can be run in small samples without running into problems of model non-convergence or impossible values, and the second is that the priors play a more important role as the sample size decreases. Non-informative priors may over-estimate factor loadings and reduce statistical power, especially in small samples. Therefore, one should seriously consider using well specified informative priors in small samples to take full advantage of the potential benefits of the Bayesian model estimation method without distorting the results. One should also pay attention to the fact that using miss specified priors may influence not only the involved parameter, but also other parts of the model, as shown in the 2-factor model. In that model, the priors on the factor loadings were miss specified, which not only influenced the estimation of the factor loadings but also the correlation between the two factors.

This Monte Carlo study remained relatively small in that only two simple CFA models were evaluated. More complex CFA models, that include more factors, more observed variables, unequal factor loadings, cross-loadings, error covariances, etc. might behave differently. Also, the simulations generated Normally distributed data. Empirical data may of course deviate from the Normal distribution, which may limit the findings from this study. Another limitation of the current study is that it did not perform a direct comparison between ML and Bayes at factor loadings of 0.60 or 0.80 with correctly specified informative priors. Further research is also needed to evaluate SEM models (i.e. models that contain a measurement part and a structural part, see e.g. Lee & Song, 2004).

The current study focused mainly on the small sample advantage of Bayesian CFA, but there are various other potential benefits to using a Bayesian framework to perform CFA (or SEM), which have not been addressed here. Muthén and Asparouhov (2012) highlighted the ability to specify approximate zeroes rather than exact zeroes in models. As explained in the introduction, the fact that some items are not supposed to load on specific factors is expressed in ML-CFA by fixing a factor loading at zero. Likewise, if no error covariances are expected, these are fixed at zero as well. Muthén and Asparouhov (2012) argued that these restrictions may be unnecessarily strict, leading to rejection of the model and a series of model modifications which may capitalize on chance. They proposed to replace selected exact zeroes with approximate zeroes using informative small variance priors in a substantively driven fashion.

The potential benefit of such an approach is illustrated by a study by Fong and Ho (2013). They evaluated a variety of latent structures for the Hospital Anxiety and Depression Scale (HADS) using ML-CFA and Bayesian CFA. They argue that the ML-CFA methods imposed overly strict model constraints, which triggers model modifications that capitalize on chance, and ultimately leading to ambiguous findings in the literature with very different factor structures for the HADS being proposed. They further showed that using a Bayesian method that included near-zero cross-loadings and residual correlations led to a simple two-factor structure that tapped into anxiety and depression as originally intended.

Bayesian CFA can thus bring several advantages. Besides performing well in small samples, Bayesian CFA might be an interesting tool to (re-)evaluate the latent structure of measurement instruments. The results from the current study however also raise awareness of the potential influence of miss specified priors and their potential effects on other parts of the model.

7. REFERENCES

- Asparhouhov, T. & Muthén, B. (2010a). *Bayesian Analysis Using Mplus: Technical Implementation*. Retrieved from <http://www.statmodel.com/download/Bayes3.pdf> on 13-August-2014.
- Asparhouhov, T. & Muthén, B. (2010b). *Bayesian Analysis of Latent Variable Models using Mplus*. Retrieved from <http://www.statmodel.com/download/BayesAdvantages18.pdf> on 13-August-2014.
- Centers for Medicare (2012). *Medicare Health Outcomes Survey. 2009-2011 Cohort 12 Analytic Public Use File Data User's Guide*. Technical report prepared by Health Services Advisory Group. Retrieved from <http://www.hosonline.org/Content/UsersGuide.aspx> on 05-June-2013.
- De Winter, J. C. F., Dodou, D., and P. A. Wieringa (2009). Exploratory Factor Analysis With Small Sample Sizes. *Multivariate Behavioral Research*, 44, 147-181.
- Fong, T.C.T & Ho, R.T.H. (2013). *Factor analyses of the Hospital Anxiety and Depression Scale: a Bayesian structural equation modeling approach*. Qual Life Res. DOI 10.1007/s11136-013-0429-2.
- Hatcher, L. (1994). *A Step-by-Step Approach to Using the SAS® System for Factor Analysis and Structural Equation Modeling*, Cary, NC: The SAS Institute.
- Heerwegh, D. (2013). *Around the world in three statistical models: determining the level of measurement invariance across countries of a PRO instrument*. Paper presented at PhUSE 2013, Brussels, 13-15 October 2013.
- Lee, S.-Y. & Song, X.-Y. (2004) Evaluation of the Bayesian and Maximum Likelihood Approaches in Analyzing Structural Equation Models with Small Sample Sizes, *Multivariate Behavioral Research*, 39:4, 653-686, DOI: 10.1207/s15327906mbr3904_4
- Muthén, B. & Asparouhov, T. (2012). Bayesian SEM: A more flexible representation of substantive theory. *Psychological Methods*, 17, 313-335.
- Spiro, A. III, Rogers, W. H., Qian, S., and Kazis, L. E. (2004). *Imputing physical and mental summary scores (PCS and MCS) for the Veterans SF-12 Health Survey in the context of missing data*. Technical Report prepared by: The Health Outcomes Technologies Program, Health Services Department, Boston University School of Public Health, Boston, MA and The Institute for Health Outcomes and Policy, Center for Health Quality, Outcomes and Economic Research, Veterans Affairs Medical Center, Bedford, MA. 2004. Retrieved from www.hosonline.org/surveys/hos/download/HOS_Veterans_12_Imputation.pdf on 25 July 2013.
- 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*, 85(3), pp. 842-860.

8. RESULT TABLES FROM THE MONTE CARLO STUDIES

The results from the Monte Carlo studies are summarized in the tables below. Note that the models contained multiple observed variables (4 in Model 1 and 6 in Model 2), and therefore several lambdas. As described earlier, all factor loadings were set to the same population lambda in each condition. Therefore, the results are summarized across the lambdas. For instance, the “average estimated lambda” column reports the average of estimated parameter values for the in the model lambdas. More detailed results, which include the results for each individual lambda, can be obtained from the author.

Table 4. Results from study 1

Estimation Method	Prior	Condition	Sample Size	Population Lambda	Average estimated lambda	Average percent bias on lambda	Average MSE on lambda	Average power on lambda (%)	Error-free runs (%)
ML		1	200	0.8	0.80	-0.11	0.004	100	100
BAYES	N(0.000,infinity)	2	200	0.8	0.81	1.19	0.004	100	100
BAYES	N(0.400,0.050)	3	200	0.8	0.75	-5.76	0.005	100	100
BAYES	N(0.400,0.010)	4	200	0.8	0.66	-18.08	0.023	100	100
ML		5	200	0.6	0.60	-0.02	0.006	100	100
BAYES	N(0.000,infinity)	6	200	0.6	0.60	0.70	0.007	100	100
BAYES	N(0.400,0.050)	7	200	0.6	0.58	-3.97	0.006	100	100
BAYES	N(0.400,0.010)	8	200	0.6	0.52	-13.46	0.009	100	100
ML		9	200	0.4	0.40	0.90	0.017	92	99.2
BAYES	N(0.000,infinity)	10	200	0.4	0.39	-2.61	0.015	91	100
BAYES	N(0.400,0.050)	11	200	0.4	0.40	-0.98	0.008	98	100
BAYES	N(0.400,0.010)	12	200	0.4	0.40	-0.04	0.002	100	100
ML		13	100	0.8	0.80	-0.54	0.008	100	100
BAYES	N(0.000,infinity)	14	100	0.8	0.82	2.38	0.009	100	100
BAYES	N(0.400,0.050)	15	100	0.8	0.72	-9.81	0.011	100	100
BAYES	N(0.400,0.010)	16	100	0.8	0.60	-24.90	0.042	100	100
ML		17	100	0.6	0.60	-0.43	0.013	100	100
BAYES	N(0.000,infinity)	18	100	0.6	0.61	1.30	0.015	100	100
BAYES	N(0.400,0.050)	19	100	0.6	0.56	-7.17	0.010	100	100
BAYES	N(0.400,0.010)	20	100	0.6	0.49	-18.99	0.015	100	100
ML		21	100	0.4	0.42	3.76	0.055	62	91.1
BAYES	N(0.000,infinity)	22	100	0.4	0.38	-4.38	0.026	61	100
BAYES	N(0.400,0.050)	23	100	0.4	0.40	-1.00	0.011	90	100

Estimation Method	Prior	Condition	Sample Size	Population Lambda	Average estimated lambda	Average percent bias on lambda	Average MSE on lambda	Average power on lambda (%)	Error-free runs (%)
BAYES	N(0.400,0.010)	24	100	0.4	0.40	0.02	0.002	100	100
ML		25	50	0.8	0.79	-1.14	0.016	100	100
BAYES	N(0.000,infinity)	26	50	0.8	0.84	4.86	0.019	100	100
BAYES	N(0.400,0.050)	27	50	0.8	0.68	-15.48	0.023	100	100
BAYES	N(0.400,0.010)	28	50	0.8	0.54	-31.93	0.068	100	100
ML		29	50	0.6	0.60	-0.68	0.029	95	98.6
BAYES	N(0.000,infinity)	30	50	0.6	0.62	2.88	0.029	94	100
BAYES	N(0.400,0.050)	31	50	0.6	0.53	-11.85	0.016	99	100
BAYES	N(0.400,0.010)	32	50	0.6	0.45	-24.19	0.023	100	100
ML		33	50	0.4	0.44	10.41	0.118	38	70.7
BAYES	N(0.000,infinity)	34	50	0.4	0.38	-4.71	0.041	30	100
BAYES	N(0.400,0.050)	35	50	0.4	0.40	-0.53	0.011	77	100
BAYES	N(0.400,0.010)	36	50	0.4	0.40	0.13	0.001	100	100
ML		37	25	0.8	0.78	-2.61	0.031	99	98.2
BAYES	N(0.000,infinity)	38	25	0.8	0.88	10.13	0.045	99	100
BAYES	N(0.400,0.050)	39	25	0.8	0.62	-22.83	0.043	100	100
BAYES	N(0.400,0.010)	40	25	0.8	0.49	-38.64	0.098	100	100
ML		41	25	0.6	0.60	-0.13	0.072	74	85.1
BAYES	N(0.000,infinity)	42	25	0.6	0.64	6.74	0.060	61	100
BAYES	N(0.400,0.050)	43	25	0.6	0.49	-17.91	0.022	94	100
BAYES	N(0.400,0.010)	44	25	0.6	0.43	-28.31	0.030	100	100
ML		45	25	0.4	0.47	16.49	0.233	25	50.4
BAYES	N(0.000,infinity)	46	25	0.4	0.40	0.28	0.065	12	100
BAYES	N(0.400,0.050)	47	25	0.4	0.40	0.12	0.009	61	100
BAYES	N(0.400,0.010)	48	25	0.4	0.40	0.17	0.001	100	100

Table 5. Results from study 2

Estimation Method	Prior	Condition	Sample Size	Population psi	Estimated psi	Percent bias psi	MSE psi	Power psi (%)	Population lambda	Average estimated lambda	Average percent bias lambda	Average MSE lambda	Power lambda (%)	Error free runs (%)
ML		1	200	0.25	0.25	0.25	0.0061	87.3	0.8	0.80	-0.27	0.004	100	100
BAYES	N(0.000,infinity)	2	200	0.25	0.25	0.25	0.0062	85.5	0.8	0.81	0.92	0.004	100	100
BAYES	N(0.400,0.050)	3	200	0.25	0.26	0.26	0.0069	86.9	0.8	0.76	-5.03	0.005	100	100
BAYES	N(0.400,0.010)	4	200	0.25	0.30	0.30	0.0114	89.9	0.8	0.67	-16.87	0.020	100	100
ML		5	200	0.25	0.25	0.25	0.0119	63.5	0.6	0.60	-0.11	0.009	100	100
BAYES	N(0.000,infinity)	6	200	0.25	0.25	0.25	0.0117	60.8	0.6	0.60	0.17	0.009	100	100
BAYES	N(0.400,0.050)	7	200	0.25	0.27	0.27	0.0167	65	0.6	0.57	-4.37	0.007	100	100
BAYES	N(0.400,0.010)	8	200	0.25	0.51	0.51	0.187	72.3	0.6	0.49	-17.91	0.015	100	100
ML		9	200	0.25	0.25	0.25	0.0646	27.7	0.4	0.42	6.00	0.081	73	82.2
BAYES	N(0.000,infinity)	10	200	0.25	0.23	0.23	0.0344	21	0.4	0.38	-4.76	0.019	81	100
BAYES	N(0.400,0.050)	11	200	0.25	0.27	0.27	0.0427	24.4	0.4	0.39	-2.54	0.009	96	100
BAYES	N(0.400,0.010)	12	200	0.25	0.29	0.29	0.0518	27.7	0.4	0.40	-0.90	0.002	100	100
ML		13	200	0.4	0.40	0.40	0.0051	99.8	0.8	0.80	-0.26	0.004	100	100
BAYES	N(0.000,infinity)	14	200	0.4	0.40	0.40	0.0052	99.8	0.8	0.81	0.96	0.004	100	100
BAYES	N(0.400,0.050)	15	200	0.4	0.42	0.42	0.0061	99.8	0.8	0.76	-5.04	0.005	100	100
BAYES	N(0.400,0.010)	16	200	0.4	0.50	0.50	0.0298	100	0.8	0.66	-17.67	0.023	100	100
ML		17	200	0.4	0.40	0.40	0.0107	95.7	0.6	0.60	-0.14	0.008	100	100
BAYES	N(0.000,infinity)	18	200	0.4	0.40	0.40	0.0109	94.1	0.6	0.60	0.27	0.009	100	100
BAYES	N(0.400,0.050)	19	200	0.4	0.46	0.46	0.0316	94.8	0.6	0.57	-5.01	0.007	100	100
BAYES	N(0.400,0.010)	20	200	0.4	0.87	0.87	0.2823	96.6	0.6	0.46	-22.81	0.022	100	100
ML		21	200	0.4	0.40	0.40	0.0461	50.4	0.4	0.41	3.29	0.047	79	91.5
BAYES	N(0.000,infinity)	22	200	0.4	0.36	0.36	0.0346	47.9	0.4	0.38	-3.90	0.018	84	100
BAYES	N(0.400,0.050)	23	200	0.4	0.42	0.42	0.0418	55.3	0.4	0.39	-2.36	0.009	97	100
BAYES	N(0.400,0.010)	24	200	0.4	0.46	0.46	0.058	59.8	0.4	0.40	-1.18	0.002	100	100
ML		25	100	0.25	0.25	0.25	0.0125	62	0.8	0.80	-0.47	0.009	100	100
BAYES	N(0.000,infinity)	26	100	0.25	0.25	0.25	0.0129	55.7	0.8	0.81	1.80	0.009	100	100
BAYES	N(0.400,0.050)	27	100	0.25	0.27	0.27	0.0173	59.9	0.8	0.73	-8.94	0.011	100	100
BAYES	N(0.400,0.010)	28	100	0.25	0.50	0.50	0.1866	69.4	0.8	0.57	-28.40	0.059	100	100
ML		29	100	0.25	0.26	0.26	0.0268	41	0.6	0.60	0.01	0.022	99	97.4
BAYES	N(0.000,infinity)	30	100	0.25	0.24	0.24	0.0245	33.2	0.6	0.60	0.21	0.018	99	100

Estimation Method	Prior	Condition	Sample Size	Population psi	Estimated psi	Percent bias psi	MSE psi	Power psi (%)	Population lambda	Average estimated lambda	Average percent bias lambda	Average MSE lambda	Power lambda (%)	Error free runs (%)
BAYES	N(0.400,0.050)		31	100	0.25	0.43	0.43	0.1595	43	0.6	0.53	-11.79	99	100
BAYES	N(0.400,0.010)		32	100	0.25	0.77	0.77	0.4329	72	0.6	0.44	-27.39	100	100
ML			33	100	0.25	0.29	0.29	0.5248	16.2	0.4	0.44	10.94	44	55.7
BAYES	N(0.000,infinity)		34	100	0.25	0.21	0.21	0.0613	8.1	0.4	0.37	-7.88	48	100
BAYES	N(0.400,0.050)		35	100	0.25	0.31	0.31	0.1278	17.4	0.4	0.38	-4.39	83	100
BAYES	N(0.400,0.010)		36	100	0.25	0.35	0.35	0.1492	22.2	0.4	0.39	-1.58	100	100
ML			37	100	0.4	0.40	0.40	0.0105	95.2	0.8	0.80	-0.42	100	100
BAYES	N(0.000,infinity)		38	100	0.4	0.40	0.40	0.0108	92.8	0.8	0.81	1.85	100	100
BAYES	N(0.400,0.050)		39	100	0.4	0.45	0.45	0.0258	94.3	0.8	0.72	-9.40	100	100
BAYES	N(0.400,0.010)		40	100	0.4	0.88	0.88	0.287	97.2	0.8	0.52	-34.49	100	100
ML			41	100	0.4	0.40	0.40	0.0241	76.5	0.6	0.60	-0.17	99	98.9
BAYES	N(0.000,infinity)		42	100	0.4	0.39	0.39	0.0224	68.5	0.6	0.60	0.29	99	100
BAYES	N(0.400,0.050)		43	100	0.4	0.71	0.71	0.2005	77.9	0.6	0.51	-14.93	99	100
BAYES	N(0.400,0.010)		44	100	0.4	0.95	0.95	0.3288	92.8	0.6	0.44	-27.48	100	100
ML			45	100	0.4	0.46	0.46	0.2291	29.7	0.4	0.43	7.36	50	65.5
BAYES	N(0.000,infinity)		46	100	0.4	0.33	0.33	0.0586	21.2	0.4	0.37	-6.59	51	100
BAYES	N(0.400,0.050)		47	100	0.4	0.48	0.48	0.112	35.2	0.4	0.38	-4.10	85	100
BAYES	N(0.400,0.010)		48	100	0.4	0.53	0.53	0.1395	42.7	0.4	0.39	-1.78	100	100
ML			49	50	0.25	0.25	0.25	0.026	38.6	0.8	0.79	-0.88	100	98.1
BAYES	N(0.000,infinity)		50	50	0.25	0.25	0.25	0.0285	29.5	0.8	0.83	3.71	100	100
BAYES	N(0.400,0.050)		51	50	0.25	0.47	0.47	0.211	41.5	0.8	0.64	-20.09	99	100
BAYES	N(0.400,0.010)		52	50	0.25	0.82	0.82	0.566	88.9	0.8	0.46	-42.30	100	100
ML			53	50	0.25	0.25	0.25	0.0556	28.1	0.6	0.61	1.67	86	78.9
BAYES	N(0.000,infinity)		54	50	0.25	0.25	0.25	0.0584	16.8	0.6	0.60	0.16	85	100
BAYES	N(0.400,0.050)		55	50	0.25	0.62	0.62	0.4869	59.9	0.6	0.46	-23.69	89	100
BAYES	N(0.400,0.010)		56	50	0.25	0.68	0.68	0.6024	79.5	0.6	0.42	-30.68	100	100
ML			57	50	0.25	0.37	0.37	0.4681	10.5	0.4	0.46	13.97	27	33.3
BAYES	N(0.000,infinity)		58	50	0.25	0.17	0.17	0.0954	4.6	0.4	0.37	-8.08	22	100
BAYES	N(0.400,0.050)		59	50	0.25	0.35	0.35	0.3714	24.3	0.4	0.38	-6.10	65	100
BAYES	N(0.400,0.010)		60	50	0.25	0.37	0.37	0.3736	31.3	0.4	0.39	-1.83	100	100
ML			61	50	0.4	0.40	0.40	0.0221	77.5	0.8	0.79	-0.85	100	98.5

Estimation Method	Prior	Condition	Sample Size	Population psi	Estimated psi	Percent bias psi	MSE psi	Power psi (%)	Population lambda	Average estimated lambda	Average percent bias lambda	Average MSE lambda	Power lambda (%)	Error free runs (%)
BAYES	N(0.000,infinity)	62	50	0.4	0.40	0.40	0.024	65.4	0.8	0.83	3.75	0.020	100	100
BAYES	N(0.400,0.050)	63	50	0.4	0.77	0.77	0.2453	77.8	0.8	0.60	-24.68	0.059	100	100
BAYES	N(0.400,0.010)	64	50	0.4	0.97	0.97	0.3532	97.7	0.8	0.47	-41.37	0.114	100	100
ML		65	50	0.4	0.42	0.42	0.0964	52.3	0.6	0.60	0.71	0.051	88	85
BAYES	N(0.000,infinity)	66	50	0.4	0.39	0.39	0.0511	38.4	0.6	0.60	0.37	0.032	86	100
BAYES	N(0.400,0.050)	67	50	0.4	0.84	0.84	0.3255	76.4	0.6	0.46	-22.73	0.033	93	100
BAYES	N(0.400,0.010)	68	50	0.4	0.89	0.89	0.3757	89.5	0.6	0.42	-29.80	0.034	100	100
ML		69	50	0.4	0.54	0.54	0.4153	15.5	0.4	0.45	12.70	0.184	31	39.3
BAYES	N(0.000,infinity)	70	50	0.4	0.27	0.27	0.1005	8	0.4	0.37	-6.94	0.039	24	100
BAYES	N(0.400,0.050)	71	50	0.4	0.53	0.53	0.2865	32.4	0.4	0.38	-5.42	0.011	67	100
BAYES	N(0.400,0.010)	72	50	0.4	0.55	0.55	0.2955	40.7	0.4	0.39	-1.66	0.001	100	100
ML		73	25	0.25	0.26	0.26	0.0577	31.6	0.8	0.79	-1.81	0.042	98	80.6
BAYES	N(0.000,infinity)	74	25	0.25	0.25	0.25	0.0623	14.2	0.8	0.87	8.29	0.047	97	100
BAYES	N(0.400,0.050)	75	25	0.25	0.65	0.65	0.6506	85	0.8	0.50	-38.01	0.118	84	100
BAYES	N(0.400,0.010)	76	25	0.25	0.68	0.68	0.6491	91	0.8	0.43	-46.01	0.137	100	100
ML		77	25	0.25	0.29	0.29	0.1688	26	0.6	0.63	5.56	0.174	59	44.7
BAYES	N(0.000,infinity)	78	25	0.25	0.24	0.24	0.1038	5.6	0.6	0.61	1.64	0.060	47	100
BAYES	N(0.400,0.050)	79	25	0.25	0.54	0.54	0.703	68.9	0.6	0.43	-28.50	0.041	77	100
BAYES	N(0.400,0.010)	80	25	0.25	0.52	0.52	0.7348	79.8	0.6	0.41	-31.96	0.038	100	100
ML		81	25	0.25	0.57	0.57	2.4522	13.1	0.4	0.45	12.70	0.287	21	18.5
BAYES	N(0.000,infinity)	82	25	0.25	0.13	0.13	0.1048	0.7	0.4	0.40	-0.82	0.054	8	100
BAYES	N(0.400,0.050)	83	25	0.25	0.31	0.31	0.7012	36.3	0.4	0.38	-5.15	0.008	51	100
BAYES	N(0.400,0.010)	84	25	0.25	0.33	0.33	0.6817	46.5	0.4	0.39	-1.36	0.001	100	100
ML		85	25	0.4	0.41	0.41	0.0488	54.3	0.8	0.79	-1.85	0.038	98	82.1
BAYES	N(0.000,infinity)	86	25	0.4	0.41	0.41	0.0531	33	0.8	0.87	8.23	0.047	97	100
BAYES	N(0.400,0.050)	87	25	0.4	0.88	0.88	0.4019	91.1	0.8	0.51	-35.84	0.103	91	100
BAYES	N(0.400,0.010)	88	25	0.4	0.89	0.89	0.3963	94.9	0.8	0.44	-45.03	0.132	100	100
ML		89	25	0.4	0.45	0.45	0.1475	37.2	0.6	0.62	3.71	0.133	63	52.7
BAYES	N(0.000,infinity)	90	25	0.4	0.38	0.38	0.0907	14.5	0.6	0.61	2.30	0.060	48	100
BAYES	N(0.400,0.050)	91	25	0.4	0.75	0.75	0.4949	76.3	0.6	0.44	-26.99	0.038	81	100
BAYES	N(0.400,0.010)	92	25	0.4	0.76	0.76	0.4943	86.2	0.6	0.41	-31.45	0.037	100	100

Estimation Method	Prior	Condition	Sample Size	Population psi	Estimated psi	Percent bias psi	MSE psi	Power psi (%)	Population lambda	Average estimated lambda	Average percent bias lambda	Average MSE lambda	Power lambda (%)	Error free runs (%)
ML			93	25	0.4	0.67	0.9326	17.1	0.4	0.46	15.01	0.307	23	21.2
BAYES	N(0.000,infinity)		94	25	0.4	0.20	0.128	1.2	0.4	0.40	-0.04	0.055	8	100
BAYES	N(0.400,0.050)		95	25	0.4	0.47	0.5909	39.7	0.4	0.38	-4.53	0.009	52	100
BAYES	N(0.400,0.010)		96	25	0.4	0.47	0.5833	51	0.4	0.40	-1.18	0.001	100	100