

Bayesian SEM:
A more flexible representation of substantive theory

Bengt Muthén & Tihomir Asparouhov

April 14, 2011

Abstract

This paper proposes a new approach to factor analysis and structural equation modeling using Bayesian analysis. The new approach replaces parameter specifications of exact zeros with approximate zeros based on informative, small-variance priors. It is argued that this produces an analysis that better reflects substantive theories. The proposed Bayesian approach is particularly beneficial in applications where parameters are added to a conventional model such that a non-identified model is obtained if maximum-likelihood estimation is applied. This approach is useful for measurement aspects of latent variable modeling such as with CFA and the measurement part of SEM. Two application areas are studied, cross-loadings and residual correlations in CFA. The approach encompasses three elements: Model testing, model estimation, and model modification. Monte Carlo simulations and real data are analyzed using Mplus.

1 Introduction

This paper proposes a new approach to factor analysis and structural equation modeling using Bayesian analysis. It is argued that current analyses using maximum likelihood (ML) and likelihood-ratio χ^2 testing apply unnecessarily strict models to represent hypotheses derived from substantive theory. This often leads to rejection of the model (see, e.g., Marsh et al, 2009) and a series of model modifications that may capitalize on chance (see, e.g. McCallum et al. 1992). The hypotheses are reflected in parameters fixed at zero. Examples include zero cross-loadings and zero residual correlations in factor analysis.

The new approach is intended to produce an analysis that better reflects substantive theories. It does so by replacing the parameter specification of exact zeros with approximate zeros. The new approach uses Bayesian analysis to specify informative priors for such parameters. In key applications, freeing these parameters in a conventional analysis, the model would not be identified. The Bayesian analysis, however, identifies the model by substantively-driven small-variance priors. Model testing is carried out using posterior predictive checking which is found to be less sensitive than likelihood-ratio χ^2 testing to ignorable degrees of model misspecification. A side product of the proposed approach is information to modify the model in line with the use of modification indices in ML analysis. ML modification indices inform about model improvement when a single parameter is freed and can lead to a long series of modifications. In contrast, the proposed approach informs about model modification when all parameters are freed and does so in a single-step analysis.

Section 2 presents a brief overview of the Bayesian analysis framework that is used. Sections 4 and 5 present two studies that illustrate the new approach. Each study consists of a real-data example showing the problem, the proposed Bayesian solution for the real-data problem, and simulations showing how well the method works. Study 1 considers factor analysis where cross-loadings make simple structure CFA inadequate. As

an example, a re-analysis is made of the classic Holzinger-Swineford mental abilities data, where a simple structure does not fit well by ML CFA standards. Study 2 considers residual correlations in factor analysis which make a factor model inadequate. As an example, the big-five factor model is analyzed using an instrument administered in a British household survey, where the hypothesized five-factor pattern is not well recovered by ML CFA or EFA due to many minor factors. All analyses are carried out by Bayesian analysis in Mplus (Muthén & Muthén, 1998-2010) and scripts are available at www.statmodel.com. Section 6 concludes.

2 Bayesian analysis

There are many books on Bayesian analysis and most are quite technical. Gelman et al. (2004) provides a good general statistical description, whereas Kruschke (2010) and Lynch (2010) give somewhat more introductory accounts. Press (2003) discusses Bayesian factor analysis. Lee (2007) gives a discussion from a structural equation modeling perspective. Schafer (1997) gives a statistical discussion from a missing data and multiple imputation perspective, whereas Enders (2010) gives an applied discussion of these same topics. Statistical overview articles include Gelfand et al. (1990) and Casella and George (1992). Overview articles of an applied nature and with a latent variable focus include Scheines et al. (1999), Rupp et al. (2004), and Yuan and MacKinnon (2009).

Bayesian analysis is firmly established in mainstream statistics and its popularity is growing. Part of the reason for the increased use of Bayesian analysis is the success of new computational algorithms referred to as Markov chain Monte Carlo (MCMC) methods. Outside of statistics, however, applications of Bayesian analysis lag behind. One possible reason is that Bayesian analysis is perceived as difficult to do, requiring complex statistical specifications such as those used in the flexible, but technically-oriented general Bayes

program WinBUGS. These observations were the background for developing Bayesian analysis in Mplus (Muthén & Muthén, 1998-2010). In Mplus, simple analysis specifications with convenient defaults allow easy access to a rich set of analysis possibilities. For a technical description of the Mplus implementation, see Asparouhov and Muthén (2010a).

Four key points motivate taking an interest in Bayesian analysis:

1. More can be learned about parameter estimates and model fit
2. Better small-sample performance can be obtained and large-sample theory is not needed
3. Analyses can be made less computationally demanding
4. New types of models can be analyzed

Point 1 is illustrated by parameter estimates that do not have a normal distribution. An example is an indirect effect $a \times b$ in a mediation model (MacKinnon, 2008). Maximum-likelihood (ML) gives a parameter estimate and its standard error and assumes that the distribution of the parameter estimate is normal based on asymptotic (large-sample) theory. In contrast, Bayes does not rely on large-sample theory and provides the whole distribution, referred to as the posterior distribution, not assuming that it is normal. The ML confidence interval $Estimate \pm 1.96 \times SE$ assumes a symmetric distribution, whereas the Bayesian credibility interval based on the percentiles of the posterior allows for a strongly skewed distribution. Bayesian exploration of model fit can be done in a flexible way using Posterior Predictive Checking (PPC; see, e.g., Gelman et al., 1996; Gelman et al., 2004, Chapter 6; Lee, 2007, Chapter 5; Scheines et al., 1999). Any suitable test statistics for the observed data can be compared to statistics based on simulated data obtained via draws of parameter values from the posterior distribution, avoiding statistical assumptions about the distribution of the test statistics.

Point 2 is illustrated by better Bayesian small-sample performance for factor analyses prone to Heywood cases and better performance when a small number of clusters are

analyzed in multilevel models. This, however, requires a judicious choice of prior. For examples, see Asparouhov and Muthén (2010b).

Point 3 may be of interest for an analyst who is hesitant to move from ML estimation to Bayesian estimation. Many models are computationally cumbersome or impossible using ML, such as with categorical outcomes and many latent variables resulting in many dimensions of numerical integration. Such an analyst may view the Bayesian analysis simply as a computational tool for getting estimates that are analogous to what would have been obtained by ML had it been feasible. This is obtained with diffuse priors, in which case ML and Bayesian results are expected to be close in large samples (Browne & Draper, 2006; p. 505).

Point 4 is exemplified by models with a very large number of parameters or where ML does not provide a natural approach. Examples of the former include image analysis (see, e.g., Green, 1996)) and examples of the latter include random change-point analysis (see, e.g., Dominicus et al., 2008). The Bayesian SEM approach proposed in this paper is a further example of the new type of models that can be analyzed.

2.1 Bayesian estimation

Frequentist analysis (e.g., maximum likelihood) and Bayesian analysis differ by the former viewing parameters as constants and the latter as variables. Bayesian analysis uses the term prior to refer to the parameter distribution. Priors can be diffuse (non-informative) or informative. Information about priors can be built up from a sequence of formulating hypotheses from theory, carrying out pilot studies, and revising hypotheses. An example that will be discussed in detail later on is drawn from the classic Holzinger-Swineford factor analysis study (Holzinger & Swineford, 1939). Using two new samples of subjects, the authors used a set of well-known tests to measure factors that had been derived from several previous factor analyses. A factor loading pattern was hypothesized corresponding

to a large number of ignorable cross-loadings. Although these authors did not invoke Bayesian analysis, their careful groundwork could have been used in the analysis of the two new samples to specify informative priors centered around zero for the cross-loadings.

Maximum likelihood (ML) finds estimates by maximizing a likelihood computed for the data. In Bayesian analysis data inform about a parameter and modify the prior into a posterior that gives the Bayesian estimate. This is illustrated in Figure 1 which shows distributions for a prior and a posterior for a parameter, together with the likelihood. The likelihood can be thought of as the distribution of the data given a parameter value. In Figure 1 the major portion of the prior distribution has a lower parameter value than that at the peak of the likelihood. The posterior is obtained as a compromise between the prior and the likelihood.

Priors can be non-informative or informative. A non-informative prior, also called a diffuse prior, can for example have a uniform distribution or have normal distribution with a large variance. A large variance reflects large uncertainty in the parameter value. With a large prior variance the likelihood contributes relatively more information to the formation of the posterior and the estimate is closer to a maximum-likelihood estimate.

[Figure 1 about here.]

2.1.1 Bayes theorem

Formally, the formation of a posterior draws on Bayes Theorem. Consider the probabilities of events A and B, $P(A)$ and $P(B)$. By probability theory the joint event A and B can be expressed in terms of conditional and marginal probabilities:

$$P(A, B) = P(A|B) P(B) = P(B|A) P(A). \quad (1)$$

Dividing by $P(A)$ it follows that

$$P(B|A) = \frac{P(A|B) P(B)}{P(A)}, \quad (2)$$

which is Bayes Theorem. Applied to modeling, let data take the role of A and the parameter values take the role of B. The posterior can then be expressed symbolically as

$$posterior = parameters|data \quad (3)$$

$$= \frac{data|parameters \times parameters}{data} \quad (4)$$

$$= \frac{likelihood \times prior}{data} \quad (5)$$

$$\propto likelihood \times prior, \quad (6)$$

where $\propto$ means proportional to, recognizing that the data do not contain parameters so that this term does not need updating when iteratively finding the posterior.

The prior distribution is the key element of Bayesian analysis. Priors reflect prior beliefs in likely parameter values before collecting new data. These beliefs may come from substantive theory and previous studies of similar populations. The priors modify the likelihood to obtain the posterior distribution. The Bayesian estimates are obtained as means, modes, or medians of their posterior distributions. The posterior distribution is obtained via Markov Chain Monte Carlo (MCMC) algorithms. MCMC is briefly outlined in the Appendix and will not be discussed here. The reader is instead referred to the Bayesian literature. The Appendix also discusses determination of convergence of the MCMC process. For a technical description of the Mplus Bayesian implementation, see Asparouhov and Muthén (2010a).

2.1.2 Model fit

Model fit assessment is possible using Posterior Predictive Checking (PPC) introduced in Gelman, Meng and Stern (1996). With continuous outcomes, PPC as implemented in Mplus builds on the standard likelihood-ratio χ^2 statistic in mean- and covariance-structure modeling. This PPC procedure is described in Scheines et al. (1999) and Asparouhov and Muthén (2010a, b) and is briefly reviewed here. Gelman et al. (2004) presents a more general discussion of PPC, not tied to likelihood-ratio χ^2 .

A Posterior Predictive p-value (PPP) of model fit can be obtained via a fit statistic f based on the usual likelihood-ratio χ^2 test of an H_0 model against an unrestricted H_1 model. Low PPP indicates poor fit. Let $f(Y, X, \boldsymbol{\pi}_i)$ be computed for the data Y, X using the parameter values at MCMC iteration i . Here, X denotes covariates that are conditioned on in the analysis. At iteration i , generate a new data set Y_i^* of synthetic or replicated data of the same sample size as the original data. In this generation the parameter values at iteration i are used. For these replicated data the fit statistic $f(Y_i^*, X, \boldsymbol{\pi}_i)$ is computed. This data generation and fit statistic computation is repeated over the n iterations, after which PPP is approximated by the proportion of iterations where

$$f(Y, X, \boldsymbol{\pi}_i) < f(Y_i^*, X, \boldsymbol{\pi}_i). \quad (7)$$

In the Mplus implementation (Asparouhov & Muthén, 2010a) PPP is computed using every 10th iteration among the iterations used to describe the posterior distribution of parameters. A 95% confidence interval is produced for the difference in the f statistic for the real and replicated data. A positive lower limit is in line with a low PPP and indicates poor fit. An excellent-fitting model is expected to have a PPP value around 0.5 and an f statistic difference of zero falling close to the middle of the confidence interval.

It should be noted that the PPP value does not behave like a p-value for a χ^2 test of

model fit (see also Hjort et al., 2006). The type I error is not 5% for a correct model. There is not a theory for how low PPP can be before the model is significantly ill-fitting at a certain level. In this sense, PPP is more akin to a structural equation modeling fit index rather than a χ^2 test. Empirical experience with different models and data has to be established for PPP and some simulation studies are presented here. From these simulations and further ones in Asparouhov and Muthén (2010b), however, the usual approach of using p-values of 0.05 or 0.01 appears reasonable. This warrants further investigations, however.

3 BSEM: A more flexible SEM approach

A new approach to structural equation modeling based on Bayesian analysis is described below. It is intended to produce an analysis that better reflects the researcher's theories and prior beliefs. It does so by systematically using informative priors for parameters that should not be freely estimated according to the the researcher's theories and prior beliefs. In a frequentist analysis such parameters are typically fixed at zero or are constrained to be equal to other parameters. In key applications, freeing these parameters would in fact produce a non-identified model. The Bayesian analysis, however, identifies the model by substantively-driven small-variance priors. The proposed approach is referred to as BSEM for Bayesian structural equation modeling. It should be recognized, however, that BSEM refers to the specific Bayesian approach proposed here of using informative, small-variance priors to reflect the researcher's theories and prior beliefs. Typically, this would be combined with the use of non-informative priors for parameters that would not be restricted in a corresponding ML analysis. For example, major loadings would have a normal prior with a very large variance.

The BSEM approach of using informative priors is applicable to any constrained

parameter in an SEM. This paper focuses on parameters in the measurement part, but restrictions in the structural part can also be considered. Two types of model features are considered, cross-loadings in CFA and residual correlations in CFA. Further examples are considered in the Conclusion section.

3.1 Informative priors for cross-loadings in CFA

An analyst who is used to frequentist methods such as ML may at first feel uncomfortable specifying informative priors. It is argued here, however, that a user of CFA is in a sense already engaged in specifying such priors. Consider the confirmatory factor analysis model for an observed p -dimensional vector $\mathbf{y}_i$ of factor indicators for individual i ,

$$\begin{aligned}\mathbf{y}_i &= \boldsymbol{\nu} + \mathbf{\Lambda} \boldsymbol{\eta}_i + \boldsymbol{\epsilon}_i, \\ E(\mathbf{y}_i) &= \boldsymbol{\nu} + \mathbf{\Lambda} \boldsymbol{\alpha}, \\ V(\mathbf{y}_i) &= \mathbf{\Lambda} \boldsymbol{\Psi} \mathbf{\Lambda}' + \boldsymbol{\Theta},\end{aligned}\tag{8}$$

where $\boldsymbol{\nu}$ is an intercept vector, $\mathbf{\Lambda}$ a loading matrix, $\boldsymbol{\eta}_i$ is an m -dimensional factor vector, $\boldsymbol{\epsilon}_i$ is a residual vector, $\boldsymbol{\alpha}$ is a factor mean vector, $\boldsymbol{\Psi}$ is a factor covariance matrix, and $\boldsymbol{\Theta}$ is a residual covariance matrix. Here, $\boldsymbol{\epsilon}$ and $\boldsymbol{\eta}$ are assumed normally distributed and uncorrelated.

Drawing on substantive theory, zero cross-loadings in $\mathbf{\Lambda}$ are specified for the factor indicators that are hypothesized to not be influenced by certain factors. An exact zero loading can be viewed as a prior distribution that has mean zero and variance zero. A prior that probably more accurately reflects substantive theory uses a mean of zero and a normal distribution with small variance. Figure 2 shows an example where a loading $\lambda \sim N(0, 0.01)$ so that 95% of the loading variation is between -0.2 and $+0.2$. Using standardized factor indicators and factors, a loading of 0.2 is considered a small loading,

implying that this prior essentially says that the cross-loading is close zero, but not exactly zero. The prior is strongly informative, but it is not assumed that the parameter is zero.

[Figure 2 about here.]

In frequentist analysis freeing all cross-loadings in a CFA model such as Table 2 leads to a non-identified model because the m^2 restrictions, where m is the number of factors, necessary to eliminate indeterminacies are not present (see, e.g., Hayashi & Marcoulides, 2006). Using small-variance priors for all cross-loadings, however, brings information into the analysis which avoids the non-identification problem. The choice of variance for the prior should correspond to the researcher's theories and prior beliefs. As stated above, the variance of 0.01 produces a prior where 95% lies between -0.2 and $+0.2$. Other choices are shown in Table 1.

[Table 1 about here.]

A smaller variance may not let cross-loadings escape sufficiently from their zero prior mean, producing a worse Posterior Predictive p-value. A larger variance may let a cross-loading have too large of a probability of having a substantial value. For example, a variance of 0.08 corresponds to 95% lying between -0.55 and $+0.55$, which on a standardized variable scale approaches a major loading size. When the variance is increased, the prior contributes less information so that the model gets closer to being non-identified which eventually causes non-convergence of the MCMC algorithm. It should be noted that the prior variance should be determined in relation to the scale of the observed and latent variables. A prior variance of 0.01 corresponds to small loadings for variables with unit variance, but it corresponds to a smaller loading for an observed variable with variance larger than one. This means that for convenience observed variables may be brought to a common scale either by multiplying them by constants, or standardizing if the model is scale free.

BSEM has an additional advantage. It produces posterior distributions for cross-loadings which can be used in line with modification indices to free parameters for which the credibility interval does not cover zero. Modification indices pertain to freeing only one parameter at a time and a long sequence of model modification is often needed, running the risk of capitalizing on chance (see, e.g., McCallum et al. 1992). In contrast, the small-variance prior approach provides information on model modification that considers the fixed parameters jointly in a single analysis. The relative benefits of the Bayesian approach to modifying the model compared to the use of modification indices with ML need further study, however.

3.2 Informative priors for residual covariances in CFA

An analogous idea can also be used to study residual correlations among factor indicators. In (8), the residual covariance matrix Θ is commonly assumed to be diagonal. Some residuals may, however, be correlated due to the omission of several minor factors. It is difficult to foresee which residuals should be covaried and freeing all of them leads to a non-identified model in the conventional ML framework. BSEM provides a possible approach to this problem.

Instead of assuming a diagonal residual covariance matrix, a more realistic covariance structure model may be expressed as

$$V(\mathbf{y}_i) = \mathbf{\Lambda} \mathbf{\Psi} \mathbf{\Lambda}' + \mathbf{\Omega} + \mathbf{\Theta}^*, \quad (9)$$

where $\mathbf{\Omega}$ is a covariance matrix for the minor factors, not assumed to be diagonal, and $\mathbf{\Theta}^*$ is a diagonal covariance matrix. Here, a freely estimated $\mathbf{\Omega}$ is not separately identified from $\mathbf{\Lambda}$, $\mathbf{\Psi}$, and $\mathbf{\Theta}^*$. In Bayesian analysis, however, $\mathbf{\Omega}$ can be given an informative prior using the inverse-Wishart distribution so that the posterior distribution can be obtained.

In this way, the diagonal and off-diagonal elements of $\mathbf{\Omega}$ are restricted to small values. This implies that the residual covariance matrix $\mathbf{\Omega} + \mathbf{\Theta}^*$ contains residual covariances that are allowed to deviate to a small extent from zero means. Sufficiently stringent priors for the off-diagonal elements are needed so that the essential correlations are channeled via $\mathbf{\Lambda} \mathbf{\Psi} \mathbf{\Lambda}'$. The sums on the diagonal of $\mathbf{\Omega} + \mathbf{\Theta}^*$ produce the residual variances. The inverse-Wishart distribution is described in the Appendix.

The BSEM approach for residual covariances outlined in connection with (9) will be referred to as Method 1. A more direct method, Method 2, applies an inverse-Wishart prior directly on $\mathbf{\Theta}$ in (8). This approach has been discussed in Press (2003; chapter 15). One advantage of Method 1 over Method 2 is that the prior for the total residual variance is not tied to the prior of the residual covariances because of the fact that the residual covariance has two components that have different priors. Method 2, however, is simpler to carry out. A disadvantage of both Method 1 and Method 2 is that particular residual covariance elements cannot be given their own priors. For example, an analysis may show that some residual covariances should be freely estimated with non-informative priors because they have 95% credibility intervals that do not cover zero. To this aim, Method 3 makes it possible to specify element-specific normal priors for the residual covariances. Mplus allows two different algorithms for Method 3, a random-walk algorithm (Chib & Greenberg, 1988) and a proposal prior algorithm (Asparouhov & Muthén, 2010a).

The choice of inverse-Wishart prior should be made to reflect prior beliefs in the potential magnitude of residual covariances. This is accomplished by using a sufficiently large choice for the degrees of freedom (df) of the inverse-Wishart distribution. To obtain a proper posterior where the marginal mean and variance is defined, $df \geq p + 4$ needs to be chosen, where p is the number of variables y . The prior means for the residual covariances can be chosen as zero and the degree of informativeness specified using the df which affects the marginal prior variance via $df - p$. For example, (26) of the Appendix shows that using

the inverse-Wishart prior $IW(\mathbf{I}, df)$ with $df = p + 6$ gives a prior standard deviation of 0.1, so that two standard deviations below and above the zero mean correspond to the residual covariance range of -0.2 to $+0.2$. The effect of priors is relative to the variances of the y s. For scale-free models, the variables may be standardized before analysis. For larger sample sizes, the prior needs to use a larger df to give the same effect.

Method 1 and 2 both use conjugate priors, that is, the posterior distribution of the covariance matrices is also of the Wishart family of distributions. This generally produces good convergence of the MCMC chain. Both versions of Method 3, random walk and proposal prior algorithm, are instead based on the Metropolis-Hastings algorithm and that generally yields somewhat worse convergence performance. The random walk algorithm has difficulty converging or converges very slowly when the variance-covariance matrix has a large number of parameters. However when a large number of parameters have small prior variance the convergence is fast. The proposal prior algorithm generally works well, but not when the prior variance is very small.

When estimating the above models all the algorithms are applied to extreme conditions with a nearly unidentified model and near zero prior variance. Convergence of the MCMC sequence should be carefully evaluated and automated convergence criteria such as PSR are not always reliable. Instead, the stability of the parameter values across the iterations should be studied. The models should be estimated with a large number of MCMC iterations, for example, 50,000 or 100,000.

4 Study 1: Cross-loadings in CFA

4.1 Holzinger-Swineford mental abilities example: ML analysis

The first example uses data from the classic 1939 factor analysis study by Holzinger and Swineford (1939). Twenty-six tests intended to measure a general factor and five specific factors were administered to seventh and eighth grade students in two schools, the Grant-White school ($n = 145$) and the Pasteur school ($n = 156$). Students from the Grant-White school came from homes where the parents were mostly American-born, whereas students from the Pasteur school came largely from working-class parents of whom many were foreign-born and where their native language was used at home.

Factor analyses of these data have been described e.g. by Harman (1976; pp. 123-132) and Gustafsson (2002). Of the 26 tests, nineteen were intended to measure four domains, five measured general deduction, and two were revisions/new test versions. Typically, the last two are not analyzed. Excluding the five general deduction tests, 19 tests measuring four domains are considered here, where the four domains are spatial ability, verbal ability, speed, and memory. The design of the measurement of the four domains by the 19 tests is shown in the factor loading pattern matrix of Table 2. Here, an X denotes a free loading to be estimated and 0 a fixed, zero loading. This corresponds to a simple structure CFA model with variable complexity one, that is, each variable loads on only one factor.

[Table 2 about here.]

Using maximum-likelihood estimation, the model fit using both confirmatory factor analysis (CFA) and exploratory factor analysis (EFA) is reported in Table 3 for both the Grant-White and Pasteur schools. It is seen that the CFA model is rejected by the likelihood-ratio χ^2 test in both samples. Given the rather small sample sizes, one cannot

attribute the poor fit to the χ^2 test being overly sensitive to small misspecifications due to a large sample size as is often done. The common fit indices RMSEA and CFI are also not at acceptable levels. In contrast to the CFA, Table 3 shows that the EFA model fits the data well in both schools.

[Table 3 about here.]

Table 4 shows the EFA factor solution for both schools using the Geomin rotation. The Quartimin rotation gives similar results. For a description of these rotations, see, e.g., Asparouhov and Muthén (2009) and Browne (2001). The table shows that the major loadings of the EFA correspond to the hypothesized four-factor loading pattern (bolded entries). Several of the tests, however, also have significant cross-loadings on other factors (significance on the 5% level marked with asterisks). There are six significant cross-loadings for the Grant-White solution and nine for the Pasteur solution. This explains the poor fit of the CFA model.

The question arises how to go beyond postulating only the number of factors as in EFA and maintain the essence of the hypothesized factor loading pattern without resorting to an exploratory rotation. Cross-loadings need to be allowed to some degree, but a model with freely estimated cross-loadings is not identified. The proposed Bayesian solution to this problem is presented next. As an intermediate step, however, it is instructive to consider the EFA alternative of target rotation (Browne, 2001; Asparouhov & Muthén, 2009). Here, a rotation is chosen to match certain zero target loadings using a least-squares fitting function. Target rotation is similar to BSEM in that it replaces mechanical rotation with rotation guided by the researcher’s judgement, in this case using zero targets for cross-loadings. Target rotation is also similar in that the fitting function can result in non-zero values for the targets. It is different from BSEM by not allowing user-specified stringency of closeness to zero by varying the prior variance, replacing that with least-squares fitting. It is also different from BSEM because specifying $m - 1$ zeros for each of the m factors

gives the same model fit as specifying more zeros. For the two Holzinger-Swineford data sets, applying target rotation with zero targets for all cross-loadings gives results similar to EFA using Geomin or Quartimin rotation. Four more significant cross-loadings are obtained for Grant-White, adding to ten, while Pasteur obtains six more cross-loadings, adding to fifteen. The increase in number of significant loadings is due to smaller standard errors for target rotation as compared to Geomin rotation. The smaller standard errors are a function of the more fully specified rotation criterion. As will be shown below, BSEM using small-variance cross-loading priors gives far simpler loading patterns by shrinking the cross-loadings towards their zero prior means.

[Table 4 about here.]

4.2 Holzinger-Swineford mental abilities example: Bayesian analysis

This section uses data from the Grant-White and Pasteur schools of the Holzinger-Swineford study to illustrate the Section 3.1 BSEM approach with informative cross-loading priors. The factor loading pattern of the four-factor model of Table 2 is used. Table 5 repeats the fit statistics for the ML CFA and EFA and adds the fit statistics for Bayesian analysis using both the original CFA model and the proposed CFA model with informative, small-variance priors for cross-loadings. The cross-loading priors use variances 0.01. All other parameters have non-informative priors. Standardized variables are analyzed and the factor variances are fixed at one in order for the priors to correspond to standardized loadings. In all analyses, the reported estimates are the median values of the parameter posterior (this is the Mplus default).

Table 5 shows that the Bayesian analysis of the CFA model with exact zero cross-loadings gives almost zero Posterior Predictive p-values in line with the ML CFA. In

contrast, for the proposed Bayesian CFA with cross-loadings model fit is acceptable in that the Posterior Predictive p-value is 0.361 for Grant-White and 0.162 for Pasteur.

As an aside, the Bayesian estimates can be used as fixed parameters in an ML analysis in order to get the likelihood-ratio test (LRT) value for the Bayes solution. They can be viewed as a descriptive measure of fit that can be compared to the ML likelihood-ratio χ^2 values. It is seen in Table 5 that the Bayesian LRT values for the CFA model are close to those of ML χ^2 values. In contrast, the Bayesian LRT values for the model with cross-loadings show a great improvement, falling in between the ML CFA and EFA χ^2 values although closer to the EFA values.

[Table 5 about here.]

The Bayes solutions for the two schools are shown in Table 6. It is interesting to compare the Bayes solution to the ML EFA solution of Table 4. The Bayes factor loadings are on the whole somewhat larger than those for ML and there are far fewer significant cross-loadings (marked with asterisks). For ML EFA, there are six significant cross-loadings for Grant-White and nine for Pasteur, whereof only three appear for both solutions. Because they appear for both schools, a researcher may be tempted to free these three cross-loadings. For Bayes, the Grant-White sample has only two cross-loadings that are significant (have a 95% credibility interval that does not cover zero) and Pasteur has one. These three cross-loadings are also significant in the ML EFA. In the Bayes analysis the significant cross-loadings are different in the two schools. Because of the lack of agreement, freeing these three cross-loadings could be capitalizing on chance and is also not necessary on behalf of model fit. Bayes clearly gives a simpler pattern than ML EFA for these data. This is achieved by shrinking the cross-loadings towards their zero prior means. The degree of shrinking that is possible while still obtaining reasonable model fit is gauged by the PPP.

It should be emphasized that cross-loadings that are found to be important in BSEM, in the sense that the 95% credibility interval does not cover zero and the cross-loading

has strong substantive backing, can be freely estimated with non-informative priors while keeping small-variance priors for other cross-loadings. This should improve the results because the small-variance prior gives a too small estimate for such a cross-loading. Monte Carlo simulations show that this gives better estimation.

The ML EFA factor correlations are smaller than the Bayesian factor correlations as is seen when comparing Table 4 with Table 6. The greater extent of cross-loadings in the EFA may contribute to the lower factor correlations in that less correlation among variables need to go through the factors. The Bayesian factor correlations are not excessively high, however, because the factors are expected to correlate to a substantial degree according to theory. Holzinger and Swineford (1939) hypothesized that the variables are all influenced by a deductive factor which in the current model is not partialled out of the four factors.

[Table 6 about here.]

The choice of cross-loading prior variance should be linked to the researcher's prior beliefs. It could be argued, however, that the choice of a variance of 0.01 resulting in 95% cross-loading limits of ± 0.20 is not substantially different from a variance of, say, 0.02 resulting in 95% limits of ± 0.28 ; see Table 1. It is therefore of interest to vary the prior variance to study sensitivity in the results. Increasing the prior variance tends to affect the Posterior Predictive p-value and also increase the variability of the estimates. At a certain point of increasing the prior variance, the model is not sufficiently identified and the MCMC algorithm tends to give non-convergence. Table 7 shows the effects of varying the prior variance for cross-loadings from 0.01 to 0.10 for the data from both the Grant-White and the Pasteur schools. The table gives the absolute value of the 95% limit of the prior distribution, the PPP, the largest cross-loading with its posterior standard deviation, and the range of the factor correlations. For Grant-White the largest cross-loading is observed for the straight test loading on the spatial factor, while for Pasteur the largest loading is observed for the figurer test loading on the spatial factor. The change in prior variance

does not affect the hypothesized pattern of major loadings and this is not reported. The range of factor correlations is included given that larger prior variance may lead to larger cross-loadings which in turn may have the effect of lowering the factor correlations due to correlations among the tests having to be channeled via the factors to a lesser degree.

Table 7 suggests that the prior variance of 0.01 may be on the low side in the sense that for both schools the PPP peaks at the prior variance 0.03 (95% cross-loading limit of 0.34). The change in prior variance, however, does not affect the results in important ways for these two data sets. For all prior variance choices the largest cross-loading for Grant-White and for Pasteur is detected, in the sense that it has its 95% credibility interval not covering zero. For Pasteur, the three highest prior variances result in the figurative cross-loading not being detected (getting a 95% credibility interval that does cover zero; entries shown as none in the table). This is because of the higher posterior variability at the higher prior variances. In hindsight, perhaps a prior variance of 0.02 or 0.03 would have been a slightly better choice, but this may not be true for other examples. On the other hand, when presenting results it is useful to give information on how a range of prior choices affects the results. Although the factor correlations show smaller values with increasing prior variance, the decrease is small and of little substantive importance. All in all, these results are reassuring in that the exact degree of informativeness does not seem critical. Also, with larger sample sizes the choice is less important in that the data provides relatively more information than the priors.

[Table 7 about here.]

In summary, BSEM provides a simpler model and a model that better fits the researcher's prior beliefs than ML. BSEM provides an approach which is a compromise between that of EFA and CFA. The ML CFA rejects the hypothesized model, presumably because it is too strict. ML EFA does not reject the model, but the model does not match the researcher's prior beliefs because it only postulates the number of factors, not where the

large and small loadings should appear. Furthermore, ML EFA provides a solution through a mechanical rotation algorithm, whereas BSEM uses priors to represent the researcher's beliefs.

4.3 Cross-loading simulations

This section discusses Monte Carlo simulations of BSEM applied to factor modeling with cross-loadings. The aim is to demonstrate that the proposed approach provides good results for data with known characteristics.

The factor loading pattern of Table 8 is considered where X denotes a major loading and x cross-loadings. The major loadings are all 0.8. The size of the three cross-loadings are varied as 0.0, 0.1, 0.2, and 0.3 in different simulations. The observed and latent variables have unit variances so the loadings are on a standardized scale. A cross-loading of 0.1 is considered to be of little importance, a cross-loading of 0.2 of some importance, and a cross-loading of 0.3 of importance (see also Cudeck & O'Dell, 1994). The correlations among the three factors are all 0.5. The factor metric is determined by fixing the first loading for each factor at 0.8. Non-informative priors are used for all parameters except for cross-loadings when those are included in the analysis. For cross-loading priors, a variance of 0.01 is chosen. Sample sizes of $n = 100$, $n = 200$ and $n = 500$ are studied.

[Table 8 about here.]

A total of 500 replications are used. The reported parameter estimate is the median in the posterior distribution for each parameter. The key result is what frequentists would refer to as the 95% coverage, that is, the proportion of the replications for which the 95% Bayesian credibility interval covers the true parameter value used to generate the data (credibility intervals obtained via percentiles of the posterior). For cross-loadings it is also of interest to study what corresponds to power in a frequentist setting. This is computed

as the proportion of the replications for which the 95% Bayesian credibility interval does not cover zero. Results are reported only for a representative set of parameters or functions of parameters: the major loading of y_2 , the cross-loading for y_6 , the variance for the first factor, and the correlation between the first and second factor.

The results are divided into three categories: Bayesian analysis using non-informative priors, model fit comparisons between ML and Bayes with non-informative priors, and Bayesian analysis with informative priors. Not all results are presented here, but some tables are instead available on the web page www.statmodel.com/examples/penn.shtml#baysem.

4.3.1 Bayes, non-informative priors

As a check of the Bayesian analysis procedure, a first analysis is carried out with non-informative priors and ignoring cross-loadings. Results can be found in web Table 1 at www.statmodel.com/examples/penn.shtml#baysem. Data are generated both with zero and non-zero cross-loadings. For zero cross-loadings, the analysis is correctly specified and close to 95% coverage is obtained for all free parameters. Posterior Predictive p-values for the model fit assessment are 0.036, 0.032, and 0.024, respectively for the three sample sizes of $n = 100$, $n = 200$, and $n = 500$, that is, reasonably close to the nominal 5% level. Bayesian analysis with non-informative priors works well in this situation.

With cross-loadings of 0.1 the effects of model misspecification show up in that the coverage is less good. Posterior Predictive p-values for the model fit assessment are 0.056, 0.080, and 0.262, respectively for the three sample sizes. This shows limited power to reject the incorrect model. On the other hand, the misspecification is deemed of little importance given the small size of the cross-loadings.

With cross-loadings of 0.2 (not shown) the Posterior Predictive p-value is 0.196 for $n = 100$, 0.474 for $n = 200$, and 0.984 for $n = 500$, showing excellent power at higher sample sizes. With cross-loadings of 0.3 the Posterior Predictive p-value is 0.544 for $n = 100$, 0.944

for $n = 200$, and 1.000 for $n = 500$, showing that the power is excellent when the cross-loading is of an important magnitude.

4.3.2 Comparing model fit for ML versus Bayes with non-informative priors

Model fit assessment comparing ML to Bayesian analysis with non-informative priors is shown in Table 9. The correctly specified model with zero cross-loadings shows an inflated ML p-value of 0.172 at $n = 100$. This small-sample bias is well-known for ML χ^2 testing (see, e.g., Scheines et al., 1999). The Posterior Predictive p-value of 0.036 based on the Bayesian analysis does not show the same problem. For the 0.1 size of the cross-loadings, which is deemed of little substantive importance, ML rejects the model 46% of the time at $n = 500$. This reflects the common notion that the ML LRT χ^2 can be oversensitive to small degrees of model misspecification. For the important degree of misspecification with cross-loading 0.3, the ML test is more powerful than Bayes, but the Bayes power is sufficient for sample sizes of at least $n = 200$.

[Table 9 about here.]

Web Table 2 shows ML model estimation results as a comparison to the Bayesian analysis with non-informative priors presented earlier. For both the correctly specified model with zero cross-loadings and for the misspecified model with cross-loadings 0.1 the ML coverage is close to that of Bayes. The mean-square-error (MSE) is also similar for Bayes and ML. Based on this, there is no reason to prefer one method over the other.

4.3.3 Bayes, informative priors

As the next step, the proposed BSEM approach of using Bayesian analysis with informative, small-variance priors for the cross-loadings is applied. The informative priors are applied to not only the three cross-loadings used to generate the data, but to all cross-loadings

to reflect a real-data analysis situation. The prior variance is chosen as 0.01. All other parameters are given non-informative priors. Table 10 shows good coverage and for the top part of the table corresponding to the correctly specified analysis with zero cross-loadings the coverage remains largely the same as in web Table 1. For cross-loadings of 0.1, however, the bottom part of the table shows that coverage has improved by the introduction of informative, small-variance priors for the cross-loadings. The coverage is acceptable also for the cross-loading. The power to detect the cross-loading is, however, small at this low cross-loading magnitude, 0.038, 0.098 and 0.176, respectively for the three sample sizes. The Posterior Predictive p-value is on the low side in all four cases.

It is interesting to compare the coverage results for the four parameters in the case of cross-loadings 0.1 given in the Table 10 for BSEM with the results from the ML approach in web Table 2. ML is outperformed by BSEM by its use of informative priors.

[Table 10 about here.]

Table 11 shows the results of BSEM where data have been generated with larger cross-loadings of 0.2 and 0.3. Here the coverage is also good with the exception of the cross-loadings. For the cross-loadings, however, the focus is on power as shown in the last columns. For a cross-loading of 0.2 a sample size of $n = 500$ is needed to obtain power above 0.8. For a cross-loading of 0.3 a sample size of $n = 200$ is sufficient to obtain power above 0.8. This shows that the approach of using informative, small-variance priors for cross-loadings leads to a successful way to modify the model, allowing free estimation of the indicated cross-loadings. When freed and estimated using non-informative priors, these cross-loadings are well estimated.

The point estimates indicate that the key parameter of factor correlation is overestimated. Note, however, that given the power to detect cross-loadings, estimating the cross-loadings freely results in good point estimates for factor correlations.

For the case of 0.3 cross-loadings in Table 11, the alternative prior variance of 0.02 was also tried, yielding improved results. The average estimates for the four entries was 0.8060, 0.2117, 1.0965, and 0.5620, while the coverage was 0.956, 0.886, 0.954, and 0.982.

In summary, the cross-loading simulation study shows that the Bayesian analysis performs well. It also shows that in terms of parameter coverage and for the case of small cross-loadings ML is inferior to BSEM. In terms of model testing, BSEM avoids the small-sample inflation of the ML χ^2 and also avoids the ML χ^2 sensitivity to rejecting a model with an ignorable degree of misspecification.

[Table 11 about here.]

5 Study 2: Residual correlations in CFA

5.1 British Household Panel Study (BHPS) big-five personality example: ML analysis

A second example uses data from the British Household Panel Study (BHPS) of 2005 and 2006. A 15-item, five-factor instrument uses three items to measure each of the "big-five" personality factors: agreeableness, conscientiousness, extraversion, neuroticism, and openness. Each item uses the question "I see myself as someone who . . ." followed by a statement. There are seven response categories ranging from 1 "does not apply" to 7 "applies perfectly". A total of 14,021 subjects are included. The big-five factors are expected a priori to have low correlations and are known to be related to gender and age; see, e.g. Marsh, Nagengast, and Morin (2010). For simplicity, the current analyses hold age constant by considering the subgroup of ages 50-55. This produces a sample of $n = 691$ females and $n = 589$ males.

The item wording and hypothesized loading pattern are shown in Table 12. For all

factors except openness, there are two positively-worded items and one negatively-worded item. Marsh, Nagengast, and Morin (2010) suggest that the four negatively-worded items may a priori have correlated residuals (correlated uniquenesses) when applying factor analysis.

[Table 12 about here.]

Using maximum-likelihood estimation, model fit using confirmatory factor analysis (CFA), CFA with correlated uniquenesses (CU) among the negatively-worded items, and exploratory factor analysis (EFA) is reported in Table 13. It is seen that the fit is not acceptable for either of the two CFA models as judged by χ^2 or the two model fit indices. The EFA model is also rejected by χ^2 and only marginally acceptable for males when judged by RMSEA or CFI.

[Table 13 about here.]

An interesting finding is that the EFA solutions for females and males do not fully capture the hypothesized factors. This is the case using the Geomin rotation as well as using Quartimin and Varimax. The Geomin rotation for each gender is shown in Table 14. The bolded entries are loadings that are the largest for the item. Comparing to Table 12 it is seen that only the factors extraversion, neuroticism, and openness are found, not the agreeableness and conscientiousness factors. A possible reason for this is the existence of correlated residuals among the items. As the CFA with CUs model showed, however, allowing residual correlations among the reverse-coded items is not sufficient. It is likely that in addition to the big-five factors the personality instrument measures many minor factors.

The question arises how correlated residuals can be accounted for while maintaining the hypothesized factor loading pattern. A model with all residual correlations freely estimated is not identified. The proposed BSEM solution to this problem is presented next.

[Table 14 about here.]

5.2 British Household Panel Study (BHPS) big-five personality example: Bayesian analysis

This section uses the big-five personality data in the British Household Panel Study (BHPS) to illustrate the BSEM approach of Section 3.2 with an informative prior for the residual covariance matrix. Because of its relative simplicity, Method 2 is used. The simulation studies to be presented also favor Method 2. An inverse-Wishart prior $IW(\mathbf{I}, df)$ with $df = p + 6 = 21$ is used, corresponding to prior means and standard deviations for residual covariances of zero and 0.1, respectively (see Appendix, (26)). Standardized variables are analyzed. Because of high auto-correlation among the MCMC iterations, only every 10th iterations is used with a total of 100,000 iterations to describe the posterior distribution. Informative cross-loading priors are also used with prior distributions $N(0, 0.01)$.

The Posterior Predictive p-values are 0.534 and 0.518, respectively for females and males indicating a good match between the model and the data. For the two samples 17 and 37 residual covariances, respectively, were found significant in the sense of the 95% Bayesian credibility interval not covering zero. The average absolute residual correlation (range) is 0.329 (−0.462 to 0.647) for females and 0.285 (−0.484 to 0.590) for males. For both genders only one residual correlation exceeds 0.5 in absolute value. This suggests that many small residual correlations need to be included in the factor model, as was expected. The fact that these residual correlations are left out in the ML analyses may contribute to the poor ML fit and the poor ML EFA loading pattern recovery.

Table 15 gives the results for the female and male samples. Standardized loadings are presented so that the results can be compared to the ML EFA of Table 14. The hypothesized major loadings are all recovered at substantial values with no significant cross-loadings. The factor correlations are all small to moderate as was expected. The

extraversion, neuroticism, and openness factors that were recovered in the ML EFA of Table 14 have lower correlations in the Bayesian solution than in the ML EFA.

In summary, BSEM provides a solution that better fits the researcher's prior beliefs than ML. The ML CFA rejects the hypothesized model, presumably because it is too strict in terms of requiring exactly zero residual covariances. ML EFA does not recover the researcher's hypothesized big-five factor pattern.

[Table 15 about here.]

5.3 Residual correlations simulations

This section discusses Monte Carlo simulations of the BSEM approach to factor modeling with residual correlations. A factor model with 10 variables and two factors is used, where the first five variables load on only the first factor and the second five variable load only on the second factor. The loadings are all 0.8, the factor variances are 1, and the factor correlation is 0.5. The residual variances are 0.36 so that observed variables all have variances 1. Two residual covariances (correlations) are included, one for the first and sixth variables and one for the second and seventh variables. In this way, ignoring the residual covariances in the modeling tends to inflate the factor correlation. An example would be an instrument administered at two time points, where some variables have residuals that are correlated over time. Residual correlations of 0.0, 0.1, and 0.3 are considered together with sample sizes $n = 200$ and $n = 500$. A total of 500 replications are used and the results presented in the format used for the cross-loading simulations. The simulations present results for all three methods discussed in Section 3.2. For Method 1, both a more informative prior with $df = 30$ and a less informative prior with $df = 14$ ($= p + 4$) is studied. For Method 2, $df = 30$ is used. For Method 3, the normal prior variance is set at 0.001. Tables with results are available at the web page www.statmodel.com/examples/penn.shtml#baysem.

5.3.1 Comparing ML to Bayes

As a first step, model testing using ML and Bayes with non-informative priors is compared for both correctly and misspecified models. With residual correlations of 0.0, both ML and Bayes give acceptable rejection rates with the correctly specified model (results presented in web Table 3). Both ML and Bayes reject the model with ignorable residual correlations of 0.1, although ML is more sensitive to this misspecification. For the larger misspecification with residual correlations of 0.3, both ML and Bayes show sufficient power to reject the model at both $n = 200$ and $n = 500$.

With BSEM Method 1 and residual correlations 0.1, good coverage is found for all parameters except the residual covariance (web Table 4). There is sufficient power to detect the residual covariance of 0.1 at $n = 500$. There is no important difference between using $df = 30$ and $df = 14$, except that the point estimate and the power for the residual covariance is slightly better for the less informative prior with $df = 14$.

With BSEM Method 1 and residual correlations 0.3, acceptable coverage is found when using the less informative prior with $df = 14$, except for the residual correlation (web Table 5). The power to detect the residual correlations is, however, excellent in all cases. For the more informative prior with $df = 30$, the coverage is less good. The key parameter of the factor correlation shows an important overestimation, which is also seen with $df = 14$.

The simulation results for BSEM Methods 2 and 3 using residual correlations of 0.3 are studied next (web Table 6). The 5% reject proportion for the Posterior Predictive p-value is 0 in all cases. For Method 2 the results are very good except for the residual covariance being underestimated and having poor coverage. The power to detect it is, however, excellent. The factor correlation is also somewhat underestimated. Method 2 performs considerably better than Method 1. The Method 3 results are somewhat worse than those of Method 1 for $df = 14$, with poorer performance for the residual covariance and the factor correlation. The power to detect the residual covariance is, however, excellent also

for Method 3. It should be noted that Method 3 is the only one of the three methods that can let such a residual covariance be freely estimated, that is, using a non-informative prior. Using a less informative Method 3 prior with a larger variance of 0.01 did not alter the results very much. In summary, Method 2 performs the best of the three methods and Method 3 the worst for this simulation setting.

Method 3 works well when the two residual covariances are freely estimated, that is, using non-informative priors (web Table 7). The remaining residual covariances are using the same informative priors as before. Results are shown for $n = 200$ and $n = 500$.

6 Conclusions

This paper proposes a new approach to factor analysis and structural equation modeling using Bayesian analysis. The new approach represents hypotheses in a new way, replacing parameter specifications of exact zeros with approximate zeros based on informative, small-variance priors. It is argued that this produces an analysis that better reflects substantive theories. The proposed Bayesian approach with informative priors for hypothesized parameter restrictions, labeled BSEM, is particularly beneficial in applications where if those parameters are added to a conventional model, a non-identified model is obtained using ML. The extra model parameters can be viewed as nuisance parameters that based on substantive theory and previous studies are hypothesized to be close to zero although perhaps not exactly zero. This approach is useful for measurement aspects of latent variable modeling such as with CFA and the measurement part of SEM. Two application areas are studied, cross-loadings in CFA and residual correlations in CFA. The approach encompasses three elements: Model testing, model estimation, and model modification. The first two are evaluated by Monte Carlo simulation studies, whereas the third warrants further studies. The Monte Carlo and real-data results can be summarized as follows.

6.1 Summary of findings

Model testing uses a Posterior Predictive Probability (PPP) approach that has not previously been investigated this extensively. It is found that PPP works well both for models with only non-informative priors and for the proposed BSEM approach where some parameters have informative priors. PPP is found to perform better than the ML likelihood-ratio χ^2 test at small sample sizes where ML typically inflates χ^2 , and is found to be less sensitive than ML to ignorable deviations from the correct model. PPP is found to have sufficient power to detect important model misspecifications.

Bayesian model estimation is shown to perform well with both non-informative and informative priors. Using BSEM with both ignorable and non-ignorable degrees of model misspecification, key parameters are well estimated in terms of their coverage. BSEM outperforms ML estimation with misspecified models.

BSEM also provides a counterpart to ML-based model modification. ML modification indices inform about model improvement when a single parameter is freed and can lead to a long series of modifications. In contrast, BSEM informs about model modification when all parameters are freed and does so in a single step. The simulations show sufficient power to detect model misspecification in terms of 95% Bayesian credibility intervals not covering zero.

An example for each of the two application areas show the promise of BSEM. For the Holzinger-Swineford example a well-fitting factor model is found that is superior to ML-based models. Instead of choosing between an ill-fitting ML CFA model and a well-fitting but unnecessarily weakly specified ML EFA model, BSEM maintains the spirit of CFA while allowing small cross-loadings. A comparison is also made with target rotation (Browne, 2001; Asparouhov & Muthén, 2009). Target rotation is similar to BSEM in that it replaces mechanical rotation with rotation guided by the researcher’s judgement, in this case using zero targets for cross-loadings. It is different from BSEM by not allowing user-

specified stringency of closeness to zero by varying the prior variance, replacing that with least-square fitting. For the Holzinger-Swineford example, applying target rotation with zero targets for all cross-loadings gives results similar to EFA using Geomin or Quartimin rotation, except yielding more significant cross-loadings. BSEM using small-variance cross-loading priors gives far simpler loading patterns, shrinking the cross-loadings toward the prior mean. A check of the degree of shrinkage that matches the data is provided by the BSEM PPP approach. Table 7 shows that for these data the prior variance choice does not have important impact on the results.

For the big-five personality example a well-fitting factor model is found that recovers the hypothesized factor loading pattern by allowing for a large number of small residual correlations. In contrast, ML CFA is ill-fitting even when allowing for a priori residual correlations, and ML EFA does not recover the hypothesized factor loading pattern.

Applying BSEM is easy and fast for analyses of cross-loadings. Analysis with residual covariances leads to heavier computations due to slow MCMC convergence. A further benefit of the Bayesian analysis is that estimation works well also for models that are large relative to the sample size (see also Asparouhov & Muthén, 2010b).

6.2 Related approaches

BSEM with its adjoining PPP model test is similar in spirit to the frequentist conceptualization of "close fit" (Browne & Cudeck, 1993). ML model testing of close fit rather than conventional exact χ^2 fit is expressed by the root mean square error of approximation (RMSEA) fit index. In assessing differences between models, McCallum et al. (2006) also argue against exact fit as being of limited empirical interest given that it is never true in practice. RMSEA uses an overall approximate fit level deemed sufficient. In contrast, BSEM allows informative priors to reflect notions of closeness for each parameter.

Press (2003; chapter 15) discusses a Bayesian factor analysis approach that has some

similarities to the one proposed in this paper. The MCMC algorithm is not used but instead estimates are obtained as expected values in the posterior distributions. Press (2003) specifies a prior for the loading matrix with a mean that uses a specific "target" pattern of large and zero loadings. All loadings have the same prior variances. In the example (Press, 2003; pp. 368-372) the variances are chosen to give weakly informative priors. In contrast, the current approach has zero prior means for all loadings, with small prior variances for non-target loadings and large prior variances for target loadings so that target loadings are solely determined by the data. In this sense, the Press (2003) approach is closer to EFA and the current approach is closer to CFA.

In BSEM the ability to free all loadings in a measurement model can be viewed as the ability to form an EFA with the rotation guided by the priors. BSEM is, however, more general than EFA and essentially has the flexibility of ESEM (Asparouhov & Muthén, 2009) because it can accommodate correlated residuals in an EFA model, it can accommodate covariates in an EFA model, and it can accommodate an EFA model as part of a larger model. In terms of the measurement model, while ESEM is exploratory in nature, BSEM has more of a confirmatory flavor. BSEM also generalizes ESEM in the following way. In ESEM the optimal rotation is determined based only on the unrotated loadings as in EFA, that is, the optimal rotation does not consider residual covariances or covariate direct effects in the optimal rotations. In contrast, in BSEM the optimal rotation is determined by all parts of the model.

6.3 BSEM extensions

BSEM can be extended to include equality constraints. A typical SEM example is multiple-group analysis with measurement invariance. It is common to find small deviations from exact invariance that cause rejection by the ML LRT. Group differences in measurement intercept vectors and loading matrices can be given zero-mean, small-variance priors. The

special case of intercept non-invariance can be handled by letting the grouping variables be covariates that influence the factors, also referred to as MIMIC modeling (see, e.g., Muthén, 1989). Here, non-invariance is defined as direct effects from covariates to the factor indicators. With ML, including all direct effects results in a non-identified model, whereas BSEM solves the problem using zero-mean, small-variance priors for the direct effects. For a study of this extension, see

www.statmodel.com/examples/penn.shtml#baysem.

The BSEM approach is not limited to measurement modeling, but is also applicable to restrictions on structural coefficients in SEM. For example, in a mediational model, a hypothesized absence of a substantively important direct effect may be specified as an informative normal prior with zero mean and variance that corresponds to a negligible effect. While the two application areas studied in this paper show the particular advantage of BSEM when ML estimation produces a non-identified model, this non-identification aspect is not a requirement for BSEM.

6.4 Reflections on analysis strategies

This paper discussed several factor analysis alternatives: Exploratory factor analysis using both mechanical and target rotation; confirmatory (ML and Bayes) factor analysis; BSEM with informative cross-loading priors; and BSEM with informative residual covariances. It is worthwhile to consider the different choices made with these different types of analyses in order to gain further understanding of the epistemological implications of BSEM.

One key aspect of factor analysis is the resulting factor correlations. The analysis of the Holzinger-Swineford data provides an illustration of the different factor correlation findings obtained by the different analysis alternatives. In EFA using oblique rotation, the

non-zero correlations among the factors typically reduce the size of cross-loadings relative to orthogonal rotation because correlations among the factor indicators on different factors can be channeled via the factors. From the point of view of BSEM with informative cross-loading priors the factor correlations from EFA with oblique rotation may be too low because too many non-zero cross-loadings are allowed. BSEM shrinks the cross-loadings toward their prior means of zero and the BSEM PPP gauges if a certain degree of shrinking, corresponding to a certain prior variance, is compatible with the data. In the Holzinger-Swineford data the EFA factor correlations were lower than the BSEM factor correlations and this is expected to generally be the case. EFA with target rotation did not change this picture.

ML CFA with correlated factors fixes many cross-loadings to zero so that the rotation of EFA is avoided. Because of the many cross-loadings fixed at zero, CFA tends to require higher factor correlations than EFA with oblique rotation in order to represent the correlations among the factor indicators (see also Asparouhov & Muthén, 2009; Marsh et al., 2009, 2010). From the point of view of BSEM with informative cross-loading priors, these CFA factor correlations are too high. This is because BSEM postulates cross-loadings that are not exactly zero, which in turn leads to lower BSEM factor correlations. In this sense, factor correlations from BSEM with informative priors for cross-loadings are expected to be a compromise between EFA and CFA factor correlations. This is the case for the Holzinger-Swineford data.

In the case of informative priors for residual covariances, BSEM is expected to result in smaller factor correlations than CFA with zero residual covariances given that less of the correlation among factor indicators need to be channeled via the factors.

Given these observations, a possible strategy is to use EFA with mechanical rotation in early pilot studies of a measurement instrument until a body of knowledge about the factor indicators and the factors has been built up. Although EFA was here carried out by ML,

it could also be carried out by Bayesian analysis with non-informative priors. A switch can then be made from EFA to BSEM with informative priors, where the informative priors can be chosen with smaller and smaller variances. In this sense, the Bayesian approach provides a continuum of analyses to be carried out in a series of studies, choosing priors to reflect increasing knowledge about the measurement situation. Here, ML CFA is the frequentist counterpart to the far end of the continuum. The Bayesian approach avoids the big increase in model parsimony going from an ML EFA to an ML CFA, which often leads to an ill-fitting CFA model. Similarly, it avoids the big jump in going directly to an ML CFA without preceding EFA steps as is currently often advocated, also typically leading to an ill-fitting CFA model.

As a final point, a devil's advocate may argue that the BSEM approach adds "junk parameters" to permit model fit. A first response in the context of cross-loadings is that EFA potentially adds more such parameters and unlike BSEM does not test statistically if they are needed. A more important response is that unlike CFA, BSEM allows the researcher to specify the degree of precision with which he or she wants to portray prior beliefs. For CFA the only choice is what corresponds to a prior variance of exactly zero, whereas with BSEM an exact zero is not required. For models and data where the choice of prior variance makes a difference to the interpretation of the results, this informs the researcher that the data does not carry enough information on the model. A more comforting situation is illustrated in Table 7, showing ignorable dependence on the prior variance choice.

6.5 Caveats

Several warnings are important for using BSEM. This is especially the case regarding the use of BSEM with informative priors and residual covariances. First, it may be difficult to balance the need for small residual covariances against small cross-loadings in that both aid

in representing correlations among factor indicators. Second, allowing for small residual covariances may obscure the need to add minor but still important factors. Third, medium-sized residual covariances may obscure that the postulated factor pattern is misspecified.

Furthermore, this paper presents only a beginning of the study of BSEM. Much more experience is needed. For example, the PPP approach to model checking needs further study. How much are the p-values influenced by the number of variables and other model features? Preliminary investigations of moderate departures from the assumed multivariate normality of the observed variables does not seem to have a critical impact, but this needs to be studied further. Another question is to which extent the maximum PPP value should guide which prior the results should be reported for. Although priors should be decided on before the data are analyzed, often a range of priors are equally motivated. Also, it would be worthwhile to offer several posterior predictive tests, extending PPP beyond merely using the likelihood-ratio test statistic for the overall model and also focus on a particularly important part of the model implications. Furthermore, the idea of the BSEM-derived counterpart to modification indices needs to be evaluated. It is of interest to see if this is more likely to lead to the correct model when the initial model needs several modifications. More needs to be learned about the performance of BSEM parameter posterior estimation using different informative priors for different types of models, sample sizes, and variable distributions. Hopefully, this paper will stimulate such further research.

References

- [1] Asparouhov, T. & Muthén, B. (2009). Exploratory structural equation modeling. *Structural Equation Modeling*, 16, 397-438.
- [2] Asparouhov, T. & Muthén, B. (2010a). Bayesian analysis using Mplus: Technical implementation. Technical appendix. Los Angeles: Muthén & Muthén. www.statmodel.com
- [3] Asparouhov, T. & Muthén, B. (2010b). Bayesian analysis of latent variable models using Mplus. Technical report. Los Angeles: Muthén & Muthén. www.statmodel.com
- [4] Barnard, J., McCulloch, R. E., & Meng, X. L. (2000). Modeling covariance matrices in terms of standard deviations and correlations, with application to shrinkage. *Statistica Sinica* 10, 1281-1311.
- [5] Browne, M. W. (2001). An overview of analytic rotation in exploratory factor analysis. *Multivariate Behavioral Research*, 36, 111-150.
- [6] Browne, M. W. & Cudeck, R. (1993). Alternative ways of assessing model fit. In: Bollen, K. A. & Long, J. S. (Eds.) *Testing Structural Equation Models*. pp. 136-162. Beverly Hills, CA: Sage.
- [7] Browne, W.J. & Draper, D. (2006). A comparison of Bayesian and likelihood-based methods for fitting multilevel models.
- [8] Casella, G. & George, E.I. (1992). Explaining the Gibbs sampler. *The American Statistician*, 46, 167-174.
- [9] Chib, S. & Greenberg, E. (1998). Bayesian analysis of multivariate probit models. *Biometrika* 85, 347-361.
- [10] Cudeck, R. & O'Dell, L.L. (1994). Applications of standard error estimates in unrestricted factor analysis: Significance tests for factor loadings and correlations. *Psychological Bulletin*, 115, 475-487.

- [11] Dominicus, A., Ripatti, S., Pedersen, N.L., & Palmgren, J. (2008). A random change point model for assessing the variability in repeated measures of cognitive function. *Statistics in Medicine*, 27, 5786-5798.
- [12] Enders, C.K. (2010). *Applied missing data analysis*. New York: Guilford Press.
- [13] Gelfand, A.E., Hills, S.E., Racine-Poon, A., & Smith, A.F.M. (1990). Illustration of Bayesian inference in normal data models using Gibbs sampling. *Journal of the American Statistical Association*, 85, 972-985.
- [14] Gelman, A. & Rubin, D.B. (1992). Inference from iterative simulation using multiple sequences. *Statistical Science*, 7, 457-511.
- [15] Gelman, A., Meng, X.L., Stern, H.S. & Rubin, D.B. (1996). Posterior predictive assessment of model fitness via realized discrepancies (with discussion). *Statistica Sinica*, 6, 733-807.
- [16] Gelman, A., Carlin, J.B., Stern, H.S. & Rubin, D.B. (2004). *Bayesian data analysis*. Second edition. Boca Raton: Chapman & Hall.
- [17] Green, P. (1996). MCMC in image analysis. In Gilks, W.R., Richardson, S., & Spiegelhalter, D.J. (eds.), *Markov chain Monte Carlo in Practice*. London: Chapman & Hall.
- [18] Gustafsson, J-E. (2002). Measurement from a hierarchical point of view. In H. I. Braun, D. N. Jackson, & D. E. Wiley (Eds.) *The role of constructs in psychological and educational measurement* (pp. 73-95). London: Lawrence Erlbaum Associates, Publishers.
- [19] Harman, H.H. (1976). *Modern factor analysis*. Third edition. Chicago: The University of Chicago Press.
- [20] Hayashi, K. & Marcoulides, G. A. (2006). Examining Identification Issues in Factor Analysis. *Structural Equation Modeling*, 13, 631-645.

- [21] Hjort, N.L., Dahl, F.A., & Steinbakk, G.H. (2006). Post-processing posterior predictive p values. *Journal of the American Statistical Association*, 101, 1157-1174.
- [22] Holzinger, K.J. & Swineford, F. (1939). A study in factor analysis: The stability of a bi-factor solution. Supplementary Educational Monographs. Chicago.: The University of Chicago Press.
- [23] Kruschke, J.K. (2010). Doing Bayesian data analysis: A tutorial with R and BUGS. New York: Elsevier.
- [24] Lee, S.Y. (2007). Structural equation modelng. A Bayesian approach. Chichester: John Wiley & Sons.
- [25] Lynch, S.M. (2010). Introduction to applied Bayesian statistics and estimation for social scientists. New York: Springer.
- [26] MacKinnon, D.P. (2008). Introduction to statistical mediation analysis. New York: Erlbaum.
- [27] Marsh, H.W., Muthén, B., Asparouhov, A., Ldtke, O., Robitzsch, A., Morin, A.J.S., & Trautwein, U. (2009). Exploratory structural equation modeling, Integrating CFA and EFA: Application to Students' Evaluations of University Teaching. *Structural Equation Modeling*, 16, 439-476.
- [28] Marsh, H.W., Nagengast, B., & Morin, A.J.S. (2010). Measurement invariance of big-five factors over the lifespan: ESEM tests of gender, age, plasticity, maturity and la dolce vita effects.
- [29] McCallum, R.C., Roznowski, M., & Necowitz, L.B. (1992). Model modifications in covariance structure analysis: The problem of capitalizing on chance. *Psychological Bulletin*, 111, 490-504.

- [30] McCallum, R. C., Browne, M.W., & Cai, L. (2006). Testing differences between nested covariance structure models: Power analysis and null hypotheses. *Psychological Methods*, 11, 19-35.
- [31] Muthén, B. (1989). Latent variable modeling in heterogeneous populations. *Psychometrika*, 54:4, 557-585.
- [32] Muthén, B. & Muthén, L. (1998-2010). *Mplus User's Guide*. Sixth Edition. Los Angeles, CA: Muthén & Muthén.
- [33] Press, S.J. (2003). *Subjective and objective Bayesian statistics. Principles, models, and applications*. Second edition. New York: John Wiley & Sons.
- [34] Rupp, A.A., Dey, D.K., & Zumbo, B.D. (2004). To Bayes or not to Bayes, from whether to when: Applications of Bayesian methodology to modeling. *Structural Equation Modeling*, 11, 424-451.
- [35] Schafer, J.L. (1997). *Analysis of incomplete multivariate data*. London: Chapman & Hall.
- [36] Scheines, R., Hoijtink, H., & Boomsma, A. (1999). Bayesian estimation and testing of structural equation models. *Psychometrika*, 64, 37-52.
- [37] Yuan, Y. & MacKinnon, D.P. (2009). Bayesian mediation analysis. *Psychological Methods*, 14, 301-322.

7 Appendix

7.1 Obtaining the posterior distribution

Bayesian estimation uses Markov Chain Monte Carlo (MCMC) algorithms. The idea behind MCMC is that the conditional distribution of one set of parameters given other sets can be used to make random draws of parameter values, ultimately resulting in an approximation of the joint distribution of all the parameters. For a technical discussion, see, e.g., Gelman et al. (2004). For the technical implementation in Mplus, see Asparohov and Muthén (2010b). Denote by $\boldsymbol{\pi}_i$ a vector of unknowns consisting of parameters, latent variables, and missing observations at iteration i . The vector is divided into several sets, $\boldsymbol{\pi} = (\boldsymbol{\pi}_{1i}, \boldsymbol{\pi}_{2i}, \dots, \boldsymbol{\pi}_{Si})'$. For example, in an application without latent variables and missing data, the parameters may be divided into means, intercepts, and slopes in one set and variance and residual variances in another set. Normal priors are commonly used for the first set while inverse-Gamma and inverse-Wishart priors are commonly used for the second set. The conditional distribution for the first set is normal and for the second set inverse-Gamma or inverse-Wishart.

The MCMC sequence of random draws can be described as follows. Using a set of parameter starting values, new $\boldsymbol{\pi}$ values are obtained by the following steps over $i = 1, 2, \dots, n$ iterations, in each step drawing from a conditional posterior parameter distribution:

$$\text{Step 1 : } \boldsymbol{\pi}_{1,i} | \boldsymbol{\pi}_{2,i-1}, \dots, \boldsymbol{\pi}_{S,i-1}, \text{ data, priors} \quad (10)$$

$$\text{Step 2 : } \boldsymbol{\pi}_{2,i} | \boldsymbol{\pi}_{1,i}, \boldsymbol{\pi}_{3,i-1}, \dots, \boldsymbol{\pi}_{S,i-1}, \text{ data, priors} \quad (11)$$

$$\dots \quad (12)$$

$$\text{Step } S : \boldsymbol{\pi}_{S,i} | \boldsymbol{\pi}_{1,i}, \dots, \boldsymbol{\pi}_{S-1,i-1}, \text{ data, priors.} \quad (13)$$

For the step 1 iteration 1, the parameter values for iteration $i - 1 = 0$ are starting values. Step 1 produces values for the parameters of π_1 . In the step 2 iteration 1 those values and the starting values for the other parameters produce values for the parameters of π_2 , and so on up to the step S iteration 1. Iterations $2, \dots, n$ go through the same steps in the same fashion. Typically, several MCMC chains are used, starting from different starting values and using different random seeds when making the random draws. The chains form independent sequences of iterations and gives an opportunity to monitor convergence.

In certain cases it is not possible to draw from the above conditional posterior distributions because they do not exist in explicit form. In such cases the Metropolis-Hastings algorithm (Gelman et al., 2004) is used instead. Suppose that in Step 1, $\pi_{1,i}$ can not be explicitly drawn. Then $\pi_{1,i}^*$ is drawn from a jumping distribution J and this draw is accepted as $\pi_{1,i}$ with probability

$$R = \frac{J(\pi_{1,i-1})}{J(\pi_{1,i}^*)} \frac{P(\pi_{1,i}^*|*)}{P(\pi_{1,i-1}|*)}.$$

Otherwise $\pi_{1,i-1}$ is used as the next draw $\pi_{1,i}$.

7.2 Assessing convergence

In the analyses of this paper convergence is investigated in the following way. Consider n iterations in m chains, where π_{ij} is the value of parameter π in iteration i of chain j .

Define the within- and between-chain variation as

$$\bar{\pi}_{.j} = \frac{1}{n} \sum_{i=1}^n \pi_{ij}, \quad (14)$$

$$\bar{\pi}_{..} = \frac{1}{m} \sum_{j=1}^m \bar{\pi}_{.j}, \quad (15)$$

$$W = \frac{1}{m} \sum_{j=1}^m \frac{1}{n} \sum_{i=1}^n (\pi_{ij} - \bar{\pi}_{.j})^2, \quad (16)$$

$$B = \frac{1}{m-1} \sum_{j=1}^m (\bar{\pi}_{.j} - \bar{\pi}_{..})^2. \quad (17)$$

Convergence is determined using the Gelman-Rubin convergence diagnostic (Gelman & Rubin (1992); Gelman et al., 2004). This considers the potential scale reduction factor (PSR),

$$PSR = \sqrt{\frac{W+B}{W}}, \quad (18)$$

where a PSR value not much larger than 1 is considered evidence of convergence. Gelman et al. (2004) suggests 1.1, or smaller values for all parameters. This means that convergence is achieved when the between-chain variation is small relative to the within-chain variation. Gelman et al. (2004) use a slightly different definition of their potential scale reduction $\hat{R}$, but the difference relative to (18) is a negligible factor of $n/(n-1)$. It may be the case, however, that PSR convergence observed after n iterations may be negated when using more iterations. Because of this, a longer chain should be run to check that PSR values are close to 1 in a long sequence of iterations.

7.3 Priors

The density of the inverse-Wishart distribution $IW(\mathbf{S}, d)$ with d degrees of freedom is given by

$$\frac{|\mathbf{S}|^{d/2} |X|^{-(d+p+1)/2} \text{Exp}(-\text{Tr}(\mathbf{S}X^{-1})/2)}{2^{dp/2} \Gamma_p(d/2)}, \quad (19)$$

where Γ_p is the multivariate gamma function and the argument X of the density is a positive density function. To use an informative prior with a certain expected value one can use the fact that the mean of the distribution is

$$\frac{\mathbf{S}}{d - p - 1}. \quad (20)$$

The mean exist and is finite only if $d > p + 1$. If $d \leq p + 1$ then one can use the fact that the mode of the distribution is

$$\frac{\mathbf{S}}{d + p + 1}. \quad (21)$$

The variance, i.e., the level of informativeness is controlled exclusively by the parameter d . The larger the value of d the more informative the prior is.

To evaluate the informativeness of the prior one can consider the marginal distribution of the diagonal elements. The marginal distribution of the j -th diagonal entry is

$$IG((d - p + 1)/2, \mathbf{S}_{jj}/2). \quad (22)$$

Thus the marginal mean is

$$\frac{\mathbf{S}_{jj}}{d - p - 1} \quad (23)$$

if $d > p + 1$ and the marginal variance is

$$\frac{2\mathbf{S}_{jj}^2}{(d - p - 1)^2(d - p - 3)} \quad (24)$$

if $d > p + 3$. To use an informative prior with a certain variance one can multiply the desired expected value by $(d - p - 1)$ to get $\mathbf{S}$.

The marginal distribution of the off-diagonal elements can not be expressed in closed form, but the marginal mean for the (i, j) off-diagonal element is

$$\frac{\mathbf{S}_{ij}}{d - p - 1} \quad (25)$$

if $d > p + 1$ and the marginal variance is

$$\frac{(d - p + 1)\mathbf{S}_{ij}^2 + (d - p - 1)\mathbf{S}_{ii}\mathbf{S}_{jj}}{(d - p)(d - p - 1)^2(d - p - 3)} \quad (26)$$

if $d > p + 3$. As an example, using an identity matrix $\mathbf{S} = \mathbf{I}$ and $d = p + 6$ for $IW(\mathbf{S}, d)$ gives mean zero and variance = 0.0111 (standard deviation = 0.1054).

It is clear that stating the level of informativeness using inverse-Wishart priors is rigid as the informativeness of one parameter in the matrix determines the informativeness of all other parameters. A special case is of particular interest. Setting the prior to $IW(\mathbf{D}, p+1)$ where $\mathbf{D}$ is a diagonal matrix, the marginal distribution for all correlations is uniform on the interval $(-1, 1)$ while the marginal distributions of the variance is $IG(1, d_{jj}/2)$. The values of the diagonal elements d_{jj} can be set to match the mode of the desired prior with the mode of $IG(1, d_{jj}/2)$ which is $d_{jj}/4$. Note however that the mean can not be used for this purpose since the mean of $IG(1, d_{jj}/2)$ is infinity. Only the mode is defined for this distribution. In this case the marginal distribution of the diagonal elements has infinite mean and variance. The marginal for the covariance elements has mean zero by symmetry but also has an infinite variance. The marginal mean for the correlation parameter is zero and the marginal variance for the correlation parameter is $1/3$.

More generally setting the prior to $IW(\mathbf{D}, d)$ where $\mathbf{D}$ is a diagonal matrix, the marginal distribution for all correlations is the beta distribution $B((d-p+1)/2, (d-p+1)/2)$

on the interval $(-1,1)$, if $d \geq p$ with mean 0 and variance

$$\frac{1}{d - p + 2}. \quad (27)$$

Note also that the posterior distribution in the MCMC generation for the variance covariance parameter with prior $IW(\mathbf{S}, d)$ is a weighted average of $\mathbf{S}/d$ and the sample variance where the weights are $d/(n + d)$ and $n/(n + d)$ respectively, where n is the sample size.. Thus one can interpret the degrees of freedom parameter d as the number of observations added to the sample with the prior variance covariance matrix. Naturally as the sample size increases the weight $d/(n + d)$ will converge to 0 and the effect of the prior matrix $\mathbf{S}$ will diminish. To maintain the same effect of the prior on the estimation for larger sample sizes the degrees of freedom parameter should be chosen proportionally larger.

More information on the inverse-Wishart distribution and the marginal distributions of all the entries in the matrix can be found in Barnard et al. (2000).

List of Figures

1	Prior, likelihood, and posterior for a parameter	50
2	Informative prior for a factor loading parameter	51

Figure 1: Prior, likelihood, and posterior for a parameter

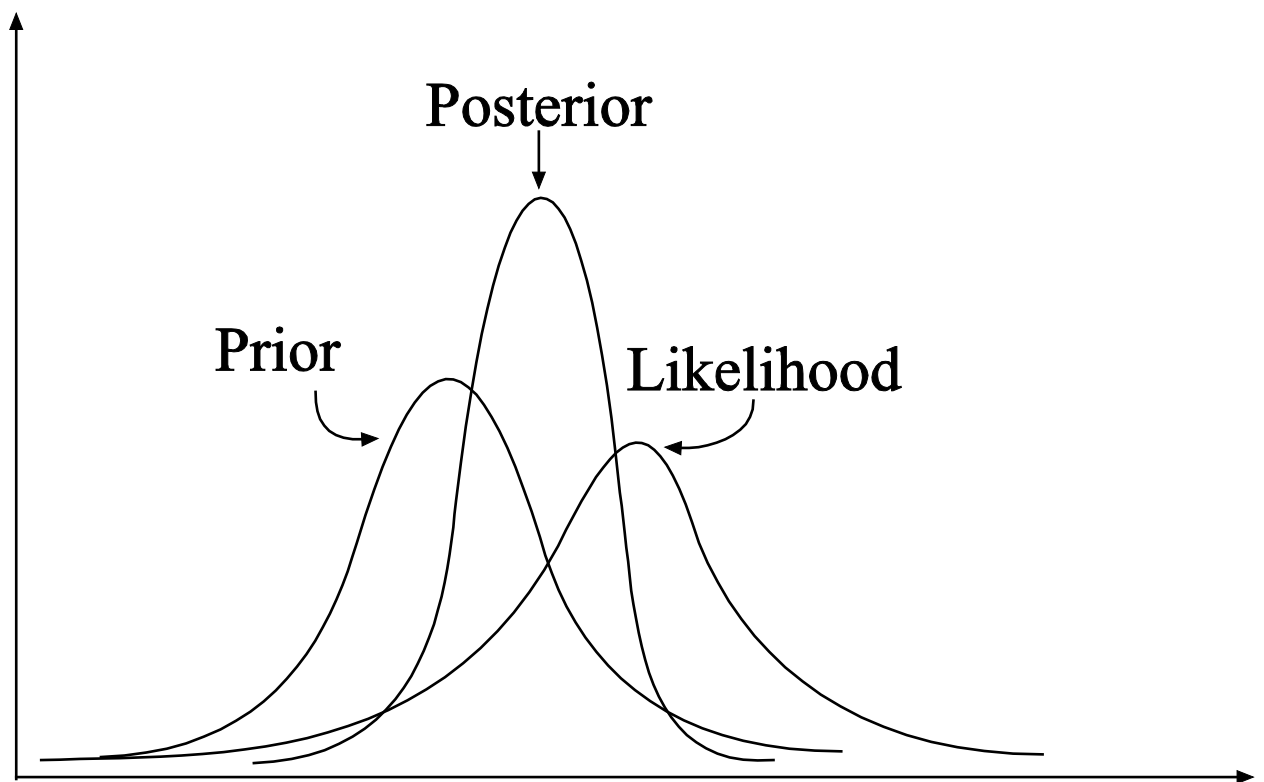
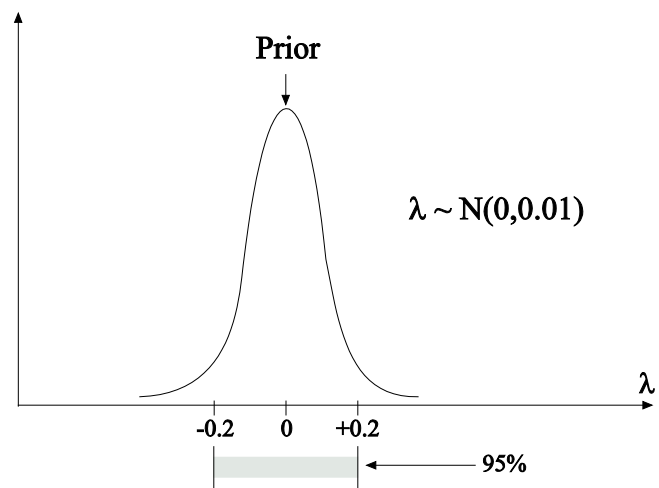


Figure 2: Informative prior for a factor loading parameter



List of Tables

1	Choice of variance for a normal prior with mean zero	53
2	Holzinger-Swineford's hypothesized four domains measured by 19 tests	54
3	ML model testing results for Holzinger-Swineford data for Grant-White ($n = 145$) and Pasteur ($n = 156$)	55
4	Holzinger-Swineford ML EFA using 19 variables and Geomin rotation: Four-factor solution	56
5	ML versus Bayes model testing results for Holzinger-Swineford data for Grant-White ($n = 145$) and Pasteur ($n = 156$)	57
6	Bayes for Holzinger-Swineford example: Four-factor solution using informative priors for cross-loadings	58
7	Effects of using different variances for the informative priors of the cross-loadings for the Holzinger-Swineford data	59
8	Factor loading pattern for simulation study of cross-loadings	60
9	Rejection rates for ML CFA and Bayes CFA with non-informative priors	61
10	Bayesian analysis using informative, small-variance priors for cross-loadings 0.0 and 0.1	62
11	Bayesian analysis using informative, small-variance priors for cross-loadings 0.2 and 0.3	63
12	Wording and hypothesized factor loading pattern for the 15 items used to measure the big-five personality factors in the British Household Panel data ("I see myself as someone who ...")	64
13	ML model testing results for big-five personality factors using British Household Panel data for females ($n = 691$) and males ($n = 589$) . . .	65
14	ML exploratory factor analysis of the big-five personality factors using British Household Panel data	66
15	Bayesian analysis using informative, small-variance priors for residual correlations using data for BHPS females and males, Method 2	67

Table 1: Choice of variance for a normal prior with mean zero

Variance	90% limits	95% limits
0.001	± 0.05	± 0.06
0.005	± 0.12	± 0.14
0.01	± 0.16	± 0.20
0.02	± 0.23	± 0.28
0.03	± 0.28	± 0.34
0.04	± 0.33	± 0.39
0.05	± 0.37	± 0.44
0.06	± 0.40	± 0.48
0.07	± 0.44	± 0.52
0.08	± 0.47	± 0.55
0.09	± 0.49	± 0.59
0.10	± 0.52	± 0.62

Table 2: Holzinger-Swineford's hypothesized four domains measured by 19 tests

Factor loading pattern				
	Spatial	Verbal	Speed	Memory
visual	X	0	0	0
cubes	X	0	0	0
paper	X	0	0	0
flags	X	0	0	0
general	0	X	0	0
paragrap	0	X	0	0
sentence	0	X	0	0
wordc	0	X	0	0
wordm	0	X	0	0
addition	0	0	X	0
code	0	0	X	0
counting	0	0	X	0
straight	0	0	X	0
wordr	0	0	0	X
numberr	0	0	0	X
figurer	0	0	0	X
object	0	0	0	X
numberf	0	0	0	X
figurew	0	0	0	X

Table 3: ML model testing results for Holzinger-Swineford data for Grant-White ($n = 145$) and Pasteur ($n = 156$)

Model	χ^2	df	p-value	RMSEA	CFI
Grant-White					
CFA	216	146	0.000	0.057	0.930
EFA	110	101	0.248	0.025	0.991
Pasteur					
CFA	261	146	0.000	0.071	0.882
EFA	128	101	0.036	0.041	0.972

Table 4: Holzinger-Swineford ML EFA using 19 variables and Geomin rotation: Four-factor solution

Loadings

	Grant-White				Pasteur			
	Spatial	Verbal	Speed	Memory	Spatial	Verbal	Speed	Memory
visual	0.628*	0.065	0.091	0.085	0.580*	0.307*	-0.001	0.053
cubes	0.485*	0.050	0.007	-0.003	0.521*	0.027	-0.078	-0.059
paper	0.406*	0.107	0.084	0.083	0.484*	0.101	-0.016	-0.229*
flags	0.579*	0.160	0.013	0.026	0.687*	-0.051	0.067	0.101
general	0.042	0.752*	0.126	-0.051	-0.043	0.838*	0.042	-0.118
paragrap	0.021	0.804*	-0.056	0.098	0.026	0.800*	-0.006	0.069
sentence	-0.039	0.844*	0.085	-0.057	-0.045	0.911*	-0.054	-0.029
wordc	0.094	0.556*	0.197*	0.019	0.098	0.695*	0.008	0.083
wordm	0.004	0.852*	-0.074	0.069	0.143*	0.793*	0.029	-0.023
addition	-0.302*	0.029	0.824*	0.078	-0.247*	0.067	0.664*	0.026
code	0.012	0.050	0.479*	0.279*	0.004	0.262*	0.552*	0.082
counting	0.045	-0.159	0.826*	-0.014	0.073	-0.034	0.656*	-0.166
straight	0.346*	0.043	0.570*	-0.055	0.266*	-0.034	0.526*	-0.056
wordr	-0.024	0.117	-0.020	0.523*	-0.005	0.020	-0.039	0.726*
numberr	0.069	0.021	-0.026	0.515*	-0.026	-0.057	-0.057	0.604*
figurer	0.354*	-0.033	-0.077	0.515*	0.329*	0.042	0.168	0.403*
object	-0.195	0.045	0.154	0.685*	-0.123	-0.005	0.333*	0.469*
numberf	0.225	-0.127	0.246*	0.450*	-0.014	0.092	0.092	0.427*
figurew	0.069	0.099	0.058	0.365*	0.139	0.013	0.237*	0.291*

Factor Correlations

	Grant-White				Pasteur			
	Spatial	Verbal	Speed	Memory	Spatial	Verbal	Speed	Memory
Spatial	1.000				1.000			
Verbal	0.378*	1.000			0.186*	1.000		
Speed	0.372*	0.386*	1.000		0.214	0.326*	1.000	
Memory	0.307*	0.380*	0.375*	1.000	0.190*	0.100	0.242*	1.000

Table 5: ML versus Bayes model testing results for Holzinger-Swineford data for Grant-White ($n = 145$) and Pasteur ($n = 156$)

ML analysis					
Model	χ^2	df	p-value	RMSEA	CFI
Grant-White					
CFA	216	146	0.000	0.057	0.930
EFA	110	101	0.248	0.025	0.991
Pasteur					
CFA	261	146	0.000	0.071	0.882
EFA	128	101	0.036	0.041	0.972
Bayesian analysis					
Model	Sample LRT	2.5% PP limit	97.5% PP limit	PP p-value	
Grant-White					
CFA	219	12	112	0.006	
CFA w. cross-loadings	142	-39	61	0.361	
Pasteur					
CFA	264	56	156	0.000	
CFA w. cross-loadings	156	-28	76	0.162	

Table 6: Bayes for Holzinger-Swineford example: Four-factor solution using informative priors for cross-loadings

Loadings

	Grant-White				Pasteur			
	Spatial	Verbal	Speed	Memory	Spatial	Verbal	Speed	Memory
visual	0.640*	0.012	0.050	0.047	0.633*	0.145	0.027	0.039
cubes	0.521*	-0.008	-0.010	-0.012	0.504*	-0.027	-0.041	-0.030
paper	0.456*	0.040	0.041	0.047	0.515*	0.018	-0.024	-0.118
flags	0.672*	0.046	-0.020	0.005	0.677*	-0.095	0.026	0.093
general	0.037	0.788*	0.049	-0.040	-0.056	0.856*	0.027	-0.084
paragrap	-0.001	0.837*	-0.053	0.030	0.015	0.801*	-0.011	0.050
sentence	-0.045	0.885*	0.021	-0.055	-0.063	0.925*	-0.032	-0.036
wordc	0.053	0.612*	0.096	0.029	0.055	0.694*	0.013	0.063
wordm	-0.012	0.886*	-0.086	0.020	0.092	0.803*	0.001	0.012
addition	-0.172*	0.030	0.795*	0.004	-0.147	-0.004	0.655*	0.010
code	-0.002	0.054	0.560*	0.130	-0.004	0.111	0.655*	0.049
counting	0.013	-0.092	0.828*	-0.049	0.025	-0.058	0.616*	-0.057
straight	0.189*	0.043	0.633*	-0.035	0.132	-0.067	0.558*	0.001
wordr	-0.040	0.044	-0.031	0.556*	-0.058	0.006	-0.090	0.731*
numberr	0.003	-0.004	-0.038	0.552*	0.006	-0.098	-0.106	0.634*
figurer	0.132	-0.024	-0.049	0.573*	0.156*	0.027	0.064	0.517*
object	-0.139	0.014	0.029	0.724*	-0.097	0.007	0.122	0.545*
numberf	0.099	-0.071	0.095	0.564*	-0.029	0.041	0.003	0.474*
figurew	0.012	0.045	0.007	0.445*	0.049	0.018	0.085	0.397*

Factor Correlations

	Grant-White				Pasteur			
	Spatial	Verbal	Speed	Memory	Spatial	Verbal	Speed	Memory
Spatial	1.000				1.000			
Verbal	0.535*	1.000			0.348*	1.000		
Speed	0.471*	0.443*	1.000		0.307	0.457*	1.000	
Memory	0.526*	0.515*	0.557*	1.000	0.324*	0.179	0.405*	1.000

Table 7: Effects of using different variances for the informative priors of the cross-loadings for the Holzinger-Swineford data

Prior variance	95% cross-loading limit	PPP	Cross-loading (Posterior SD)	Factor corr. range
Grant-White				
0.01	0.20	0.361	0.189 (.078)	0.443 - 0.557
0.02	0.28	0.441	0.248 (.096)	0.439 - 0.542
0.03	0.34	0.457	0.275 (.109)	0.423 - 0.530
0.04	0.39	0.455	0.292 (.120)	0.413 - 0.521
0.05	0.44	0.453	0.303 (.130)	0.404 - 0.513
0.06	0.48	0.447	0.309 (.139)	0.400 - 0.510
0.07	0.52	0.439	0.315 (.148)	0.395 - 0.508
0.08	0.55	0.439	0.319 (.156)	0.387 - 0.508
0.09	0.59	0.435	0.323 (.163)	0.378 - 0.506
1.00	0.62	0.427	0.327 (.171)	0.369 - 0.504
Pasteur				
0.01	0.20	0.162	0.132 (.076)	0.179 - 0.457
0.02	0.28	0.205	0.201 (.088)	0.184 - 0.441
0.03	0.34	0.219	0.223 (.098)	0.188 - 0.431
0.04	0.39	0.218	0.237 (.106)	0.189 - 0.424
0.05	0.44	0.205	0.247 (.115)	0.175 - 0.408
0.06	0.48	0.196	0.255 (.122)	0.175 - 0.402
0.07	0.52	0.195	0.261 (.128)	0.176 - 0.397
0.08	0.55	0.192	none	0.176 - 0.394
0.09	0.59	0.187	none	0.177 - 0.391
0.10	0.62	0.185	none	0.177 - 0.388

Table 8: Factor loading pattern for simulation study of cross-loadings

Factor loading pattern			
	F1	F2	F3
y1	X	0	x
y2	X	0	0
y3	X	0	0
y4	X	0	0
y5	X	0	0
y6	x	X	0
y7	0	X	0
y8	0	X	0
y9	0	X	0
y10	0	X	0
y11	0	x	X
y12	0	0	X
y13	0	0	X
y14	0	0	X
y15	0	0	X

Table 9: Rejection rates for ML CFA and Bayes CFA with non-informative priors

Cross-loading	Sample size	ML LRT rejection rate	Bayes PPP rejection rate
0.0	100	0.172	0.036
	200	0.090	0.032
	500	0.060	0.024
0.1	100	0.226	0.056
	200	0.228	0.080
	500	0.460	0.262
0.2	100	0.488	0.196
	200	0.726	0.474
	500	0.998	0.984
0.3	100	0.830	0.544
	200	0.996	0.944
	500	1.000	1.000

Table 10: Bayesian analysis using informative, small-variance priors for cross-loadings 0.0 and 0.1

Parameter	Estimates			S.E.	M.S.E.	95%	% Sig
	Population	Average	Std. Dev.	Average		Cover	Coeff
Cross-loading = 0.0, n=100, 5% reject proportion for the PPP = 0.006							
Major loading	0.800	0.8472	0.1212	0.1310	0.0169	0.950	1.000
Cross-loading	0.000	0.0141	0.0455	0.0732	0.0023	0.998	0.002
Factor variance	1.000	0.9864	0.2595	0.2624	0.0674	0.930	1.000
Factor correlation	0.500	0.4967	0.0869	0.1076	0.0075	0.980	0.982
Cross-loading = 0.0, n=200, 5% reject proportion for the PPP = 0.002							
Major loading	0.800	0.8311	0.0893	0.0894	0.0089	0.942	1.000
Cross-loading	0.000	0.0079	0.0421	0.0633	0.0018	1.000	0.000
Factor variance	1.000	0.9662	0.1799	0.1840	0.0335	0.948	1.000
Factor correlation	0.500	0.4962	0.0605	0.0860	0.0037	0.990	1.000
Cross-loading = 0.0, n=500, 5% reject proportion for the PPP = 0.010							
Major loading	0.800	0.8161	0.0520	0.0552	0.0030	0.962	1.000
Cross-loading	0.000	0.0033	0.0313	0.0530	0.0010	1.000	0.000
Factor variance	1.000	0.9741	0.1096	0.1293	0.0127	0.958	1.000
Factor correlation	0.500	0.4960	0.0406	0.0708	0.0017	1.000	1.000
Cross-loading = 0.1, n=100, 5% reject proportion for the PPP = 0.006							
Major loading	0.800	0.8218	0.1177	0.1263	0.0143	0.950	1.000
Cross-loading	0.100	0.0594	0.0449	0.0728	0.0037	0.982	0.038
Factor variance	1.000	1.0600	0.2738	0.2808	0.0784	0.934	1.000
Factor correlation	0.500	0.5206	0.0850	0.1047	0.0076	0.976	0.992
Cross-loading = 0.1, n=200, 5% reject proportion for the PPP = 0.006							
Major loading	0.800	0.8151	0.0860	0.0882	0.0076	0.950	1.000
Cross-loading	0.000	0.0666	0.0424	0.0636	0.0029	0.978	0.098
Factor variance	1.000	1.0217	0.1895	0.1978	0.0363	0.942	1.000
Factor correlation	0.500	0.5204	0.0602	0.0843	0.0040	0.984	1.000
Cross-loading = 0.1, n=500, 5% reject proportion for the PPP = 0.008							
Major loading	0.800	0.8089	0.0517	0.0551	0.0027	0.964	1.000
Cross-loading	0.100	0.0732	0.0316	0.0532	0.0017	0.990	0.176
Factor variance	1.000	1.0169	0.1160	0.1371	0.0137	0.972	1.000
Factor correlation	0.500	0.5229	0.0404	0.0688	0.0022	0.998	1.000

Table 11: Bayesian analysis using informative, small-variance priors for cross-loadings 0.2 and 0.3

Parameter	Estimates			S.E.	M.S.E.	95%	% Sig
	Population	Average	Std. Dev.	Average		Cover	Coeff
Cross-loading = 0.2, n=100, 5% reject proportion for the PPP = 0.010							
Major loading	0.800	0.7952	0.1152	0.1217	0.0133	0.952	1.000
Cross-loading	0.200	0.1024	0.0453	0.0731	0.0116	0.840	0.188
Factor variance	1.000	1.1522	0.2990	0.3018	0.1124	0.908	1.000
Factor correlation	0.500	0.5439	0.0819	0.1023	0.0086	0.966	0.996
Cross-loading = 0.2, n=200, 5% reject proportion for the PPP = 0.004							
Major loading	0.800	0.7979	0.0835	0.0859	0.0070	0.940	1.000
Cross-loading	0.200	0.1239	0.0418	0.0638	0.0075	0.856	0.492
Factor variance	1.000	1.0850	0.1978	0.2109	0.0463	0.942	1.000
Factor correlation	0.500	0.5424	0.0581	0.0823	0.0052	0.974	1.000
Cross-loading = 0.2, n=500, 5% reject proportion for the PPP = 0.006							
Major loading	0.800	0.8010	0.0514	0.0554	0.0026	0.964	1.000
Cross-loading	0.200	0.1427	0.0327	0.0541	0.0044	0.922	0.854
Factor variance	1.000	1.0595	0.1222	0.1473	0.0184	0.966	1.000
Factor correlation	0.500	0.5445	0.0382	0.0669	0.0034	0.986	1.000
Cross-loading = 0.3, n=100, 5% reject proportion for the PPP = 0.012							
Major loading	0.800	0.7671	0.1104	0.1166	0.0139	0.924	1.000
Cross-loading	0.300	0.1428	0.0460	0.0734	0.0268	0.364	0.470
Factor variance	1.000	1.2532	0.3158	0.3281	0.1636	0.860	1.000
Factor correlation	0.500	0.5650	0.0810	0.0999	0.0108	0.952	0.996
Cross-loading = 0.3, n=200, 5% reject proportion for the PPP = 0.012							
Major loading	0.800	0.7790	0.0807	0.0836	0.0069	0.948	1.000
Cross-loading	0.300	0.1755	0.0419	0.0640	0.0173	0.518	0.856
Factor variance	1.000	1.1623	0.2134	0.2257	0.0718	0.890	1.000
Factor correlation	0.500	0.5642	0.0577	0.0804	0.0075	0.950	1.000
Cross-loading = 0.3, n=500, 5% reject proportion for the PPP = 0.006							
Major loading	0.800	0.7891	0.0493	0.0553	0.0025	0.958	1.000
Cross-loading	0.300	0.2077	0.0318	0.0545	0.0095	0.640	1.000
Factor variance	1.000	1.1116	0.1252	0.1573	0.0281	0.930	1.000
Factor correlation	0.500	0.5636	0.0368	0.0661	0.0054	0.960	1.000

Table 12: Wording and hypothesized factor loading pattern for the 15 items used to measure the big-five personality factors in the British Household Panel data ("I see myself as someone who ...")

	Agreeableness	Conscientiousness	Extraversion	Neuroticism	Openness
y1: Is sometimes rude to others (reverse-scored)	X	0	0	0	0
y2: Has a forgiving nature	X	0	0	0	0
y3: Is considerate and kind to almost everyone	X	0	0	0	0
y4: Does a thorough job	0	X	0	0	0
y5: Tends to be lazy (reverse-scored)	0	X	0	0	0
y6: Does things efficiently	0	X	0	0	0
y7: Is talkative	0	0	X	0	0
y8: Is outgoing, sociable	0	0	X	0	0
y9: Is reserved (reverse-scored)	0	0	X	0	0
y10: Worries a lot	0	0	0	X	0
y11: Gets nervous easily	0	0	0	X	0
y12: Is relaxed, handles stress well (reverse-scored)	0	0	0	X	0
y13: Is original, comes up with new ideas	0	0	0	0	X
y14: Values artistic, aesthetic experiences	0	0	0	0	X
y15: Has an active imagination	0	0	0	0	X

Table 13: ML model testing results for big-five personality factors using British Household Panel data for females ($n = 691$) and males ($n = 589$)

Model	χ^2	df	p-value	RMSEA	CFI
Females					
CFA	552	80	0.000	0.092	0.795
CFA + CUs	432	74	0.000	0.084	0.845
EFA	183	40	0.000	0.072	0.938
Males					
CFA	516	80	0.000	0.096	0.795
CFA + CUs	442	74	0.000	0.092	0.826
EFA	113	40	0.000	0.056	0.965

Table 14: ML exploratory factor analysis of the big-five personality factors using British Household Panel data

Loadings

	Females					Males				
	F1	F2	Extraver	Neurot	Open	F1	F2	Extraver	Neurot	Open
y1	0.827*	0.000	0.014	-0.011	-0.005	0.389*	0.010	-0.016	-0.083	-0.294*
y2	0.147*	0.215*	0.323*	0.033	0.020	0.188	0.447*	0.123*	-0.030	-0.011
y3	0.103*	0.569*	0.280*	0.046	-0.095*	0.506*	0.469*	0.026	0.042	-0.030
y4	0.018	0.455*	-0.003	-0.025	0.270*	0.406*	-0.011	0.119*	0.010	0.272*
y5	0.365*	0.220*	-0.039	-0.068	0.009	0.654*	-0.449*	-0.037	0.009	0.004
y6	-0.016	0.852*	0.001	-0.052	0.087	0.656*	0.077	0.020	-0.090	0.141*
y7	-0.154*	0.053	0.541*	0.015	0.129*	0.047	0.015	0.629*	0.045	0.123
y8	-0.041	-0.024	0.748*	-0.049	-0.002	0.012	0.024	0.795*	-0.051	0.032
y9	0.064	-0.416*	0.346*	-0.116*	0.031	-0.156	-0.380*	0.396*	-0.010	-0.049
y10	-0.045	0.061	0.063	0.727*	0.036	0.023	0.020	-0.029	0.698*	0.294*
y11	0.021	0.001	-0.039	0.670*	0.013	-0.075	0.279*	-0.052	0.519*	0.021
y12	0.022	-0.250*	-0.061	0.547*	-0.063	-0.004	-0.311*	0.051	0.648*	-0.109
y13	0.024	0.011	-0.036	-0.107	0.764*	0.069	-0.007	-0.001	-0.023	0.734*
y14	0.011	-0.088*	0.038	0.125*	0.659*	-0.073	0.057	0.035	0.036	0.486*
y15	-0.054	0.140*	0.087	-0.014	0.541*	0.002	0.006	0.038	-0.146*	0.671*
Factor correlations										
	Females					Males				
	F1	F2	Extraver	Neurot	Open	F1	F2	Extraver	Neurot	Open
F1	1.000					1.000				
F2	0.151*	1.000				0.399*	1.000			
Extraver	-0.024	0.362*	1.000			0.268*	0.200*	1.000		
Neurot	-0.085	0.100*	-0.108*	1.000		-0.311*	0.065	-0.252*	1.000	
Open	-0.142*	0.229*	0.473*	-0.175*	1.000	0.344*	0.397*	0.454*	-0.180*	1.000

Table 15: Bayesian analysis using informative, small-variance priors for residual correlations using data for BHPS females and males, Method 2

Loadings

	Females					Males				
	Agreeab	Conscien	Extraver	Neurot	Open	Agreeab	Conscien	Extraver	Neurot	Open
y1	0.772*	-0.006	-0.026	0.000	-0.012	0.842*	-0.013	-0.011	-0.010	-0.018
y2	0.575	-0.014	0.021	-0.013	0.028	0.394	-0.006	0.024	-0.006	0.018
y3	0.503*	0.034	0.023	0.012	-0.010	0.479*	0.040	0.005	0.021	0.013
y4	-0.029	0.704*	0.014	-0.003	0.024	-0.040	0.683*	0.027	0.019	0.017
y5	0.017	0.657*	-0.001	0.006	-0.028	0.014	0.708*	-0.020	0.002	-0.018
y6	0.032	0.548*	-0.015	-0.007	0.015	0.043	0.579*	0.000	-0.036	0.007
y7	-0.008	0.014	0.685*	0.024	0.006	-0.005	0.005	0.748*	0.016	-0.005
y8	0.023	0.002	0.702*	-0.017	0.003	0.024	0.011	0.754*	-0.023	0.013
y9	-0.016	-0.021	0.622*	-0.008	0.002	-0.025	-0.015	0.575*	0.005	-0.005
y10	-0.003	0.025	0.022	0.791*	0.023	0.001	0.009	0.013	0.801*	0.044
y11	0.016	-0.007	-0.024	0.736*	-0.008	0.012	-0.010	-0.022	0.708*	-0.024
y12	-0.012	-0.022	-0.004	0.695*	-0.027	-0.017	-0.006	0.003	0.613*	-0.034
y13	0.006	0.020	0.006	-0.047	0.780*	0.004	0.023	-0.008	0.007	0.732*
y14	-0.008	-0.010	-0.007	0.046	0.738*	0.004	-0.021	-0.013	0.023	0.672*
y15	0.006	-0.006	0.011	-0.003	0.660*	-0.011	0.001	0.031	-0.035	0.651*

Factor correlations

	Females					Males				
	Agreeab	Conscien	Extraver	Neurot	Open	Agreeab	Conscien	Extraver	Neurot	Open
Agreeab	1.000					1.000				
Conscien	0.366*	1.000				0.319*	1.000			
Extraver	0.081	0.119	1.000			0.025	0.197	1.000		
Neurot	-0.059	-0.093	-0.163	1.000		-0.133	-0.238*	-0.160	1.000	
Open	0.041	0.201	0.321*	-0.158	1.000	0.040	0.250	0.297*	-0.091	1.000