

Psychological Methods

A Factored Regression Approach to Modeling Latent Variable Interactions and Nonlinear Effects

Brian T. Keller

Online First Publication, October 30, 2025. <https://dx.doi.org/10.1037/met0000777>

CITATION

Keller, B. T. (2025). A factored regression approach to modeling latent variable interactions and nonlinear effects. *Psychological Methods*. Advance online publication. <https://dx.doi.org/10.1037/met0000777>

A Factored Regression Approach to Modeling Latent Variable Interactions and Nonlinear Effects

Brian T. Keller

College of Education, The University of Missouri

Abstract

Interaction effects are common in the behavioral sciences, especially in psychology, as they help explore how various factors influence human behavior. This article introduces a factored regression framework designed to estimate latent variable interactions and nonlinear effects, providing a flexible approach for modeling complex data structures, accommodating diverse data types, and handling missing data on any variable. The factored regression framework also allows graphical diagnostics to probe interactions effectively. Monte Carlo simulations were conducted to compare the performance of factored regression with existing maximum likelihood methods, such as latent moderated structural equations and product indicators. Results indicate that factored regression performs comparably to, if not better than, these traditional methods. The factored regression framework is implemented in Blimp software, offering an accessible and user-friendly syntax for specifying the models. Through practical examples and syntax excerpts, this article demonstrates the application of factored regression for estimating latent interactions, making it more approachable for a wide audience in behavioral research.

Translational Abstract

Interaction or moderation effects occur when the relationship between two variables depends on a third variable. Latent variable models involve unobserved variables, the values of which depend on a set of observed scores that measure that construct. This article describes an approach that allows researchers to easily estimate numerous types of interactions involving combinations of latent and observed variables. Free software is provided, and numerous illustrative examples demonstrate its application.

Keywords: Bayesian analysis, latent variable interactions, three-way latent interactions, structural equation modeling, factored regression

Interaction or moderation effects are ubiquitous in the behavioral sciences. Within psychology, these interactions stand as captivating phenomena because they illuminate the intricacies inherent in human behavior, shedding light on how a relationship is influenced by additional factors. Without surprise, interactions play a pivotal role in regression modeling, and substantial methodological work has been devoted to extending moderation effects to latent variable models.

Latent variable approaches to modeling interactions are typically favored over manifest variable interactions used in traditional linear regression contexts, primarily because they explicitly account for measurement error. Ignoring measurement error can lead to biased

estimates, weakened statistical power, and potentially erroneous conclusions regarding moderating effects (Baron & Kenny, 1986). By conceptualizing constructs as latent variables inferred from multiple observed indicators, latent variable models significantly mitigate these biases, providing more precise and accurate interaction estimates.

Several methodological strategies have been developed to effectively incorporate interaction terms into latent variable frameworks. Early advancements relied on product indicator (PI) methods, which involve creating interaction terms by multiplying observed indicators from different latent constructs. While straightforward and intuitively appealing, these methods often encounter limitations such as increased model complexity and strict distributional assumptions.

Other techniques, including latent moderated structural equations (LMS), address some limitations inherent to PI methods by directly modeling interactions between latent variables without explicitly forming product terms. LMS methods thus offer greater flexibility and improved estimation efficiency. Additionally, recent interest has grown in Bayesian methodologies due to their versatility in specifying intricate models, accommodating nonnormal distributions, managing categorical moderators, and effectively handling missing data. The present article introduces a novel Bayesian generalization termed the factored regression (FR) specification.

The subsequent sections will explore these three distinct methodological traditions, highlighting their strengths, practical applications, and limitations. Following these discussions, the article

Samantha F. Anderson served as action editor.

Brian T. Keller  <https://orcid.org/0009-0007-6670-2864>

This work was supported by the Institute of Educational Sciences (Award R305D190002). There are no conflicts of interest to report. The ideas in this article have not been presented elsewhere. The additional online materials are available on the Open Science Framework page at <https://doi.org/10.17605/OSF.IO/2GHNYY>.

Brian T. Keller served as lead for methodology, software, writing—original draft, and writing—review and editing.

Correspondence concerning this article should be addressed to Brian T. Keller, College of Education, The University of Missouri, 16 Hill Hall, Columbia, MO 65211, United States. Email: btceller@missouri.edu

will present the FR approach, detailing its unique features and advantages.

PI Methods

Earlier methods for estimating latent variable interactions originated from pioneering research by Kenny and Judd (1984). Known as “product indicator” (PI) methods, these techniques follow a structured approach involving multiple clearly defined steps. Initially, PIs are created by multiplying select manifest indicators from different latent variables. These generated manifest products serve as indicators for a new latent variable explicitly representing the interaction between the original latent variables (Cham et al., 2012; Kelava & Brandt, 2023).

Several strategies have been proposed for selecting PIs, and the choice significantly impacts the method’s performance. As the number of original indicators grows, so does the number of possible PIs, complicating model specification. A commonly recommended approach is the three-matched pairs strategy, which pairs the three most reliable indicators from each construct. This approach has consistently shown good results in simulation studies and enjoys empirical support (Aytürk et al., 2020; Cham et al., 2012; Marsh et al., 2004).

The statistical properties of latent product variables under normality assumptions are important to consider. When original indicators are normally distributed, the mean, variance, and covariances of product variables are fully determined (Bohrnstedt & Goldberger, 1969). Marsh et al. (2004) introduced an unconstrained PI approach to address potential nonnormality, relaxing traditional parameter constraints. Building on this, Kelava’s extended unconstrained PI (EXT) method explicitly estimates all latent product parameters, incorporating additional covariances among indicators to address nonnormality (Kelava, 2009; Kelava & Brandt, 2009). Lonati et al. (2024) also support the robustness of the EXT method, highlighting its performance in handling nonnormal latent variables.

Despite these methodological advancements, PI approaches have both strengths and potential limitations. On the positive side, PI methods provide unbiased estimates with good statistical properties, even with moderate deviations from normality (Aytürk et al., 2020; Cham et al., 2012). Additionally, they offer transparency compared to alternative latent interaction modeling methods, such as LMS (Cham et al., 2012). However, incorporating multiple PIs can lead to less parsimonious models, potentially affecting estimation efficiency and statistical power (Aytürk et al., 2020; Kim et al., 2024). Specifically, the matched pairs strategy might yield weaker test statistics and larger standard errors than methods that use more PIs, reflecting a balance between simplicity and precision (Aytürk et al., 2021).

PI methods typically assume continuous manifest variables, limiting their application to binary or ordinal data. Although research shows adequate performance for PI methods when categorical indicators are symmetrically distributed with at least five scale points (Aytürk et al., 2020, 2021), performance tends to deteriorate when indicators are asymmetrically distributed, potentially leading to biased estimates or reduced statistical power. The issue of categorical variables is particularly significant in the context of this study, as the FR approach outlined in this article specifically addresses and readily incorporates binary and multicategorical moderator variables.

Finally, missing data presents another crucial consideration. PI methods align with the “just another variable” approach to missing data, which should be viewed as a last resort because it requires the

restrictive missing completely at random (MCAR) assumption (Enders et al., 2014; Lüdtke et al., 2020a, 2020b; Zhang & Wang, 2017). Violations of this assumption are common and can result in substantial biases. Therefore, PI methods may yield biased estimates under conditions such as missing at random (MAR) or missing not at random (MNAR; Cham et al., 2017). In contrast, the FR approach presented in this article effectively accommodates MAR mechanisms and can integrate selection models to address MNAR mechanisms (Du et al., 2022).

LMS

LMS modeling is another widely utilized approach for estimating latent variable interactions, belonging to a broader class of models that Kelava et al. (2011) define as distribution analytic methods. LMS was introduced as a likelihood-based estimator by Klein and Moosbrugger (2000), alongside a closely related approach, quasi-maximum likelihood, introduced by Klein and Muthén (2007). Extensions of these methodologies also exist for hierarchical data structures, such as multilevel latent interaction models described by B. Muthén and Asparouhov (2009).

Conceptually, LMS differs notably from traditional PI methods. Instead of explicitly creating interaction terms from observed variables, LMS incorporates interactions by simultaneously modeling the joint distribution of latent variables involved. Distribution analytic approaches, such as LMS, directly specify the statistical distribution of latent predictors and their interactions, thus allowing for the estimation of interaction terms without the creation of separate PIs. Specifically, LMS uses techniques like Cholesky decomposition to orthogonalize latent predictors, simplifying the estimation of nonlinear latent interactions. This strategy effectively integrates the interaction-induced nonlinearity into the structural model without additional manifest product variables (Aytürk et al., 2020; Lonati et al., 2024). For readers seeking an accessible yet comprehensive technical introduction to LMS methods, Cham et al. (2017) provide an excellent overview.

Simulation studies have established conditions under which LMS yields approximately unbiased parameter estimates. For example, Aytürk et al. (2020) summarize research indicating that LMS performs adequately when assumptions of multivariate normality and accurate measurement models are satisfied. Lonati et al. (2024) support this conclusion with simulation findings demonstrating that LMS estimates remain approximately unbiased under mild nonnormality and accurately measured latent constructs. LMS generally provides superior statistical power, efficiency, and precision relative to PI approaches, particularly under conditions approximating normality or mild nonnormality. However, this advantage diminishes notably under severe violations of the normality assumption, highlighting contexts where LMS’s relative performance advantage is reduced or reversed.

A potential challenge for LMS methods involves modeling categorical indicators, which are frequently encountered in psychological research due to the widespread use of ordinal Likert-type scales or binary response formats. Traditional LMS approaches assume continuous indicators, which can lead to biased estimates when indicators are ordinal or binary. Simulation studies suggest that ordinal indicators with at least five response categories and normally distributed latent traits can reasonably be treated as continuous, but substantial bias emerges when these conditions are not met, particularly under asymmetric thresholds or fewer categories (Aytürk et al., 2020;

Kelava & Brandt, 2023). Recent developments specifically designed for ordinal indicators, handle these variables by using numerical integration methods with probit or logit link functions (Aytürk et al., 2021; Lonati et al., 2024; Shen et al., 2025).

Finally, current LMS implementations possess two substantial practical limitations that limit their broader use. First, the appropriate way to model interactions between latent variables and binary manifest moderators remains an open question. Research suggests that binary moderators inherently violate the normality assumption of LMS, which can lead to biased interaction estimates and incorrect standard errors unless alternative estimation strategies, such as categorical confirmatory factor models, are employed (Aytürk et al., 2021; Shen et al., 2025). Second, model specification becomes substantially cumbersome when applying the approach to higher-order effects like three-way interactions.

Bayesian Approaches to Latent Variable Interactions

Bayesian approaches to latent interactions represent a third major methodological development in the field. A defining characteristic of these methods is that latent variables are treated as missing data to be iteratively estimated, an approach known as data augmentation (Lee & Shi, 2000; Merkle & Rosseel, 2018; Palomo et al., 2007). This differs from the more conventional Bayesian structural equation modeling (SEM) framework, which places a joint multivariate distribution over both latent and manifest variables and integrates over the latent variables rather than explicitly imputing them (Kaplan & Depaoli, 2012; Merkle et al., 2020; B. Muthén & Asparouhov, 2012). The data augmentation strategy enables Bayesian SEMs to incorporate latent interactions by repeatedly sampling and updating the latent variable distributions within a Markov chain Monte Carlo (MCMC) process, facilitating more flexible modeling of nonlinear effects (Lee & Song, 2003; Lee et al., 2007).

Lee and colleagues extended the data augmentation framework specifically to nonlinear latent variable models, including interactions and polynomial effects (Lee, 2007; Lee & Song, 2003, 2012; Lee et al., 2007). In this approach, latent variables are iteratively updated, and their realized values are multiplied within the MCMC framework to estimate interaction terms. When an interaction effect is present, the distribution of the imputed latent variables takes the form of a scale mixture of normal distributions, leading to heteroscedastic variation in its distribution (Lee et al., 2004). Recent work by Asparouhov and Muthén (2021b) described a related data augmentation procedure for multilevel structural equation models with latent interactions. While these methods allow for flexible estimation of interactions, they are primarily developed under the general interaction model, which only allows for latent-by-latent interactions. Existing implementations do not explicitly address interactions involving categorical moderators or more complex nonlinear functions. The FR specification introduced in the next section builds directly upon Lee's approach, extending its core principles to accommodate a broader range of variable types and nonlinear transformations while retaining the flexibility of data augmentation-based estimation.

Although Bayesian estimation provides key advantages, practical implementation remains a challenge. Most applications require general MCMC estimation libraries, such as JAGS (Plummer, 2017) or Stan (Stan Development Team, 2023). These tools provide flexibility but often require researchers to write custom sampling

code and specify detailed prior distributions, making Bayesian SEM less accessible to applied researchers (Merkle et al., 2020). The FR specification in the next section addresses these limitations by providing a user-friendly software program that uses MCMC estimation to model a variety of interaction effects that are not currently supported in existing Bayesian SEM software.

Purpose of Current Study

The purpose of this article is to extend and generalize data augmentation-based Bayesian estimation for latent variable interactions. To this end, I describe a general FR framework for estimating a variety of nonlinear effects involving latent variables. A major issue with interaction and nonlinear effects is that they require distributional assumptions that are incompatible with a multivariate normal distribution. The FR framework avoids this problem by specifying multivariate distributions as a collection of simpler (often univariate) distributions that allow the specification of complex data structures required by interactive effects. This approach is related to distribution analytic methods in that it specifies the distribution of the interacting latent variables, albeit using a piecewise function comprised of multiple, simpler distributions. As a result, it should be asymptotically equivalent to LMS methods when the latent distributions are normal while offering the additional advantage of accommodating a wide range of manifest-by-latent variable interactions.

In the manifest missing data literature, the use of factorization strategies was first introduced by Ibrahim and colleagues (Ibrahim, 1990; Ibrahim et al., 1999, 2002; Lipsitz & Ibrahim, 1996). This work spawned recent innovations with Bayesian estimation, multiple imputation, and maximum likelihood (ML) estimation (for a review, see Enders, 2025). FR has been used to handle incomplete interactive and polynomial effects with manifest variables (Du et al., 2022; Enders et al., 2020; Erler et al., 2016; Goldstein et al., 2014; Keller & Enders, 2023; Lüdtke et al., 2020a, 2020b; Zhang & Wang, 2017). Several software packages offer factored specifications (Bartlett et al., 2021; Erler et al., 2021; Keller & Enders, 2022; Robitzsch & Lüdtke, 2021), although they differ in user-friendliness and the breadth of models they can fit.

By treating latent variables as missing data to be imputed (data augmentation), established FR specifications can be extended to a broad range of latent variable models. As explained previously, MCMC produces realized values of the latent variables at each computational cycle. When the latent interaction is nonzero, the latent imputes are sampled from a normal conditional distribution with heteroscedastic variation (a scale-mixture distribution), just like they are with incomplete manifest variables. This article builds on ideas from related work that views latent variables as missing data. Historically, Dempster et al. (1977) discussed latent variables as missing data when developing the EM algorithm to handle incomplete data, and this view has further been discussed in the missing data literature (Blackwell et al., 2017; Little & Rubin, 2002; Mealli & Rubin, 2015). Similarly, treating latent variables as missing data has also been discussed in areas of the Bayesian SEM literature (Lee, 2007; Lee & Shi, 2000; Lee et al., 2007; Merkle & Rosseel, 2018; Palomo et al., 2007). More recently, Lüdtke et al. (2020b) briefly discussed extending FR to latent variable models and, in Supplemental J of the additional online materials on the Open Science Framework (OSF) page at <https://doi.org/10.17605/OSF>

.IO/2GHNY, demonstrated the approach's ability to handle a single latent factor in a moderated regression model. In addition, Keller (2022) illustrated using the FR approach to estimate a moderated mediation model with two latent factors.

Developing a FR approach for interactive and nonlinear latent variable effects provides important methodological contributions. First, the procedure is comprehensive. Structural equation models can include two-way effects, three-way effects, quadratic and other polynomial effects, combinations of interactive and polynomial effects, and other nonlinear functions (e.g., sine, cosine, and exponentiation). Furthermore, nonlinear effects are not limited to the structural models; measurement models can include similar effects (e.g., squared loadings and moderated loadings). Second, FR expands the types of variables and interactions that can be included in a model. Because a structural model is decomposed into a series of submodels, each with unique distributional assumptions, an analysis may include combinations of binary, ordinal, multicategorical, continuous normal, skewed continuous, and count variables. This means that latent variables can interact seamlessly with any variable in a model. As an example, a researcher can readily specify an interaction where a multicategorical variable moderates the influence of a latent variable on a downstream outcome or one of its indicators. Finally, FR specifications leverage the MAR assumption where missingness is systematically related to one's observed data, but not to the unseen score values. This framework also accommodates MNAR processes using either the selection modeling or pattern mixture modeling framework (Du et al., 2022; Enders, 2022; Keller & Enders, 2022). Although this article's focus is single-level regression, these features all extend to multilevel models as well.

The real data analysis example in the article illustrates the FR modeling framework in Blimp (Keller & Enders, 2022), a free standalone software with syntax similarities to many modern-day SEM software packages. From the researcher's perspective, applying the FR approach to latent interactions is no different from standard moderated regression. The syntax looks like a regression equation, and this equation includes an explicit product operator where interacting variables are joined with an asterisk. Specifying latent variable interactions in a way that resembles standard moderated regression analysis offers an important practical benefit—familiarity and integration with established and familiar procedures. For example, I later illustrate how to apply pick-a-point approaches and Johnson–Neyman regions of significance plots (P. O. Johnson & Neyman, 1936) to probe latent interactions interaction. Finally, because data augmentation yields latent variable imputations at each MCMC cycle, researchers can readily apply graphical regression diagnostics. The analysis examples illustrate this functionality.

The structure of this article is as follows. First, I begin by discussing the FR approach to modeling and how to use it to incorporate measurement models. Second, I illustrate how factorization incorporates latent variable interactions. This includes a two-way latent interaction and a latent by manifest interaction. Third, I provide a high-level overview of the MCMC algorithm used to estimate the models in a Bayesian paradigm. Fourth, I provide a series of Monte Carlo simulation studies that evaluate the FR model's performance. To provide a comparison, the simulations also include the PI and LMS methods. Finally, I use a series of data analysis examples to demonstrate a model-building procedure that includes two-way and three-way latent variable effects.

FR Modeling

A key feature of the FR specification is that it avoids working with off-the-shelf multivariate distributions like the multivariate normal distribution. Instead, the FR framework aims to chain several regression models (i.e., conditional densities) to characterize the hypothesized joint distribution. Throughout the article, I represent models via a so-called “functional” notation. To illustrate, suppose I am modeling the joint distribution of three variables, X , Y , and Z . The functional notation representation of the joint distribution is as follows.

$$f(X, Y, Z), \quad (1)$$

where $f(\dots)$ represents a generic probability distribution for X , Y , and Z . This unspecified distribution could have any form. For example, I may assume that the three variables follow a multivariate normal distribution with a mean vector and covariance matrix

$$f(X, Y, Z) = \mathcal{N}_3(\boldsymbol{\mu}, \boldsymbol{\Sigma}), \quad (2)$$

where $\boldsymbol{\mu}$ is a 3×1 mean vector and $\boldsymbol{\Sigma}$ is a 3×3 unstructured covariance matrix for a total of nine parameters. Traditionally, SEM approaches work from a multivariate distribution like Equation 2. Ultimately, this is a limiting factor because interactions and nonlinear terms induce characteristics often at odds with the multivariate normal distribution (Liu et al., 2014; Seaman et al., 2012).

Alternatively, FR models the joint density as a sequential product of conditional and marginal densities. By applying the chain rule for probability, I can factor the joint density into a product of univariate distributions. Reconsider the joint density of the three variables, FR framework models this joint density as a sequential product of conditional and marginal densities.

$$f(X, Y, Z) = f(Y | X, Z) \times f(X, Z). \quad (3)$$

Note, I refer to this as a partially factored specification because the bivariate distribution $f(X, Z)$ can further be factored into two univariate densities.

$$f(X, Y, Z) = f(Y | X, Z) \times f(Z | X) \times f(X). \quad (4)$$

This factorization is sometimes called the “sequential” specification in the literature (Lüdtke et al., 2020a, 2020b). I will refer to this as the fully factored specification because I have wholly decomposed the joint density into a product of univariate models. In contrast to a traditional SEM approach, the fully factored specification models the multivariate distribution as a collection of univariate regression models, each potentially with its unique distribution features. This allows for greater flexibility to accurately represent special features induced by nonlinear models.

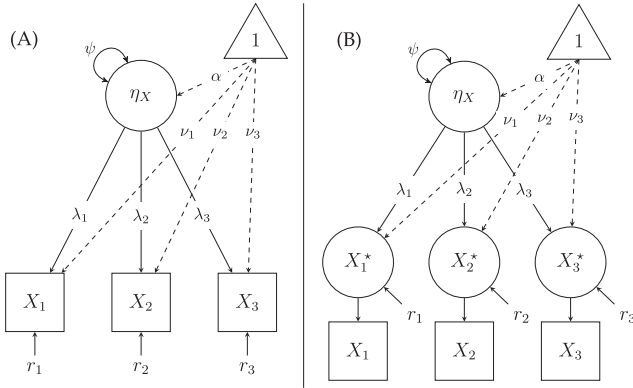
Measurement Models as FRs

To illustrate how FR applies to a measurement model, consider the model depicted in Figure 1A, where three variables (X_1 , X_2 , and X_3) load onto a single factor (η_X). To express this model in terms of a FR, I fully factor out the joint distribution of all variables into a product of four univariate distributions.

$$f(X_1, X_2, X_3, \eta_X) = f(X_1 | X_2, X_3, \eta_X) \times f(X_2 | X_3, \eta_X) \times f(X_3 | \eta_X) \times f(\eta_X) \quad (5)$$

$$= f(X_1 | \eta_X) \times f(X_2 | \eta_X) \times f(X_3 | \eta_X) \times f(\eta_X). \quad (6)$$

Figure 1
Path Diagrams for One Factor Models



Note. Panel A is the path diagram for a single-factor model with continuous indicators. Panel B is the path diagram for a single-factor model with binary indicators. The double-headed curved arrow on η represents the variance of the latent factor. The residuals, r_X , r_Y , and r_Z , also have estimated variances associated with them in (A) but are fixed to 1 in (B) for identification.

Equation 5 characterizes the full factorization of the measurement model. Based on the model in Figure 1A, I simplify the factorization by removing indicators from the right side of all vertical bars, giving Equation 6. This simplification is possible because each indicator is conditionally independent of the other indicator given the latent factor η_X . Note that the joint distribution can be factorized in different orders. For example, the first term to the right of the equals sign could be $f(X_3 | X_2, X_1, \eta_X)$.

The factorization in Equation 6 makes sense because the latent variable acts as an exogenous predictor, leaving it as the last variable in the chain and including its marginal distribution as the final distribution. The four distributions in Equation 6 map one-to-one onto a specific regression model

$$\begin{aligned} f(X_1 | \eta_X) &\rightarrow x_{1i} = v_1 + \lambda_1 \eta_{Xi} + r_{1i}, \\ f(X_2 | \eta_X) &\rightarrow x_{2i} = v_2 + \lambda_2 \eta_{Xi} + r_{2i}, \\ f(X_3 | \eta_X) &\rightarrow x_{3i} = v_3 + \lambda_3 \eta_{Xi} + r_{3i}, \\ f(\eta_X) &\rightarrow \eta_{Xi} = \alpha + \zeta_i. \end{aligned} \quad (7)$$

These equations represent the common factor model but are expressed in univariate regression models and distributions instead of the usual matrix specification and conditional multivariate normal invoked by SEM. For now, I have assumed that each indicator is continuous. However, as discussed later, FR readily accommodates nonnormal continuous, binary, ordinal, categorical, or count indicators by specifying the appropriate regression model for the data type. Finally, just like a standard confirmatory factor analysis, I must identify the model via constraints (e.g., $\lambda_1 = 1$, $\alpha = 0$) or narrow priors in a Bayesian context (B. Muthén & Asparouhov, 2012).

Some measurement model specifications require a partially factored specification. For example, one might require residual correlations between two indicators. The partially factored specification can accommodate these correlated residuals using a multivariate model. To illustrate, suppose that X_1 and X_2 have some method effect that

would induce a correlation between them outside of η_X . This correlation can be accommodated by applying a partially factored specification, thereby allowing for a conditional joint distribution between X_1 and X_2 .

$$f(X_1, X_2, X_3, \eta_X) = f(X_1, X_2 | \eta_X) \times f(X_3 | \eta_X) \times f(\eta_X). \quad (8)$$

Equation 8 characterizes the correlated residual between X_1 and X_2 by specifying a joint distribution for these two variables conditional on η_X . This joint distribution is equivalent to a multivariate regression model, allowing the residuals of X_1 and X_2 to be correlated.

In addition to continuously normally distributed indicators, FR readily extends to nonnormal indicators and exogenous predictors. For simplicity, suppose the indicators for η_X are all binary indicators. I use a probit model to specify the regression of each item onto η_X (Agresti, 2012; Albert & Chib, 1993; V. E. Johnson & Albert, 2006). The probit model conceptualizes the discrete responses (X_i) as originating from a latent normally distributed propensity (i.e., X_i^* for X_i) with a threshold dividing the latent variable into the discrete observations via a link function. I have depicted the path model in Figure 1B, with circles denoting the latent response variables. Notice that each indicator now has a latent propensity that is predicted by η_X ; therefore, each indicator must incorporate this unobserved latent propensity within the factorization.

For example, reconsider the factorization for X_1 given η_X in Equation 23. This density must now incorporate the latent propensity and the manifest indicator.

$$f(X_1 | \eta_X) = f(X_1 | X_1^*) \times f(X_1^* | \eta_X). \quad (9)$$

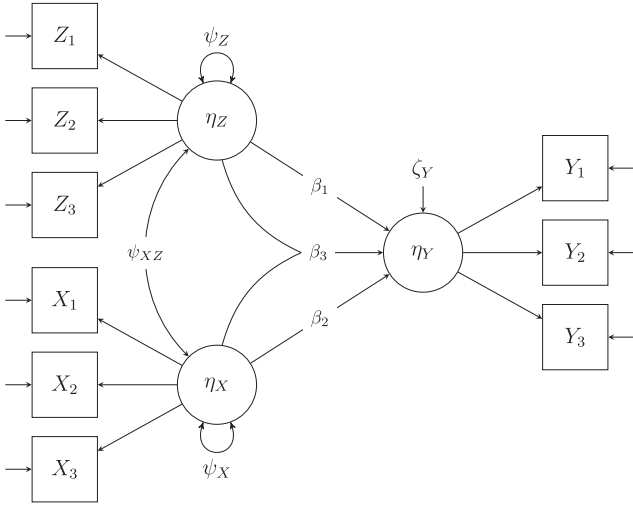
In words, the distribution of X_1 regressed on η_X is factored into two separate models, a conditional density of X_1 given the latent propensity X_1^* and a conditional density of X_1^* conditional on η_X . These two densities result in the following models:

$$\begin{aligned} f(X_1 | X_1^*) &\rightarrow x_{1i} = \mathcal{I}(x_{1i}^* > 0), \\ f(X_1^* | \eta_X) &\rightarrow x_{1i}^* = v_1 + \lambda_1 \eta_{Xi} + r_{1i}. \end{aligned} \quad (10)$$

The first line of Equation 10 is an additional “link” function that relates the observed binary variable to the latent propensity. In words, $\mathcal{I}(x_{1i}^* > 0)$ returns one if $x_{1i}^* > 0$ is true; otherwise, it returns zero. In Figure 1B, the link function is the arrow connecting the X_1^* to X_1 . The second line of the equation characterizes the endogenous normally distributed latent propensity with a mean that relates to the proportion of ones to zeros. Returning to Figure 1B, this is the measurement model where the latent factor influences the latent response. To identify this model, the variance of this propensity is constrained to one and the threshold parameter is fixed at zero. Ordinal indicators share the same model and factorization but simply require additional threshold parameters (e.g., a five-point indicator requires four thresholds with the first threshold fixed for identification). Finally, I have opted to model $f(X_1 | \eta_X)$ via the probit specification but nothing precludes using a logistic regression specification.

FR With Latent Interactions

The FR specification readily accommodates interactions and other types of nonlinearities involving latent variables. As an example, consider the latent moderation model depicted in Figure 2. The latent factor, η_Y , is regressed on η_X , η_Z , and their product (represented by the

Figure 2*Path Diagram for Latent × Latent Interaction Model*

Note. The mean structure is not illustrated in the diagram. The conjoined arrow for the β_3 path represents the product between η_X and η_Z .

conjoined arrow in the path diagram). Later in the article, I show an application mapping onto this exact example, where η_X and η_Z are latent variables measuring conscientiousness and organizational constraints, and η_Y is a latent measure of counterproductive work behavior.

To obtain the factorization for the model in Figure 2, I first separate the joint density into the manifest variables conditional on the unobserved latent factors.

$$f(X_1, X_2, X_3, Z_1, Z_2, Z_3, Y_1, Y_2, Y_3, \eta_X, \eta_Z, \eta_Y) = f(X_1, X_2, X_3, Z_1, Z_2, Z_3, Y_1, Y_2, Y_3, | \eta_X, \eta_Z, \eta_Y) \times f(\eta_X, \eta_Z, \eta_Y). \quad (11)$$

In words, Equation 11 says that the multivariate distribution of the manifest and latent variables can be expressed as the product of a conditional distribution that links the indicators to their respective latent variables and a multivariate distribution for the latent variables. This factorization directly maps onto Figure 2 because the manifest indicators are all conditionally independent given their respective latent variables. Said differently, the indicators are all predicted by latent variables only. Notice that this factorization maps directly onto how SEM models are conceptualized, with the first distribution corresponding to a measurement density and the second distribution corresponding to a structural density (Keller, 2022; Lüdtke et al., 2020b; Lee & Song, 2003, 2012).

Because the model implies each indicator is conditionally independent given their respective latent variable and there are no cross-loadings, I factor and simplify the measurement density into three separate distributions.

$$f(X_1, X_2, X_3, Z_1, Z_2, Z_3, Y_1, Y_2, Y_3, | \eta_X, \eta_Z, \eta_Y) = f(X_1, X_2, X_3 | \eta_X) \times f(Z_1, Z_2, Z_3 | \eta_Z) \times f(Y_1, Y_2, Y_3 | \eta_Y). \quad (12)$$

As before, the above factorization maps directly onto the path diagram in Figure 2. Notice how each distribution maps directly onto

the setup of the single-factor model. That is, three indicators are conditional on their respective latent variable. I further factor each of these densities just like I did in Equation 6, but with each respective grouping of variables—for example, $f(X_1, X_2, X_3 | \eta_X) = f(X_1 | \eta_X) \times f(X_2 | \eta_X) \times f(X_3 | \eta_X)$. Moreover, each distribution is equivalent to the first three lines in Equation 7. Note that the above specification readily extends to more than three indicators, and as discussed in the previous section, these indicators are not required to be continuous. The Blimp software also accommodates binary, ordinal, multicategorical, count, and skewed continuous indicators.

The factorization presented in Equation 12 assumes the absence of cross-loadings. However, this assumption does not always hold. For example, if X_1 exhibits a cross-loading with η_Z , the factorization in Equation 12 must be modified accordingly, as demonstrated below:

$$f(X_1, X_2, X_3, Z_1, Z_2, Z_3, Y_1, Y_2, Y_3, | \eta_X, \eta_Z, \eta_Y) = f(X_1 | \eta_X, \eta_Z) \times f(X_2, X_3 | \eta_X) \times f(Z_1, Z_2, Z_3 | \eta_Z) \times f(Y_1, Y_2, Y_3 | \eta_Y).$$

The factorization includes the cross-loading by having X_1 conditional on η_X and η_Z . In terms of a regression, this maps onto a simple linear regression with two predictors, η_X and η_Z : $f(X_1 | \eta_X, \eta_Z) \rightarrow x_{1i} = v_1 + \lambda_{11}\eta_{Xi} + \lambda_{12}\eta_{Zi} + r_{1i}$.

Turning now to the structural density, I apply the partial factorization in the same manner as in Equation 3, resulting in a conditional and joint distribution for η_X and η_Z .

$$f(\eta_X, \eta_Z, \eta_Y) = f(\eta_Y | \eta_X, \eta_Z) \times f(\eta_X, \eta_Z). \quad (13)$$

In words, Equation 13 says that the multivariate distribution of the latent variables can be expressed as the product of a conditional distribution for the endogenous latent variable, η_Y , and a bivariate distribution for the exogenous latent variables, η_X and η_Z . Similar to before, this factorization stems from the path diagram. In Figure 2, η_Y is predicted by the two latent variables and their product. Therefore, I model $f(\eta_Y | \eta_X, \eta_Z)$ via a moderated linear regression model.

$$f(\eta_Y | \eta_X, \eta_Z) \rightarrow \eta_{Yi} = \alpha_Y + \beta_1\eta_{Zi} + \beta_2\eta_{Xi} + \beta_3(\eta_{Xi} \times \eta_{Zi}) + \zeta_{Yi}. \quad (14)$$

Notably, the regression in Equation 14 includes the product of the two latent variables. This product is not a unique variable but a function of η_X and η_Z . Finally, the bivariate distribution for η_X and η_Z could be further factored into a conditional and marginal distribution, as I do in the next section.

Although I focus on interaction effects in this article, it is important to note that the factorization procedure readily accommodates other types of nonlinear associations. For example, suppose that the moderated regression from Equation 14 also included a quadratic effect of η_X . The regression model would be as follows:

$$f(\eta_Y | \eta_X, \eta_Z) \rightarrow \eta_{Yi} = \alpha_Y + \beta_1\eta_{Xi} + \beta_2\eta_{Xi}^2 + \beta_3\eta_{Zi} + \beta_4(\eta_{Xi} \times \eta_{Zi}) + \zeta_{Yi}. \quad (15)$$

To reiterate, the squared term is not a unique variable but a function of η_X . Similarly, a latent factor could exert a quadratic effect on any of its indicators. By allowing any latent variable to be expressed as a nonlinear

function, the FR specification affords much greater flexibility to specify complex functions while accounting for the measurement models.

Latent $\times$ Manifest Interaction With FR

In addition to accommodating latent by latent interactions, FRs readily extend to include manifest predictors (Keller, 2022). Reconsider the latent moderation example but with a single manifest predictor, denoted as X . Figure 3 depicts the path diagram for this model. Note I have opted to specify the relationship between X and η_Z via directed path (as opposed to correlation) because X is manifest, so this specification would be analogous to a multiple indicator multiple cause model.¹

To illustrate, I first provide the factorization for the model in Figure 3. I apply the chain rule and simplify the densities to obtain the factorization of the joint density.

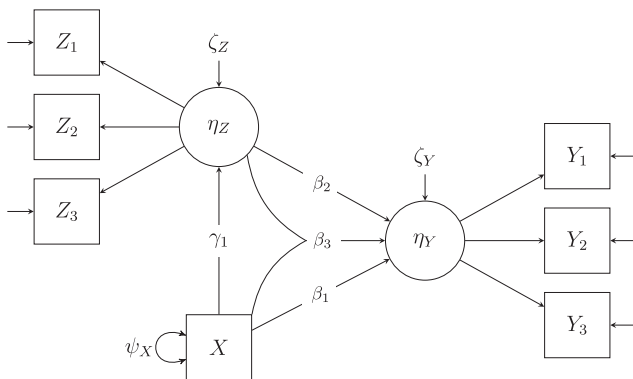
$$\begin{aligned} f(X, Z_1, Z_2, Z_3, Y_1, Y_2, Y_3, \eta_Z, \eta_Y) &= f(\eta_Y | \eta_Z, X) \times f(\eta_Z | X) \\ &\quad \times f(X) \times f(Z_1, Z_2, Z_3 | \eta_Z) \\ &\quad \times f(Y_1, Y_2, Y_3 | \eta_Y). \end{aligned} \quad (16)$$

The above factorization maps onto the previous section with only minor modifications. In words, Equation 16 says that there are three structural densities (the first three distributions) and two measurement densities (the last two distributions). Focusing on the two measurement densities, the two sets of indicators can all be expressed in the same way as the previous section because they are conditionally independent given their respective latent variable, and there are no cross-loadings.

Turning to the structural densities, I use the full factorization where the bivariate distribution of the predictors is represented as a model with η_Z conditional on X and a marginal model for X . This specification maps directly onto the path diagram, where X predicts η_Z . These structural densities result in the following regression models:

$$\begin{aligned} f(\eta_Y | \eta_Z, X) &\rightarrow \eta_{Yi} = \alpha_Y + \beta_1 \eta_{Zi} + \beta_2 x_i + \beta_3 (\eta_{Zi} \times x_i) + \zeta_{Yi}, \\ f(\eta_Z | X) &\rightarrow \eta_{Zi} = \alpha_Z + \gamma_1 x_i + \zeta_{Zi}, \\ f(X) &\rightarrow x_i = \alpha_X + r_{Xi}. \end{aligned} \quad (17)$$

Figure 3
Path Diagram for Latent $\times$ Manifest Interaction Model



Note. The mean structure is not illustrated in the diagram. The conjoined arrow for the β_3 path represents the product between X and η_Z .

Notably, the first regression in Equation 17 includes the product of the latent η_Z and manifest X . Again, this product is not a unique variable but a function of the two variables. The second regression above includes the effect of X predicting η_Z because I used the fully factored specification. Finally, the marginal model for X is an intercept-only model. If X is complete, this model may not be necessary because X could be treated as fixed (i.e., only show up right of the vertical bars). However, if X is incomplete, this model is required to handle the missing data on X appropriately.

The manifest by latent factorization readily incorporates binary, ordinal, multicategorical, and even skewed manifest variables. To illustrate one such extension, now suppose that X is a binary variable. The pathways where X predicts η_Y and η_Z representing a mean shift between the two categories, and X also moderates the effect of η_Z on η_Y . I depict this model in Figure 4. As I did previously, I use the probit specification to include the underlying distributed latent propensity X^* . Notice that the manifest X variable now has a latent propensity that predicts the binary X , and the binary X appears in the structural model.

By using this probit model, the factorization in Equation 16 only requires a slight modification of the marginal density for X to incorporate the latent propensity.

$$f(X) = f(X | X^*) \times f(X^*). \quad (18)$$

In words, the distribution of X is factored into two separate models, a conditional density of X given the latent propensity X^* and a marginal density for X^* . These two densities result in the following models:

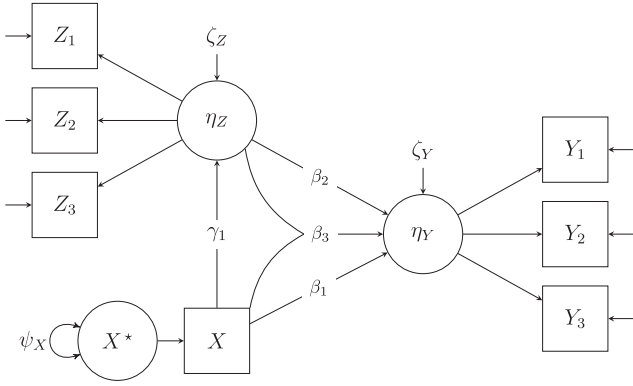
$$\begin{aligned} f(X | X^*) &\rightarrow x_i = \mathcal{I}(x_i^* > 0), \\ f(X^*) &\rightarrow x_i^* = \alpha_{X^*} + r_{X^*i}. \end{aligned} \quad (19)$$

The first line of Equation 19 is the additional link function that relates the observed binary variable to the latent propensity. In Figure 4, the link function is the arrow connecting the X^* to X . The second line of the equation characterizes the exogenous normally distributed latent propensity with a mean that relates to the proportion of ones to zeros. Returning to Figure 4, notice that the binary X remains in the model and still predicts η_Z and η_Y . Therefore, the first two lines of Equation 17 remain unchanged, providing a similar interpretation but with a binary predictor. For example, the interpretation of β_2 and β_3 provide the latent group difference in the intercept and slope of the regression.

As previously mentioned, the FR specification accommodates other types of latent by manifest variable interactions. The Blimp software described later allows for continuous, binary, ordinal, multicategorical, and skewed continuous manifest predictor variables. The factorizations for these variables mimic the one from Equation 16 but change depending on the metric of the predictor. For example, if X was a multicategorical nominal variable, the $f(X^*)$ term in the second line of Equation 19 corresponds to an empty multinomial logistic model, and the link function on the first line is a logit link. A set of

¹ In Blimp, the partial factorization from Equation 13 is also available for binary, ordinal, and normal manifest variables; however, linking exogenous variables with directed pathways accommodates a broader range of predictor metrics.

Figure 4
Path Diagram for Latent $\times$ Binary Interaction Model



Note. The mean structure is not illustrated in the diagram. The conjoined arrow for the β_3 path represents the product between X and η_Z . For identification, ψ_X is generally fixed to one.

dummy codes would appear on the right side of any equation where the discrete X is a predictor.

Bayesian Estimation of FR

Estimating FR models with MCMC methods allows for sampling from the posterior distribution of the model conditioned on the data. Accumulating many samples from the posterior distribution provides useful summaries corresponding to quantities analogous to point estimates (i.e., mean/median) and 95% interval estimates (i.e., 2.5% and 97.5% quantiles of the distribution). These MCMC methods are iterative by nature, so each iteration produces one sample from each unknown parameter's posterior distribution. This section describes the MCMC sampling steps for FR to provide the technical details of Blimp's MCMC steps. Readers not interested in these details can skip to the "Illustrative Analysis Examples Via Blimp" section, which provides a data analysis example using Blimp.

Although one can use different MCMC methods to estimate models, I have found Gibbs sampling (Geman & Geman, 1984) to provide efficient and reliable convergence. Thus, this section provides a conceptual overview of the Gibbs sampler used to estimate the models. The Gibbs sampler divides unknown quantities into blocks. Then, each block is updated in a sequence, holding all other quantities constant at their current values. The updating process for each unknown quantity involves sampling estimates from a distribution of plausible values (a conditional posterior, or a posterior predictive distribution in the case of missing data). For the FR models, MCMC estimation can be reduced down to two fundamental blocks: (1) latent variable block and (2) manifest variable block.

To illustrate, consider the entire factorization from the latent moderation model in Figure 2.

$$\begin{aligned} f(\mathbf{X}, \mathbf{Z}, \mathbf{Y}, \eta_X, \eta_Z, \eta_Y) &= f(\mathbf{X}, \mathbf{Z}, \mathbf{Y} \mid \eta_X, \eta_Z, \eta_Y) \times f(\eta_X, \eta_Z, \eta_Y) \\ &= f(\mathbf{X} \mid \eta_X) \times f(\mathbf{Z} \mid \eta_Z) \times f(\mathbf{Y} \mid \eta_Y) \\ &\quad \times f(\eta_Y \mid \eta_X, \eta_Z) \times f(\eta_X, \eta_Z). \end{aligned} \quad (20)$$

Note I use the boldface variable to denote the set of all items that load onto a single factor—for example, $\mathbf{X} = \{X_1, X_2, X_3\}$. These sets are simply a shorthand notation for the collection of univariate conditional distributions described earlier—for example, $f(\mathbf{X} \mid \eta_X) = f(X_1 \mid \eta_X) \times f(X_2 \mid \eta_X) \times f(X_3 \mid \eta_X)$. The five densities above are divided into the following blocks:

1. Latent variable block:
 - a. $f(\eta_Y \mid \eta_X, \eta_Z)$,
 - b. $f(\eta_X, \eta_Z)$.
2. Manifest variable block:
 - a. $f(\mathbf{X} \mid \eta_X) = f(X_1 \mid \eta_X) \times f(X_2 \mid \eta_X) \times f(X_3 \mid \eta_X)$,
 - b. $f(\mathbf{Y} \mid \eta_Y) = f(Y_1 \mid \eta_Y) \times f(Y_2 \mid \eta_Y) \times f(Y_3 \mid \eta_Y)$,
 - c. $f(\mathbf{Z} \mid \eta_Z) = f(Z_1 \mid \eta_Z) \times f(Z_2 \mid \eta_Z) \times f(Z_3 \mid \eta_Z)$.

There are two fundamental steps within each density, regardless of whether the variable is latent or manifest: (1) sampling the estimated parameters for the corresponding regression model and (2) imputing the missing observations.

For Step 1, the specific parameters depend on the type of model employed. For example, the $f(X_1 \mid \eta_X)$ density could be a linear regression with a normally distributed residual (e.g., see Equation 7), a probit regression for categorical data, or a regression on the transformed metric if nonnormal (Keller & Enders, 2023; Lüdtke et al., 2020b; Yeo & Johnson, 2000). To illustrate, suppose X_1 is a continuous variable that is conditionally normal given the latent variable, η_X .

$$\begin{aligned} x_{1i} &= \alpha_1 + \lambda_1 \eta_{Xi} + e_{1i}, \\ e_{1i} &\sim \mathcal{N}(0, \sigma_{e1}^2). \end{aligned} \quad (21)$$

Conceptually, constructing a Gibbs sampler involves drawing parameters conditionally on each other: (a) $f(\alpha_1, \lambda_1 \mid \sigma_{e1}^2, \eta_X, X_1)$ and (b) $f(\sigma_{e1}^2 \mid \alpha_1, \lambda_1, \eta_X, X_1)$, where the two densities correspond to a multivariate normal and inverse gamma distribution. The sampling of the parameters (i.e., α_1, λ_1 , and σ_{e1}^2) conditional on the manifest and latent data follow standard procedures for linear regression models with conjugate priors (e.g., see Gelman, Stern, et al., 2013; Keller & Enders, 2023; Lynch, 2007). This simplification follows because the imputed values of η_X function as observed/known quantities.

As previously mentioned, the sampling steps for regression models hold true regardless if the variable is latent or manifest. To illustrate, consider drawing a structural regression corresponding to the model in Block (1a) above.

$$\begin{aligned} \eta_{Yi} &= \alpha_Y + \beta_1 \eta_{Zi} + \beta_2 x_i + \beta_3 (\eta_{Zi} \times x_i) + \zeta_{Yi}, \\ \zeta_{Yi} &\sim \mathcal{N}(0, \psi_Y^2). \end{aligned} \quad (22)$$

Again, treating the current latent imputations as known data allows one to leverage standard Gibbs sampling steps for multiple regression models because the latent data have been imputed and function as observed/known quantities. For example, the conditional posterior for the structural β 's with a conjugate prior is a multivariate normal distribution with the ordinary least squares point estimates and sampling covariance matrix as its mean vector and covariance matrix, respectively.

Imputation of Variables

As previously discussed, conceptually, there is no difference between a single observation of a latent variable and an unobserved observation of a manifest variable. Both are estimated via imputation—that is, data augmentation (Tanner & Wong, 1987). The Gibbs sampler obtains imputations from the conditional distribution of the variable to be imputed given all other variables and parameters in the model. The one caveat is that to recover the parameters associated with the incomplete manifest variable without explicitly modeling the selection mechanism; one must assume that the data are MAR and that the parameters associated with the selection process are distinct from the rest of the analysis model (Rubin, 1976). If this assumption is not tenable, one would need to model this selection model directly as an additional FR model.

To illustrate, reconsider the latent moderation model in Figure 2. Suppose the manifest variable X_1 has incomplete observations that must be imputed. I determine the conditional distribution by simplifying Equation 20 to only densities that contain X_1 .

$$f(X_1 | X_2, X_3, \mathbf{Y}, \mathbf{Z}, \eta_X, \eta_Y, \eta_Z) = f(X_1 | \eta_X). \quad (23)$$

In this instance, there is only one density, and the problem simplifies nicely. If I know what η_X is, I can impute X_1 conditional on this value. I previously characterized $f(X_1 | \eta_X)$ in Equation 21 as a linear regression in the example; therefore, the conditional distribution of X_1 is as follows:

$$x_{1i} \sim \mathcal{N}(\alpha_1 + \lambda_1 \eta_{X_i}, \sigma_{\epsilon 1}^2). \quad (24)$$

The mean and variance of the above distribution are defined by the predicted score and residual variance of Equation 21, respectively. Like the latent scores, the parameter values that define the distribution's center and spread are held at their current values.

Generally speaking, obtaining the imputations of the latent variables is more complicated because they appear in more than one density. The latent moderation example has three latent variables that need to be imputed. The conditional density for any variable can be expressed up to proportionality by returning to Equation 20 and removing any density that does not contain the specific variable. For example, η_X 's conditional density is expressed up to proportionality as follows:

$$f(\eta_X | \mathbf{X}, \mathbf{Z}, \mathbf{Y}, \eta_Y, \eta_Z) \propto f(\eta_Y | \eta_X, \eta_Z) \times f(\eta_X, \eta_Z) \times f(\mathbf{X} | \eta_X). \quad (25)$$

Importantly, the density $f(\eta_Y | \eta_X, \eta_Z)$ maintains the nonlinear relationship between η_X and η_Y by including the $\eta_X \times \eta_Z$ product (i.e., β_3 in Equation 14). Similarly, the imputations of η_Z and η_Y can be obtained by sampling from their conditional densities.

$$f(\eta_Z | \mathbf{X}, \mathbf{Z}, \mathbf{Y}, \eta_X, \eta_Y) \propto f(\eta_Y | \eta_X, \eta_Z) \times f(\eta_X, \eta_Z) \times f(\mathbf{Z} | \eta_Z), \quad (26)$$

$$f(\eta_Y | \mathbf{X}, \mathbf{Z}, \mathbf{Y}, \eta_X, \eta_Z) \propto f(\eta_Y | \eta_X, \eta_Z) \times f(\mathbf{Y} | \eta_Y). \quad (27)$$

Under the assumption of normally distributed residuals, all three conditional latent densities above have closed-form solutions. Each latent variable follows a normal distribution, but η_X and η_Z are heteroscedastic conditional on η_Y .

The technical Appendix A describes the recursive algorithmic approach used by Blimp to obtain the solutions for the mean and variance of the latent variable distributions. In addition, I include an R implementation at <https://github.com/blimp-stats/Latent-Interactions> to demonstrate how the closed-form solutions are obtained. Although beyond the scope of this article, this approach readily extends to a wide range of models, including interactions, multivariate models, and multilevel models.

Alternatively, when the conditional distribution cannot be solved, I can leverage various general MCMC sampling techniques that require knowing the density only up to proportionality and embed these within the Gibbs sampler. For example, I can employ a random walk Metropolis step to sample from the density (Lynch, 2007). Such techniques are necessary when an incomplete or latent variable is squared or a categorical mediator is present.

Prior Distributions

In this section, I briefly discuss the default priors used by Blimp. For more specific technical details about Blimp's default prior distributions in the online supplemental materials from Enders et al. (2020) and Keller and Enders (2023),² as well as the user manual (Keller & Enders, 2022). Generally speaking, Blimp provides uninformative or weakly informative priors by default (when possible), and the priors for regression coefficients (i.e., structural coefficients or factor loadings), thresholds, and variances (latent or manifest) are all considered independent. Note, informative priors are possible via the `PARAMETERS` command.

Returning to Figure 2, each path coefficient (i.e., loading or regression slope) and intercept would invoke an improper uniform prior; however, Blimp can incorporate other priors that are user specified. For path coefficients, if a univariate normal distribution prior is requested, then Blimp will sample the coefficient from a full conditional distribution. For factor loadings, if the variance of a latent variable is fixed at one for identification, then a truncated prior distribution over the positive range is usually desirable. Otherwise, the posterior distributions of the loadings may be bimodal (Levy & Mislavy, 2017).

For unstructured covariance matrices, they are given a weakly informative inverse Wishart prior. For structured covariance matrices, the prior largely depends on the nature of the matrix. For example, if only a variance is fixed, then I maintain the inverse Wishart prior and use a transformation to maintain the fixed variance (Asparouhov & Muthén, 2010a; Lynch, 2007). Alternatively, if covariances are fixed, I place a uniform prior on the elements and use a Metropolis–Hasting algorithm to draw each element conditionally on all the others. Separation type specifications that place distinct priors on variances and correlations are also possible in Blimp (Keller & Enders, 2023). Finally, residual variances are given inverse gamma priors that map onto the univariate special case of the inverse Wishart priors.

Turning to regressions of categorical outcomes, Blimp imposes a default normal prior with a mean of zero and a variance of five. In my experience, this weakly informative prior is useful for stabilizing the estimates where the regression coefficient ranges are more limited due to scaling. Finally, for ordinal probit models, the thresholds are drawn according to Cowles (1996).

² Both the online supplemental materials are available at <https://www.appliedmissingdata.com/blimp-papers>.

Monte Carlo Simulation Studies

This section outlines four computer simulation that investigate the FR approach to latent interactions:

1. Latent by latent interaction with normally distributed indicators (Figure 2).
2. Latent by manifest interaction with normally distributed indicators and moderator (Figure 3).
3. Latent by binary interaction with normally distributed indicators (Figure 4).
4. Latent by binary interaction with ordinal indicators (Figure 4 with ordinal indicators).

In all four simulations, I compared the FR specification to both LMS and PI (estimated via ML). The manipulated conditions were sample size ($N = 250, 500, \text{ and } 1,000$), the number of indicators per latent factor (five or 10), and missing data rate (0% or 35%). Each condition was replicated 2,500 times. All computer simulation code is available at <https://osf.io/2ghny> (Keller, 2024c).

Data Generation

Data were generated in R (Version 4.2.2; R Core Team, 2022) under the diagrams specific to each simulation (i.e., Figures 2–4). The population parameters, along with item proportions for Simulation 4, were selected based on the summary statistics and effect sizes reported by Zhou et al. (2014), the motivating example I use in the next section to demonstrate estimating the models in Blimp. In all simulations, I focus solely on a single two-way interaction, with the proportion of variance accounted for by the three variables equal to 0.348. Figure 5 illustrates the conditional regression equations expressed by the population parameters (i.e., $\beta_1 = -.25$, $\beta_2 = .50$, and $\beta_3 = -.25$). Because the data are generated with a standard normal or binary moderator, for Simulations 1 and 2, the three lines in Figure 5 correspond to -1 SD below the mean of the moderator, at the moderator mean, and $+1$ SD above the

moderator mean. For Simulation 3, the difference between the two groups is equivalent to two lines $X = 0$ and $X = +1$.

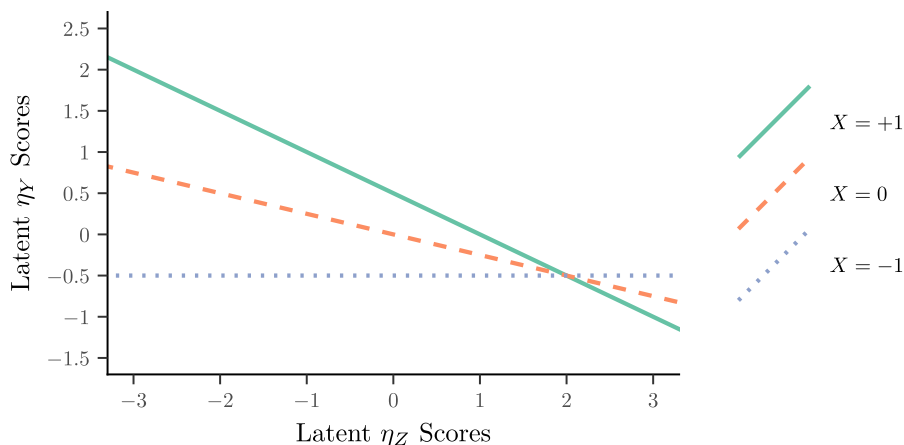
Turning to the measurement model, for Simulations 1–3, the indicators were generated using normally distributed residuals. Each indicator was generated with equal loadings and the latent variable accounted for 50% of a specific indicator's variance (i.e., a standardized loading of ~ 0.707). For Simulation 4, the ordinal indicators were generated by categorizing the normally distributed indicators based on thresholds to generate an ordered 5-point Likert scale (i.e., a probit response model). The ordinal indicators were generated to replicate the reported means and standard deviations of scale scores from Zhou et al. (2014, Table 1), leading to asymmetric indicator distributions.

For missing data conditions, missing values were induced on the manifest variables via a selection model. For Simulation 1, missing data was induced on each individual indicator for η_X based on the corresponding indicator for η_Y . For example, in Figure 2, X_1 would have missing values induced based on the values of Y_1 . The relationship between each indicator and the propensity to be missing was set to account for 60% of the propensity's variance. Because Simulations 2–4 do not have items on the manifest moderator, a different approach was required. To ensure a MAR-based mechanism, items for the endogenous latent variable (i.e., η_Y in Figures 3 and 4) were averaged and used to predict the propensity to be missing on the manifest variable (i.e., X in Figures 3 and 4). Similar to Simulation 1, the relationship between this mean score and the propensity to be missing was set to account for 60% of the propensity's variance.

Model Estimation

For each simulation, I estimated the model with FR via MCMC with Blimp, LMS via ML with Mplus (Version 8.8; L. Muthén & Muthén, 2017), and PI via ML with Mplus. The models were identified by fixing the residual variance of the latent variables to one (i.e., fixed factor scaling; Klopp & Klöbner, 2021). All estimation procedures were automated and read into R via packages.

Figure 5
Population Conditional Effects Plot for Latent Interaction



Note. The regression of latent η_Y on latent η_Z moderated at different levels of X based on the data generation. For Simulation 1, X is treated as a latent unobserved construct with observed ordinal indicators. In Simulations 2 and 3, X is partially observed. See the online article for the color version of this figure.

For Blimp, I used the `rblimp` package (Keller, 2024b) and Mplus I used the `MplusAutomation` package (Hallquist & Wiley, 2018). As a reminder, all simulation code is provided at <https://osf.io/2ghny> (Keller, 2024c).

For the MCMC method, I implemented the split-chain potential scale reduction factor diagnostic (Gelman & Rubin, 1992) and effective sample size (Gelman, Carlin, et al., 2013; Geyer, 2011) to assess the convergence of each replication. I ran pilot simulations to assess the number of burn-in and postburn-in iterations required to achieve convergence. 10,000 burn-in iterations across two chains with 20,000 postburn-in would provide potential scale reduction factor statistics generally below 1.05 (no more than 1.10) and effective sample size of at least 200 (100 per chain; Gelman, Carlin, et al., 2013). For Simulations 2 and 3, it was determined that 20,000 burn-in iterations across two chains with 40,000 postburn-in were required. Convergence was also assessed on the individual simulation replications.

The current implementation of LMS requires special considerations for Simulations 2–4. For Simulation 2, LMS requires a proxy latent variable to model the latent by manifest interaction (Asparouhov & Muthén, 2021b). This proxy variable specifies a latent factor with a factor mean, a loading of 1, and a sole indicator of the manifest variable (i.e., X). The intercept of X must be fixed to zero to estimate the factor mean. Similarly, the variance of X must then be fixed to zero so that the proxy latent variable is analogous to the manifest predictor. In Simulations 3 and 4, the LMS method cannot directly model the covariance between the latent and binary predictors. Therefore, the binary predictor was regressed on the latent predictor to keep the latent predictor itself as an exogenous variable with a variance of 1. Across all conditions of all four simulations, there were only a total of 12 failed replications (Simulation 1 = 0, Simulation 2 = 33, Simulation 3 = 0, and Simulation 4 = 12). Because of the low percentage of failed replications, these replications were dropped for every estimator.

The simulations with normally distributed indicators deploy conditions that connect with previous literature. I also implemented additional simulations using ordinal indicators (i.e., Simulation 4 and two additional simulations in the online additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY> that mimic Simulations 1 and 2). These simulations were restricted to have missing data, as the goal was to stress test the FR method under a set of extreme (but realistic) conditions. The FR and LMS models used a probit link for the discrete indicators. In addition, LMS required Monte Carlo numerical integration, and I used the default setting provided by Mplus to map onto what most researchers choose. Although it is not a direct apples-to-apples comparison, I also implemented normal-theory PI because there is no way to incorporate a categorical model with PIs. The fixed-factor scaling still allows for comparison because the latent variables are standardized.

Results

For each simulation, I focused on evaluating proportional bias and the ratios of mean square error (MSE) for the path coefficients,³ specifically targeting first-order and interaction regression coefficients. A summary of the key findings from each simulation is provided below. Comprehensive graphical summaries of all results, along with the analysis code, pooled data, and raw outputs, can be accessed through the online repository as the additional

online materials at <https://osf.io/2ghny>. Additionally, although not explicitly reported, I computed jackknife and bootstrap standard errors for the evaluation measures to estimate the Monte Carlo sampling variability (Koehler et al., 2009). These calculations confirmed that 2,500 replications per condition were sufficient to ensure minimal Monte Carlo variability.

Simulation 1: Latent by Latent Interaction With Normally Distributed Indicators

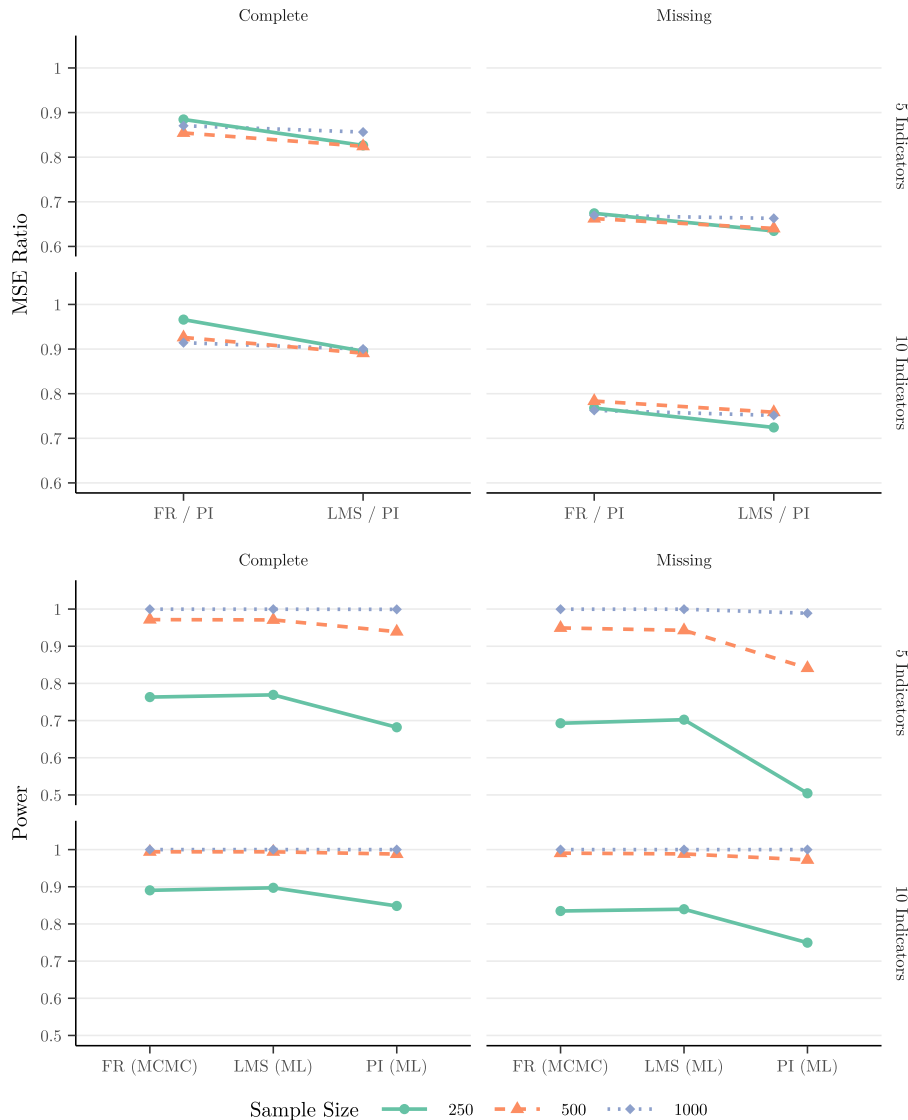
Simulation 1 evaluated latent interaction models with normally distributed indicator variables. A full graphical summary of the simulation results is available in the additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY>. Overall, the relative bias values across the three estimation methods were mostly unremarkable, with most estimates exhibiting absolute biases below 5%. There were noticeable differences in the MSE ratio and power estimates. As a reminder, MSE ratios <1 reflect an accuracy advantage for either the FR or LMS relative to the PI approach, the MSE for which is in the denominator of the ratio. With complete data, most MSE ratios for the lower order slopes were between 0.95 and 1.0, reflecting only a modest advantage for FR and LMS. In contrast, the MSE ratios for the interaction coefficient were more pronounced, ranging between 0.80 and 0.90 with five indicators and between 0.90 and 0.95 with 10 indicators. Figure 6 displays the ratios for the interaction slope. Because the MSE is inversely related to sample size, the ratios can be interpreted as the proportional amount of data needed for FR and LMS to achieve the same accuracy as the PI method. For example, a value of 0.85 means that FR or LMS can achieve the same overall accuracy using ~85% as much data. As you might expect, missing values on the moderator factor's (η_X) indicators increased the accuracy differences. For example, the MSE ratios for the interaction coefficient ranged between 0.60 and 0.70 with five indicators and between 0.70 and 0.80 with 10 indicators. There were no material differences between FR and LMS. Figure 6 also displays the power estimates for the same subset of conditions. The MSE ratios translated into power reductions between 5% and 20% for the PI method. FR and LMS power estimates were virtually identical.

Simulation 2: Normally Distributed Manifest Moderator

Simulation 2 evaluated latent interaction models with normally distributed indicator variables for η_Z and a normally distributed manifest moderator X . A full graphical summary of the simulation results is available in the additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY>. Overall, the relative bias values from the complete-data simulation approached zero for all three methods. However, with missing data, the PI's interaction slope exhibited about +10% bias, and the manifest moderator's slope exhibited about -10% bias. MSE ratios from the second simulation closely mimicked those from Figure 6, so no further displays are necessary. The only difference was that MSE ratios slightly favored FR over LMS. Presumably, this is because LMS requires a proxy latent variable with a single indicator whereas the FR model does not. In this case, the modest MSE differences did not translate into

³ Posterior medians were evaluated for the Bayesian FR procedure.

Figure 6
Simulation 1 Power and MSE Ratio for β_3 Coefficient



Note. *MSE* = mean square error; FR = factored regression estimation via MCMC; PI = product indicator estimation via ML; LMS = latent moderated structural equations estimation via ML; MCMC = Markov chain Monte Carlo; ML = maximum likelihood. The top panel is *MSE* ratio and the bottom panel is power. See the online article for the color version of this figure.

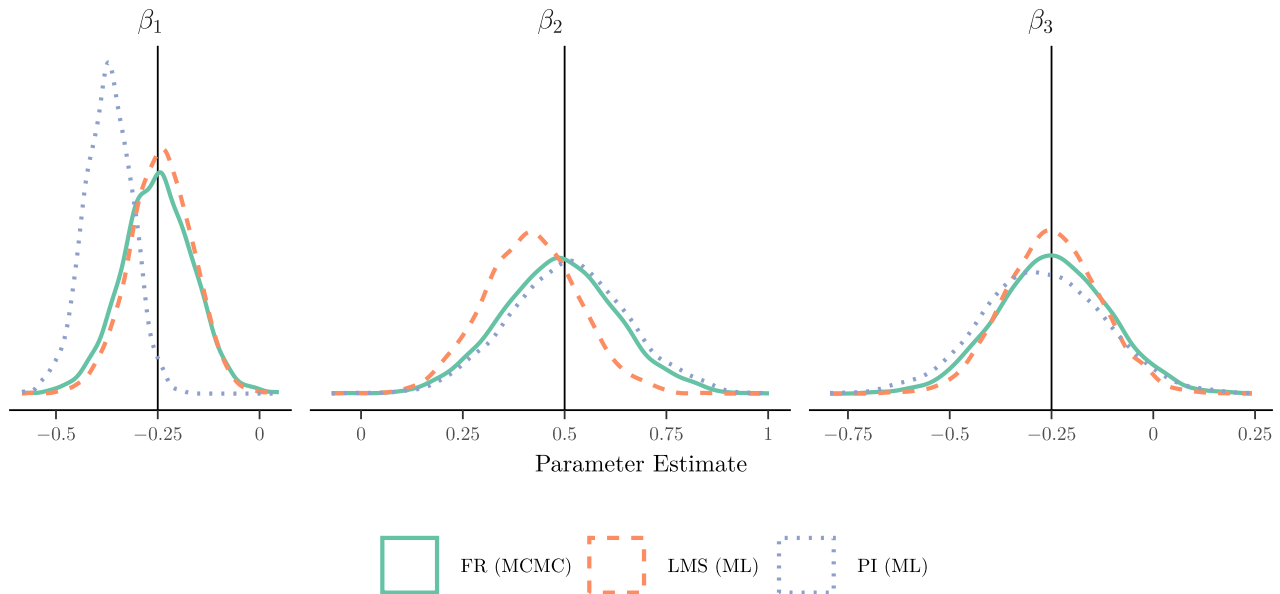
meaningful power differences (LMS estimates were about 2%–3% lower than FR).

Simulation 3: Binary Moderator With Normally Distributed Indicators

Simulation 3 evaluated latent interaction models with normally distributed indicator variables for η_Z and a binary moderator X . A full graphical summary of the simulation results is available in the additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY>. With complete data, FR's MCMC estimator and LMS ML-based estimator produced bias values close to zero.

The PI method produced substantially biased estimates of the lower-order slope for the continuous latent predictor ($\sim 50\%$ lower than the true value). Because the other coefficients were unbiased, this means that the simple slopes of both groups were dramatically distorted. With missing data, FR estimates of the structural coefficients remained unbiased, but LMS produced biased estimates of binary moderator's main effect, with stable bias values of $\sim -17\%$ across other conditions. To illustrate, Figure 7 shows the kernel density plots of the structural coefficients from the 10-item condition with $N = 500$. Consistent with the complete-data results, PI produced substantially biased estimates of the latent variable's slope. The middle row panel illustrates the negative bias in the LMS estimates of the

Figure 7
Simulation 3 Density Plots of Structural Coefficients Across All Replications



Note. FR = factored regression estimation via MCMC; MCMC = Markov chain Monte Carlo; LMS = latent moderated structural equations estimation via ML; ML = maximum likelihood; PI = product indicator estimation via ML. See the online article for the color version of this figure.

binary moderator's slope. Because this coefficient is the simple intercept for the group coded 1 (the latent mean is fixed at zero), the negative bias implies that the vertical separation of the group-specific conditional regression lines is too small by about 10% of a standard deviation unit.

As shown in Figure 7, the distribution of the FR interaction coefficient was noticeably wider than that of LMS. Figure 8 displays the corresponding *MSE* ratios for the interaction slope. With incomplete data, these precision differences translated into *MSE* ratios that favored LMS (e.g., 0.60 for LMS vs. 0.80 for FR). However, the bottom row panel of the figure illustrates that these precision differences did not translate into meaningful power disparities.

Simulation 4: Binary Moderator With Ordinal Indicators

Simulation 4 evaluated latent interaction models with five-point ordinal indicator variables for η_Z and a binary moderator X . These simulations were limited to missing data, as the goal was to thoroughly stress test FR. As a reminder, PI treated the ordinal items as normally distributed (they were substantially asymmetric), as there is no meaningful way to apply a probit link to PIs. A full graphical summary of the simulation results is available in the additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY>. The additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY> also summarize versions of Simulations 1 and 2 with ordinal indicators.

FR estimates were effectively unbiased. Both LMS and PI exhibited biases. To illustrate, Figure 9 shows the kernel density plots of the structural coefficients from the 10-item condition with $N = 500$. As seen in the figure, PI produced negatively biased estimates of both lower order terms but not the interaction. In practical terms, this bias implies that the simple slopes for both groups

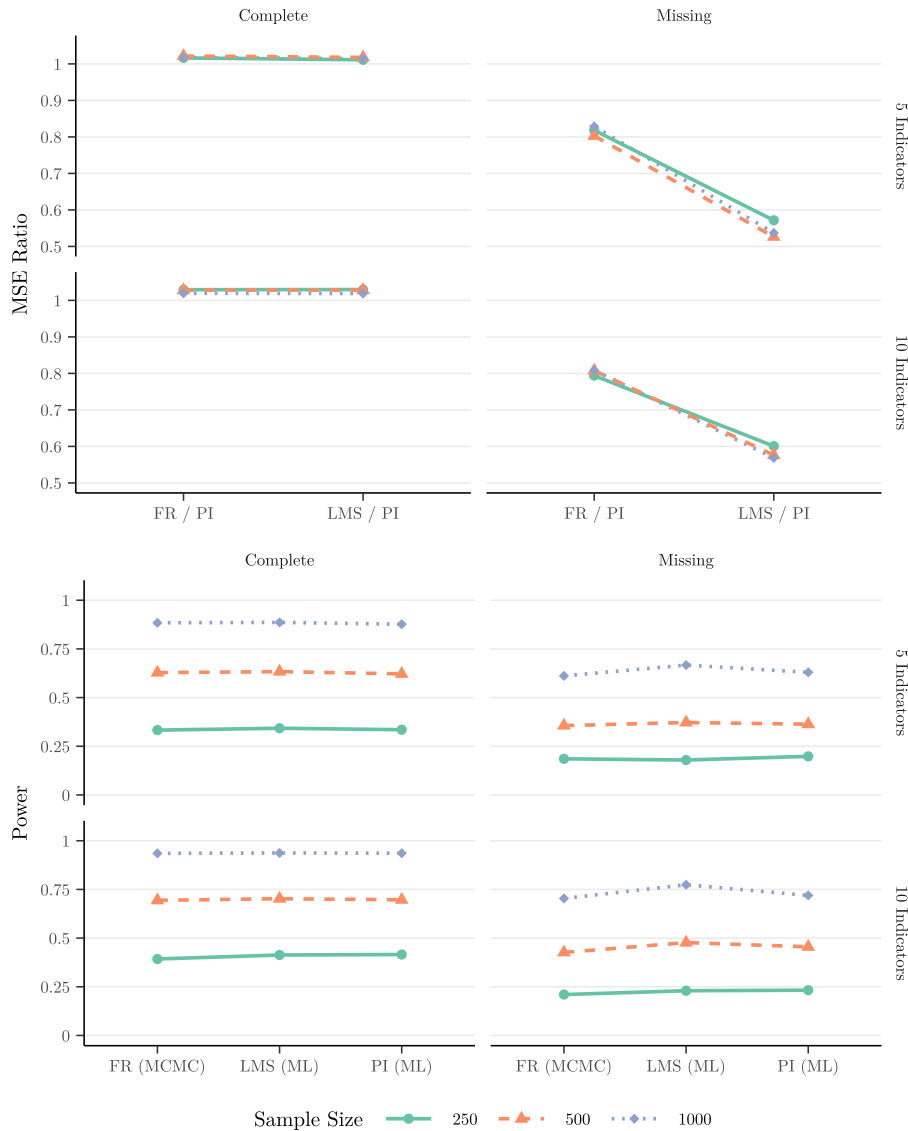
were too steep, and the vertical separation of the regression lines was too small. In contrast, LMS produced negatively biased estimates of the moderator effect and positively biased estimates of the interaction slope. The simple slope for the group coded 1 was too flat, and the vertical separation of the regression lines was too small. *MSE* ratios indicated that there was a bias-precision tradeoff. For the latent variable's slope (β_1), FR and LMS produced similar *MSE* ratios. For the moderator's slope (β_2), ratio differences were limited to the five-indicator condition, where LMS had somewhat better accuracy. Finally, FR estimates of the interaction coefficient were generally more accurate than LMS, except at the smallest sample size, where LMS had a slight advantage. For brevity, graphical summaries of *MSE* ratios are available in the additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY>.

Interpreting power differences is questionable in the presence of bias. For example, FR-based tests of the interaction effect generally had noticeably higher power than LMS, but that is likely idiosyncratic feature of the simulation design (e.g., had the population product slope been positive, positive bias would artificially inflate rather than deflate power).

Illustrative Analysis Examples Via Blimp

As previously mentioned, I provide an implementation of the FR modeling framework in Blimp (Version 3.2; Keller & Enders, 2022). Blimp is a free standalone software with syntax similarities to many modern-day SEM software packages, but internally, it is built around the FR modeling approach. In this section, I use an illustrative data example to demonstrate using Blimp to investigate a latent interaction model, including using familiar pick-a-point approaches or a Johnson–Neyman plot (P. O. Johnson & Neyman, 1936) to probe

Figure 8
Simulation 3 Power and MSE Ratio for β_3 Coefficient



Note. MSE = mean square error; FR = factored regression estimation via MCMC; PI = product indicator estimation via ML; LMS = latent moderated structural equations estimation via ML; MCMC = Markov chain Monte Carlo; ML = maximum likelihood. The top panel is MSE ratio and the bottom panel is power. The vertical line represents the true parameter value. See the online article for the color version of this figure.

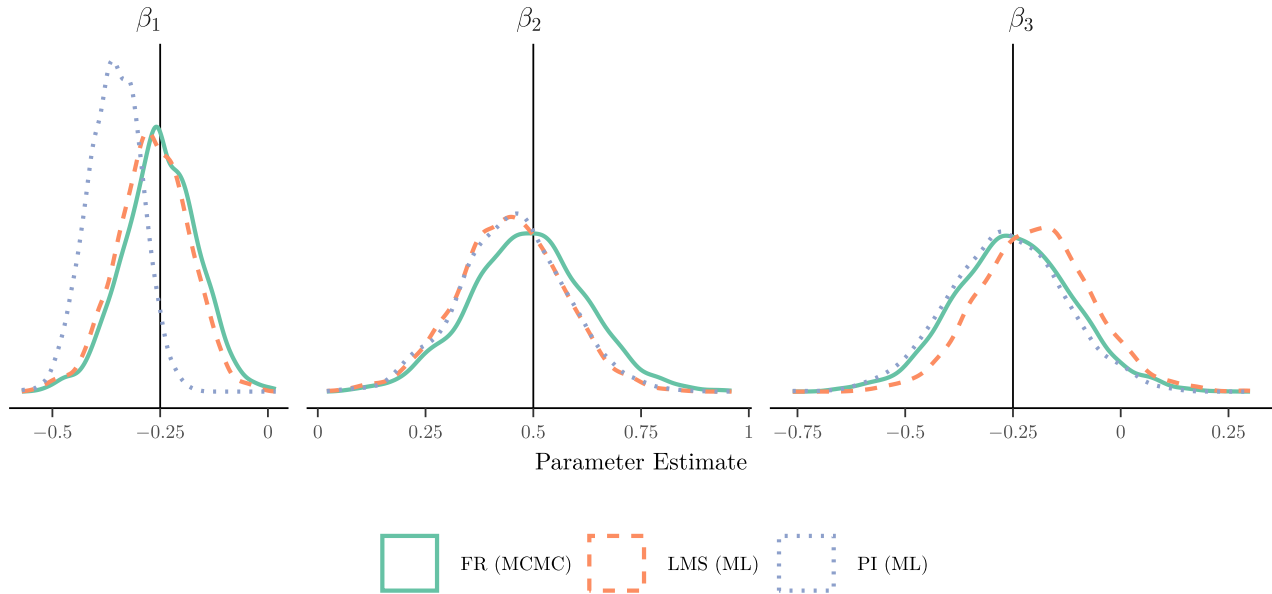
the interaction. Throughout this section, I supplement the discussion with code excerpts that show how to specify the model. All syntax files, output files, and raw data are provided at <https://github.com/blimp-stats/Latent-Interactions>. In addition, the additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY> include fitting the models using the `rblimp` package (Keller, 2024b) for R. This package allows researchers to readily generate the graphical diagnostics and conditional regression plots discussed in this section.

The substantive example is based on Zhou et al. (2014), who investigated the interactions between personality traits and

organizational constraints when predicting counterproductive work behavior. Based on the reported summary statistics, effects, and discussion of scales, I generated data to mimic theirs. I generated all items to follow a 5-point Likert scale, just as they were measured in the original article.

Initially, I focus on conscientiousness (`consci`) and organizational constraints (`orgcon`) scales consisting of 10 items each, and how they predict counterproductive work behavior directed at people (`cwb`). The conscientiousness scale measures the “Big Five” conceptualization of the personality trait, which relates to the degree of organization and self-discipline in a participant (Costa & McCrae, 1992).

Figure 9
Simulation 4 Density Plots of Structural Coefficients Across All Replications



Note. FR = factored regression estimation via MCMC; MCMC = Markov chain Monte Carlo; LMS = latent moderated structural equations estimation via ML; PI = product indicator estimation via ML; ML = maximum likelihood. The vertical line represents the true parameter value. See the online article for the color version of this figure.

The organizational constraints scale captures how difficult it is for a participant to perform their job. An example item might ask if a participant feels the organization provides poor equipment or supplies. Finally, the counterproductive work behavior scale consisted of 23 items, asking participants how often they had performed counterproductive behaviors at their jobs. The concept of counterproductive work behavior refers to intentional behaviors that harm organizations or people in organizations (Spector & Fox, 2005). An example item might ask how often they might yell at someone at work or damage an organization's property.

The initial models map directly onto the discussion of the latent moderation model (i.e., Figure 2), with latent factors of conscientiousness and constraints as η_X and η_Z , respectively, and counterproductive work behavior as the endogenous latent factor, η_Y . Because each item follows a Likert scale, I use an ordered probit model to characterize each item (Albert & Chib, 1993; V. E. Johnson & Albert, 2006). The metrics of the indicators are specified by listing the indicator names after the `ORDINAL` command. Blimp will automatically use the data to determine the number of categories and fit the appropriate ordered probit model.

Fitting a Latent Regression Model

When fitting these models, I advocate for a modeling-building approach driven by substantive theory. I begin by fitting the structural model without the interaction.

$$cwb_i = \beta_0 + \beta_1(\text{consci}_i) + \beta_2(\text{orgcon}_i) + \zeta_i. \quad (28)$$

First, I invoke the modeling section using the `MODEL` command. Each subsequent line under this command specifies a regression equation from the broader model. The following code excerpt

specifies the regression of the two latent exogenous variables onto the endogenous latent counterproductive work behavior variable. Below, the first code excerpt includes two commands (denoted as capitalized names followed by a colon):

```
## Define latent variables
LATENT: consci orgcon cwb;

## Begin Modeling Syntax
MODEL:
  ## Latent Regression Model (no moderation)
  cwb ~ consci orgcon;

  # Fix residual variance to 1 for identification
  cwb ~~ cwb@1;
```

Note that Blimp's syntax uses a hashtag or pound symbol (#) to turn the remaining line into comments and a semicolon (;) to terminate any statement. The `LATENT` command defines three new latent variables. The names themselves are arbitrary, but I use them to represent the three latent variables for conscientiousness (`consci`), organizational constraints (`orgcon`), and counterproductive work behavior (`cwb`). The `MODEL` command begins the specification of the model. The first noncomment line in the `MODEL` section specifies the regression without the interaction (Equation 28). Blimp syntax uses a tilde (~) to denote "regressed on" or "predicted by." Thus, the first equation in the `MODEL` section reads as `cwb` is predicted by `consci` and `orgcon`. Finally, the double tilde syntax (~~) is used to specify correlations or variances. In this case, I use the syntax to fix the residual variance of `cwb` to be equal to one for identification to specify a fixed factor scaling (Klopp & Klößner, 2021). I also use the double tilde syntax to specify the models for the

two exogenous factors.

```
# Correlate two exogenous variables
consci ~ orgcon;

# Fix variance to 1 for identification
consci ~ consci@1;
orgcon ~ orgcon@1;
```

Mapping onto a standard multiple regression, I correlate the two exogenous latent variables. I also set their variances to one for identification of the latent constructs.

The final section of the MODEL command is devoted to specifying the measurement models (i.e., Equation 12). I allow all slopes (i.e., loadings) to be freely estimated because I have identified the latent variables by constraining their variances. Below, I regress each item onto its specific latent variable.

```
# Measurement Models for conscientiousness items
consci -> consci1@l_con consci2:consci10;

# Measurement Models for org constraints items
orgcon -> orgcon1@l_org orgcon2:orgcon10;

# Measurement Models for CWB items
cwb -> cwb1@l_cwb cwb2:cwb23;

# Specify first loading as always positive
PARAMETERS:
  l_con ~ truncate(0, Inf);
  l_org ~ truncate(0, Inf);
  l_cwb ~ truncate(0, Inf);
```

Blimp uses a right-pointing arrow ($\rightarrow$) to specify that the variable to the left (a latent factor) predicts every variable to the right (its indicators). As a shortcut, Blimp allows for the specification of multiple variables using a colon. For example, above the “consci2:consci10” denotes a consecutive range of consci2 to consci10. As discussed in the previous section, I must identify the latent variable by setting the first loading to always be positive. I accomplish this by first specifying a parameter name for each loading using the ampersand (“&”) followed by the name of my choosing (e.g., consci1@l_con). I then used the PARAMETERS command to specify a truncated prior distribution that constrains the first loading of each factor to always be positive (Levy & Mislavy, 2017).

The Blimp script also requests 50,000 iterations after the initial warm-up period (see Appendix B for a full syntax example). Thus, each parameter is characterized by an empirical posterior distribution of 50,000 samples. In Table 1, the “Coefficient” column is the median and is analogous to a frequentist point estimate. The posterior *SD* values in the second column give the *SD* of the 50,000 parameters, which are numerically similar to frequentist *SEs*. The “Credible intervals” columns report the 95% interval, which is similar to a frequentist confidence interval, but they make no reference to repeated sampling. Said differently, the credible interval gives the 95% probability that the true estimate would lie within the interval given the observed data and model. For researchers who prefer to work with null hypothesis significance testing logic, the credible interval can be used to test such hypotheses. Specifically, suppose the null value of zero falls outside this interval. In that case, one can conclude that the effect is

Table 1
Results for Latent Regression Models

Predictor	Coefficient	SD	Credible intervals	
			2.5%	97.5%
Model 1				
consci	−0.288	0.039	−0.365	−0.213
orgcon	0.611	0.043	0.527	0.697
R ²	0.335	0.026	0.285	0.386
Model 2				
consci	−0.300	0.040	−0.379	−0.223
orgcon	0.632	0.045	0.546	0.721
consci × orgcon	−0.223	0.042	−0.307	−0.142
R ²	0.362	0.026	0.310	0.412
Model 3				
consci	−0.311	0.042	−0.396	−0.230
orgcon	0.653	0.048	0.562	0.749
emosta	−0.017	0.039	−0.094	0.060
consci × orgcon	−0.251	0.045	−0.341	−0.165
consci × emosta	−0.072	0.039	−0.148	0.004
orgcon × emosta	0.084	0.043	0.001	0.169
consci × orgcon × emosta	−0.093	0.044	−0.181	−0.008
R ²	0.377	0.026	0.325	0.427

Note. All variables are unobserved latent constructs. The “Coefficient” column is the posterior median. consci = conscientiousness construct; orgcon = organizational constraint; emosta = emotional stability.

“significant” because the probability that a parameter value of zero produced the data is $<.05$. For researchers who prefer a frequentist-like *p* value, the Bayesian Wald χ^2 tests (Asparouhov & Muthén, 2021a) are also printed in the Blimp output by default.

Under the Model 1 heading, Table 1 gives the posterior summaries for the latent regression model with no interactions. I can see that zero is not included within the 95% credible interval for both effects. Looking at the coefficient for latent conscientiousness, with the fixed factor scaling, the coefficient is interpreted as a 1 *SD* unit increase in latent conscientiousness, I would expect the latent counterproductive work behavior to decrease by 0.288 *SDs* of its residual (Klopp & Klößner, 2021, p. 193). Similarly, turning to latent organizational constraints, with a 1 *SD* unit increase in the latent organizational constraints, I would expect the latent counterproductive work behavior to increase by 0.611 *SDs* of its residual. Finally, the results suggest that the inclusion of the two latent predictors explains 33.5% of the variance of the latent counterproductive work behavior.

Graphical Diagnostics on Latent Imputations

One important practical advantage of the FR approach is that it yields latent variable imputations or “plausible values” as a byproduct of estimation. These latent variable imputations can be leveraged to produce familiar graphical diagnostics for multiple regression models. To illustrate, I used the follow code block to save 20 imputed data sets that included imputations, predicted values, and residuals for every variable in the model (latent or manifest).

```
NIMPS: 20; # Number of imputed data sets
OPTIONS: SaveLatent; # Request for saving latent imputations
         SavePredicted; # Request for saving predicted scores
         SaveResiduals; # Request for saving residuals
SAVE: stacked = m0.txt; # File to save imputations into 'm0.txt'
```

Each imputed data set will contain different imputed values. The imputations will be equally spaced across the requested post burn-in iterations and chains. Because these are imputations by nature, nothing precludes analyzing them as complete data sets and using multiple imputation pooling rules (Asparouhov & Muthén, 2010b; Mansolf, 2023).

Frequentist inference using Rubin's (1987) pooling rules is a typical application for multiply imputed data sets. One could fit the structural regression using the latent variable imputations if desired. Instead, I use the residual and predicted scores for *cwb* from the model to construct diagnostic regression plots (Cook & Weisberg, 1999). Panel A in Figure 10 plots the residuals versus predicted based on the model without the interaction and was generated using `residual_plot()` function in the `rblimp` package for R. While the plotted point represents the average values across the 20 imputations, the loess line's fitted values and ± 2 SE bars are based upon the fitting the loess line to each imputed data set and pooling the results using Rubin's rules (Rubin, 1987). The averaged data points in the plot are solely used for graphical representation in the plot and are not incorporated into the loess line estimation. Although looking at regression plots involves some subjectivity, the curvilinear form of the loess line in Panel A indicates that an unmodeled nonlinear relationship may remain in the latent regression. Paired with a substantive theory about a stressor-personality interaction (Bowling & Eschleman, 2010; Fox et al., 2001; Penney & Spector, 2005), it would make sense to model the interaction between conscientiousness (*consci*) and organizational constraints (*orgcon*).

Fitting a Latent Moderation Model

Fitting a latent moderation model like the one in Figure 2 requires only a minor change to the syntax from the Fitting a Latent Regression Model

section. Below I provide the equation for the focal structural model.

$$cwb_i = \beta_0 + \beta_1(consci_i) + \beta_2(orgcon_i) + \beta_3(consci_i \times orgcon_i) + \zeta_i. \quad (29)$$

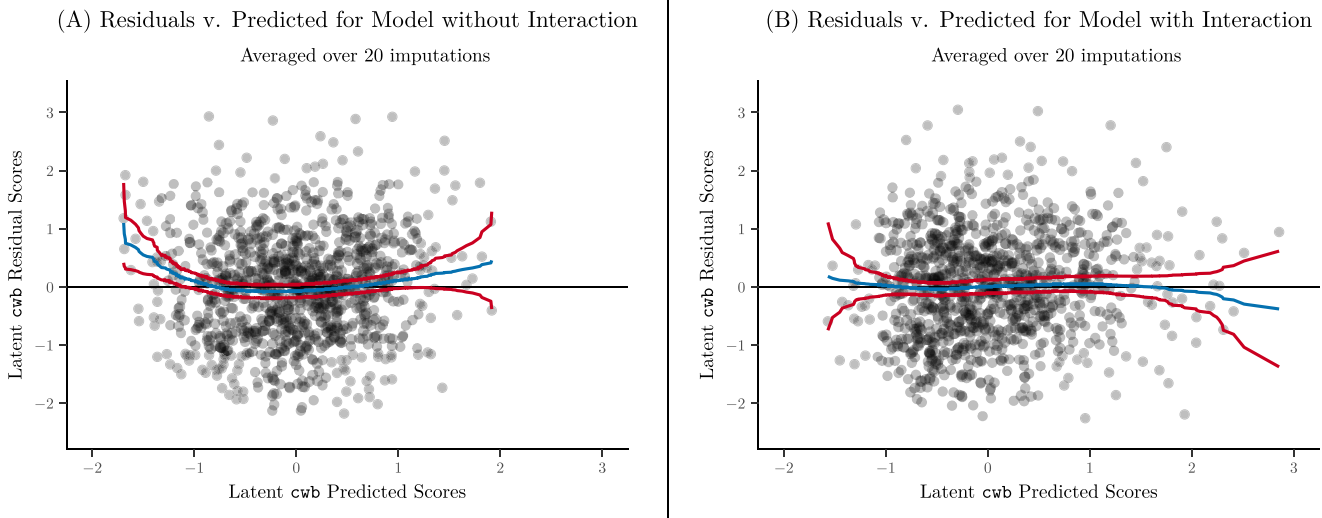
Equation 29 illustrates that I must include the interaction between the two latent variables, *consci* and *orgcon*. In Blimp's syntax, I explicitly specify the interaction within the regression model. The modified MODEL excerpt is as follows:

```
## Begin Modeling Syntax
MODEL:
  ## Latent Moderation Model
  # Explicitly write out the product with '*'
  cwb ~ consci orgcon consci*orgcon;
```

All other parts of the model specification exactly follow the Fitting a Latent Regression Model section, and Appendix B gives the full Blimp script for the example.

Under the Model 2 heading, Table 1 gives the posterior summaries for the latent moderation model. The negative interaction coefficient implies that this negative association between conscientiousness (*consci*) and counterproductive work behavior (*cwb*) strengthens (becomes more negative) as organizational constraints increase (*orgcon*). Because the 95% credible interval does not contain zero, one could conclude that the effect is "statistically significant" at $p < .05$. Turning to the lower-order effects, they have clear interpretations because the latent predictors have zero means. For example, at the average latent organizational constraints, the conditional effect of conscientiousness is negative, and zero is not within the 95% credible intervals. Thus, higher levels of conscientiousness are associated with lower

Figure 10
Regression Diagnostic Plots Based on Imputed Scores



Note. Panel A is the plot of the residuals versus predicted for *cwb* regressed on *consci* and *orgcon*. Panel B is the plot of the residuals versus predicted for *cwb* regressed on *consci*, *orgcon*, and the product of *consci* $\times$ *orgcon*. Each point represents the posterior mean of an observation's predicted and residual scores across 20 imputations. The center loess line represents the average fitted value of a loess line across 20 imputations. The outer loess lines represent ± 2 SE of the fitted value pooled across 20 imputations. *cwb* = counterproductive work behavior; *consci* = conscientiousness constraint; *orgcon* = organizational constraint. See the online article for the color version of this figure.

levels of counterproductive work behavior. Finally, note that including the product term has explained $\sim 3\%$ more variance in the latent outcome.

Like the model in the Fitting a Latent Regression Model section, I evaluated the residual diagnostics for the latent moderation model. Returning to Figure 10, Panel B plots the latent *cwb* residuals against its predicted values, both of which were now generated under the (correct) assumption that there is a latent interaction. Focusing on both panels, the diagnostic plot illustrates that including the latent interaction has most of the systematic relationship between the predicted and residual scores. Visually, the loess line now falls closer to the horizontal zero bar. Furthermore, the ± 2 *SE* bands consistently incorporate the zero within the predictive range, indicating that I have adequately modeled the functional form.

Conditional Effects Analysis for the Latent Moderation Model

In addition to estimating the latent moderation model, Blimp can perform conditional effects analysis to probe the interaction at different values of a moderator (Aiken & West, 1991; Bauer & Curran, 2005). Using the `SIMPLE` command, I requested the simple intercept and slope of the focal predictor (*consci*) at three levels of the moderator (*orgcon*): 1 *SD* below the latent *orgcon* mean, at the latent *orgcon* mean, and 1 *SD* above the latent *orgcon* mean.

```
## Conditional Effects Analysis with +/- 1 SD of orgcon
# mean = 0 by definition
# SD = 1 by definition
SIMPLE:
  consci | orgcon @ -1.0;
  consci | orgcon @ 0.0;
  consci | orgcon @ 1.0;
```

The variable to the left of the vertical bar is the focal variable, and the variable to the right is the moderator. The numeric value after the ampersand (@) is the moderator value being conditioned on. Recall that I chose identification constraints that standardized the latent predictor variables. Thus, the numeric value of the moderator listed after the ampersand reflects standard deviation units from the mean.

Because the conditional effects are simple functions of the estimated parameters, Blimp provides summary statistics of the posterior distribution for each simple intercept and simple slope. Table 2 gives the posterior summaries for the conditional effect's analysis at 1 *SD* above *orgcon*'s mean, at average *orgcon*, and 1 *SD* below *orgcon*'s mean. The conditional effect's analysis shows that the negative relationship between counterproductive work behavior and an employee's conscientiousness gets stronger as organizational constraints increase.

To further illustrate the interaction, I computed the conditional effect estimates and used them to construct plots. Figure 11 illustrates two established approaches for probing interactions. The plots were generated by the `simple_plot()` and `jn_plot()` functions in the `rblimp` package for R, and I include the commented R code to produce these plots at <https://github.com/blimp-stats/Latent-Interactions>. To begin, Panel A plots the conditional regression equations, illustrating the relationship between the latent counterproductive work behavior as a function of the latent conscientiousness factor evaluated at the different levels of organizational constraints. The bands around each regression line represent the 95%

Table 2
Results for Model 2 Conditional Effects

			Credible intervals	
Parameter	Coefficient	SD	2.5%	97.5%
consci orgcon @ -1				
Intercept	-0.632	0.045	-0.721	-0.546
Slope	-0.076	0.055	-0.184	0.031
consci orgcon @ 0				
Intercept	0.000	0.000	0.000	0.000
Slope	-0.300	0.040	-0.379	-0.223
consci orgcon @ +1				
Intercept	0.632	0.045	0.546	0.721
Slope	-0.523	0.061	-0.646	-0.406

Note. All variables are unobserved latent constructs. The "Coefficient" column is the posterior median. *consci* = conscientiousness constraint; *orgcon* = organizational constraint.

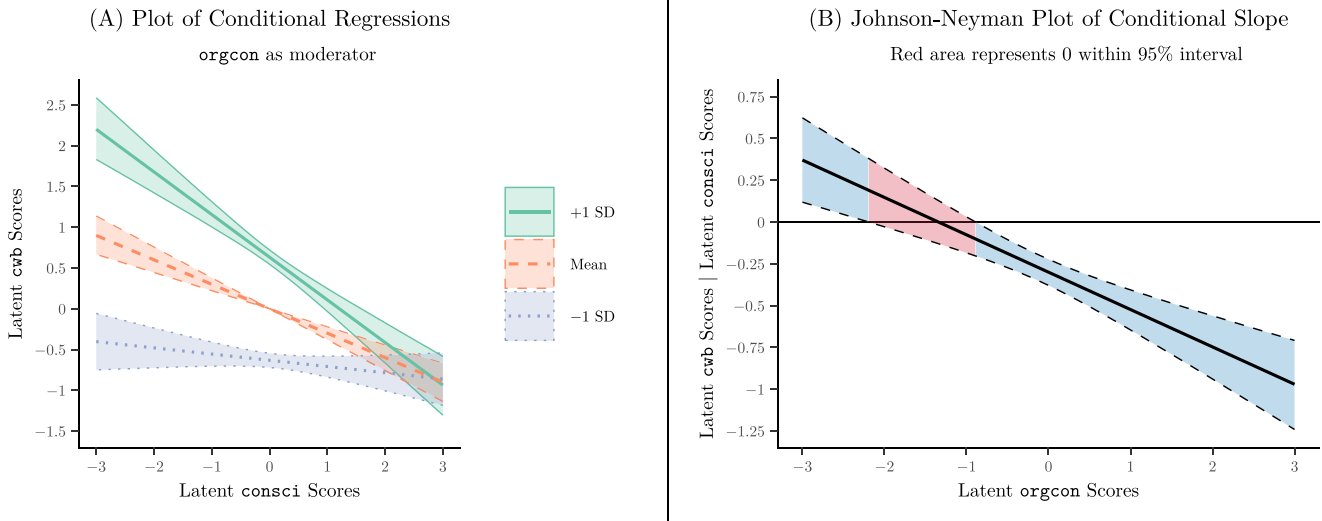
credible intervals of prediction. The flatter purple line near the bottom of the chart shows that having lower organizational constraints predicts little to no relationship between counterproductive work behavior and an employee's conscientiousness. Only as employee's perceive more constraints does lower conscientiousness predict higher counterproductive work behavior.

Panel B of Figure 11 illustrates a plot for the Johnson–Neyman approach (P. O. Johnson & Neyman, 1936) to probing interactions. On the *x* axis, I have different moderator values—that is, organizational constraints. On the *y* axis, I have the values of the conditional slope of latent counterproductive work behavior regressed on the latent conscientiousness factor. The shaded region represents the values of the moderator that produce a conditional effect where zero falls within the 95% credible intervals (i.e., a range of moderator scores that produce a nonsignificant conditional slope). The plot graphically illustrates that zero is a plausible value for the conditional relationship at about 1 *SD* below the mean of organizational constraints. Furthermore, zero remains within the interval until a bit past 2 *SDs* below the mean of organizational constraints.

Adding a Three-Way Latent Interaction

As previously noted, the general interaction model used by LMS and other procedures is restricted to the product of two latent variables. Fitting a model with a three-way interaction requires a complex model specification where two-way effects are defined by a proxy latent variable (Asparouhov & Muthén, 2021b). In contrast, the FR approach readily accommodates three-way effects. Extending the example, Zhou et al. (2014) also investigated how emotional stability (*emosta*) might moderate the two-way effects of organizational constraints and conscientiousness on counterproductive work behaviors. The emotional stability contained ten 5-point Likert items that loaded onto a common factor. Thus, the primary analysis model with a three-way latent variable interaction is as follows:

$$\begin{aligned}
 cwb_i = & \beta_0 + \beta_1(consci_i) + \beta_2(orgcon_i) \\
 & + \beta_3(emosta_i) + \beta_4(consci_i \times orgcon_i) \\
 & + \beta_5(consci_i \times emosta_i) + \beta_6(orgcon_i \times emosta_i) \\
 & + \beta_7(consci_i \times orgcon_i \times emosta_i) + \zeta_i.
 \end{aligned}
 \tag{30}$$

Figure 11*Conditional Effects Plots for *cwb* Regressed on *consci* Moderated by *orgcon**

Note. Panel A is the regression of latent *cwb* on latent *consci* moderated at different levels of latent *orgcon*. Panel B is the Johnson–Neyman plot to evaluate when the credible intervals for the conditional slope. The ribbons around each line depict the corresponding 95% credible intervals. *cwb* = counterproductive work behavior; *consci* = conscientiousness constraint; *orgcon* = organizational constraint. See the online article for the color version of this figure.

Fitting the three-way latent interaction model requires only a minor change to the code. In Blimp’s syntax, I explicitly specify all interactions within the regression model. The modified MODEL excerpt is as follows:

```
## Begin Modeling Syntax
MODEL:
  ## Threeway Latent Moderation Model
  cwb ~ consci@b1 orgcon@b2 emosta@b3
        consci*orgcon@b4 consci*emosta@b5 orgcon*emosta@b6
        consci*orgcon*emosta@b7;
```

In the code above, I have included parameter labels on each coefficient using the ampersand (@) followed by the name based on Equation 30. As demonstrated below, I use these parameter labels to probe the three-way interaction. Note that the code also specifies a measurement model (not shown here) for *emosta*, and it correlates the new latent factor with the other two exogenous latent variables. The entire script can be accessed at <https://github.com/blimp-stats/Latent-Interactions>.

Under the Model 3 heading, Table 1 gives the posterior summaries for the three-way latent interaction model. Because all latent variables have a mean of zero, the two-way interactions in the table are conditional effects at the mean of the third variable. For example, consider the latent conscientiousness by latent organizational constraints interaction from the Fitting a Latent Moderation Model section. The two-way interaction and lower-order terms are similar in magnitude and sign, so the substantive interpretations only slightly change by referencing an employee with average emotional stability. The three-way interaction captures the degree to which the two-way effect changes as a function of emotional stability. The negative slope means that the two-way product of conscientiousness and organizational constraints becomes more negative as emotional stability increases. Because zero does not fall inside the 95%

credible interval, one could conclude that the parameter differs from zero. For researchers who prefer a frequentist-like test statistic, a Bayesian Wald χ^2 test (Asparouhov & Muthén, 2021a) of the three-way interaction (β_7) was statistically significant, $\chi^2(1) = 4.502, p = .034$.

To further investigate the three-way interaction, I probe the three-way effect by computing the two-way interaction between *consci* and *orgcon* conditional on different levels of *emosta*. In Blimp, I computed these auxiliary conditional effect parameters using the PARAMETERS command and the parameter labels specified earlier.

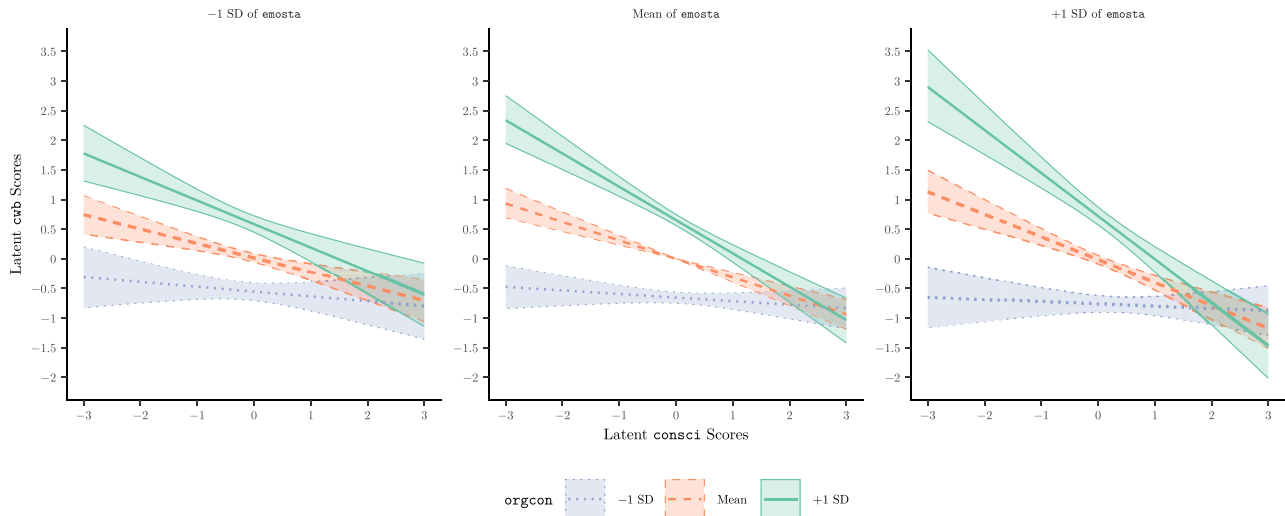
```
# Obtain two-way conditional effects
PARAMETERS:
  inter.emo_m1 = b4 + b7 * (-1);
  inter.emo_0  = b4 + b7 * ( 0);
  inter.emo_p1 = b4 + b7 * (+1);
```

The excerpt above creates three new generate quantities that represent the two-way interaction above, below and at the mean of *emosta*. Because these three quantities are functions of the estimated parameters, Blimp provides summary statistics of the posterior distribution.

To further illustrate the interaction, I compute the conditional effect estimates and used them to construct plots to probe the interactions. Figure 12 illustrates the relationship between the latent counterproductive work behavior as a function of the latent conscientiousness evaluated at the different levels of organizational constraints and employee emotional stability. The bands around each regression line represent the 95% credible intervals of prediction. Moving from the left to the right in Figure 12 illustrates that the interaction becomes strong (i.e., it’s coefficient more negative) as emotional stability increases.

Figure 12

*Conditional Effects Plots for *cwb* Regressed on *consci* Moderated by *orgcon* and *emosta**



Note. Plots of the regression of latent *cwb* on latent *consci* moderated at different levels of latent *orgcon* and latent *emosta*. The ribbons around each line depict the corresponding 95% credible intervals. *cwb* = counterproductive work behavior; *consci* = conscientiousness constraint; *orgcon* = organizational constraint; *emosta* = emotional stability. See the online article for the color version of this figure.

In addition to the conditional effects plot, Figure 13 illustrates the Johnson–Neyman plot for the two-way interaction between conscientiousness and organizational constraints moderated by emotional stability. On the x axis, I have different values of emotional stability. On the y axis, I have the values of the conditional slope for the two-way interaction between conscientiousness and organizational constraints. The shaded region represents the values of the moderator that produce a conditional effect where zero falls within the 95% credible intervals (i.e., two-way effects that would not be deemed statistically significant). The plot graphically illustrates that a zero interaction is a plausible value for the conditional relationship at about 1 *SD* below the mean of emotional stability.

Discussion

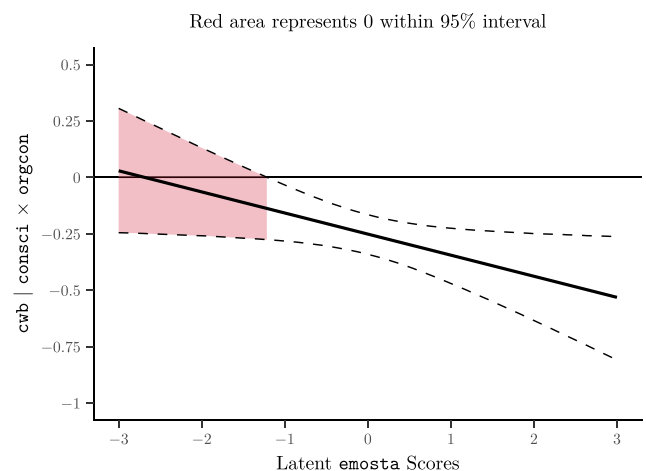
Moderated regression models are widely used across the social sciences, and significant methodological work has been dedicated to extending these models to incorporate latent variables. Several approaches have been proposed to estimate these models (see Kelava & Brandt, 2023, for more details), with prominent approaches being PI methods, LMS, and data augmentation-based Bayesian approaches.

This article extends and generalizes data augmentation-based Bayesian estimation for latent variable interactions with a FR framework. This framework expresses the multivariate distribution of variables as a collection of simpler distributions. These simpler distributions can be configured to accommodate complex distributional features (e.g., heteroscedastic variation of a latent predictor) incompatible with the multivariate normality assumption invoked by traditional SEMs. Overall, the FR approach builds on previous work both in the manifest missing data literature (Ibrahim, 1990; Ibrahim et al., 1999, 2002; Lüdtke et al., 2020b; Lipsitz & Ibrahim, 1996) and the Bayesian nonlinear SEM literature (Lee, 2007; Lee & Song, 2012; Lee et al., 2007; Lee & Zhu, 2000). I leverage these previous works to lay out the FR approach to latent variable modeling.

The framework treats latent variables as missing data that should be estimated via imputation (Blackwell et al., 2017; Lee & Shi, 2000; Little & Rubin, 2002). Once the latent missing data are imputed, the estimation steps parallel those for any regression model; therefore, allowing FR to fit models not possible with existing methods (e.g., latent by multicategorical interactions).

Figure 13

*Johnson–Neyman Plot for *consci* × *orgcon* Moderated by *emosta**



Note. A Johnson–Neyman plot to evaluate when the credible intervals for the two-way interaction between latent *consci* and *orgcon* at different levels of latent *emosta*. The ribbons around each line depict the corresponding 95% credible intervals. *consci* = conscientiousness constraint; *orgcon* = organizational constraint; *emosta* = emotional stability. See the online article for the color version of this figure.

The Monte Carlo simulations presented in this article evaluate the performance of the FR method in comparison to the LMS and PI approaches across three different interaction types: latent by latent, latent by manifest normally distributed, and latent by manifest binary. The simulations illustrated that FR performs as well or sometimes better than existing methods. In practical terms of power, the simulation results suggest that there is no material downside to imputing the latent variables instead of marginalizing over them as in ML. In contrast to a previous finding by Lüdtke et al. (2020b, Supplemental J in the additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY>), these simulations suggest that FR can produce unbiased results with smaller sample sizes; however, Lüdtke et al. (2020b) estimation algorithm was a Metropolis–Hastings algorithm as opposed to the analytically derived solutions presented in Appendix A. Overall, the Monte Carlo results reinforce the utility of FR as a framework for modeling complex latent interactions, and future studies can expand on these simulations by examining FR in other models (e.g., three-way interactions, quadratic effects, multicategorical-by-latent interactions).

The FR approach discussed in this article is implemented in the statistical software Blimp (Keller & Enders, 2022), a free standalone software with syntax similarities to many modern-day SEM software packages. Blimp allows for the specification of complex structural equation models with interactive and nonlinear effects for latent variables, including three-way or higher interactions and quadratic effects. These nonlinear effects can be specified with any variable, including a mixture of manifest and latent variables, and they are not strictly limited to the structural models. It can easily accommodate a diverse set of manifest variable types (e.g., binary, ordinal, multicategorical, continuous normal, skewed continuous, and count indicators), which can serve as indicators of latent variables or as manifest predictors with nonlinear relationships. As I demonstrated in the illustrative analysis example, the practical advantage is that these methods are plug-and-play with familiar procedures with moderated regression models—for example, centering, conditional effects analysis, regression diagnostics, Johnson–Neyman plots, making these models accessible to a much broader audience.

To link to existing research, I provided a few examples of the types of latent variable interactions this framework can handle. However, nothing precludes FR from incorporating more complex models. For example, Enders et al. (2024) illustrate using the framework to assess measurement invariance and integrative data analysis via moderated nonlinear factor analysis (Bauer, 2017). In this approach, the manifest predictor variables moderate loadings and variances without the complications of imposing nonlinear constraints, and the framework can further extend this to allow latent variables to moderate the loadings and variances. Another example is intensive longitudinal data, which has recently shown an interest in heterogeneous variances. In these models, the magnitude of the intra-individual variation is represented as a level-2 latent variable (Ernst et al., 2024; Feng & Hancock, 2024; Lester et al., 2021; McNeish, 2021). In the FR framework, the latent variable representing the heterogeneous variance can be a moderator for any other variable in the model. Other examples include interaction effects in dynamic structural equation models discussed by Asparouhov et al. (2018) and latent moderated mediation models with single-level or multilevel data (Keller, 2022). More broadly, general nonlinear structural equation models (see Harring & Zou, 2023; Lee & Song, 2012) can be estimated, allowing for nonlinear effects, and general nonlinear structural equation

models with latent and manifest variables (e.g., quadratic effects, sine-cosine models, and Jenss–Bayley functions; Jenss & Bayley, 1937).

Although not discussed in this article, FR readily extends to multilevel latent variable modeling, where sampling units are clustered within another unit. My approach easily accommodates the decomposition of a multilevel interaction effect into its within, between, and cross-level interactions via latent cluster mean centering discussed by Preacher et al. (2016). Current methods using the general interaction model are limited, often relying on unintuitive proxy variables. By requiring these proxy variables, the general interaction model cannot handle categorical moderators and mediators in the context of multilevel models (Asparouhov & Muthén, 2021b). In contrast, the FR framework provides a straightforward path for incorporating a categorical moderator or mediator. FR also readily incorporates within-cluster centering at level-2 and level-3 latent group means, and latent centering is available for continuous, binary, ordinal, and multicategorical predictors. The Blimp User's Guide (Keller & Enders, 2022) provides illustrative syntax examples for multilevel modeling.

All of the modeling possibilities described above provide avenues for future methodological work. While the simulation results in this article suggest that $N = 250$ provided a sufficient sample size to estimate the parameters reliably, the simulation only investigated a limited number of conditions. Additional methodological research is still needed to determine the data requirements to estimate these complex models reliably across various conditions and the robustness to violations of the model-specific distributions. Similarly, power studies are needed to provide researchers with guidelines linked to presumed effect sizes. Furthermore, as discussed by Lüdtke et al. (2020b, Supplemental J in the additional online materials on the OSF at <https://doi.org/10.17605/OSF.IO/2GHNY>), future studies could investigate the usage of FR to provide a richer multiple imputation model (i.e., the inclusion of auxiliary variables to increase the plausibility of missing data assumptions) to estimate the nonlinear relationships with latent variables.

In summary, this article provides researchers with a FR approach to fitting latent variable interactions of all types. It fully generalizes past work and offers interesting extensions that were not previously possible. All of this is provided in the user-friendly and free software package Blimp.

References

- Agresti, A. (2012). *Analysis of ordinal categorical data* (3rd ed.). Wiley.
- Aiken, L. S., & West, S. G. (1991). *Multiple regression: Testing and interpreting interactions*. Sage Publications.
- Albert, J. H., & Chib, S. (1993). Bayesian analysis of binary and polychotomous response data. *Journal of the American Statistical Association*, 88(422), 669–679. <https://doi.org/10.1080/01621459.1993.10476321>
- Asparouhov, T., Hamaker, E. L., & Muthén, B. (2018). Dynamic structural equation models. *Structural Equation Modeling: A Multidisciplinary Journal*, 25(3), 359–388. <https://doi.org/10.1080/10705511.2017.1406803>
- Asparouhov, T., & Muthén, B. (2010a). *Bayesian analysis using Mplus: Technical implementation* [Electronic Article]. <https://www.statmodel.com/download/Bayes3.pdf>
- Asparouhov, T., & Muthén, B. (2010b). *Plausible values for latent variables using Mplus* [Electronic Article]. <https://www.statmodel.com/download/Plausible.pdf>
- Asparouhov, T., & Muthén, B. (2021a). Advances in Bayesian model fit evaluation for structural equation models. *Structural Equation Modeling: A Multidisciplinary Journal*, 28(1), 1–14. <https://doi.org/10.1080/10705511.2020.1764360>

- Asparouhov, T., & Muthén, B. (2021b). Bayesian estimation of single and multilevel models with latent variable interactions. *Structural Equation Modeling: A Multidisciplinary Journal*, 28(2), 314–328. <https://doi.org/10.1080/10705511.2020.1761808>
- Aytürk, E., Cham, H., Jennings, P. A., & Brown, J. L. (2020). Latent variable interactions with ordered-categorical indicators: Comparisons of unconstrained product indicator and latent moderated structural equations approaches. *Educational and Psychological Measurement*, 80(2), 262–292. <https://doi.org/10.1177/0013164419865017>
- Aytürk, E., Cham, H., Jennings, P. A., & Brown, J. L. (2021). Exploring the performance of latent moderated structural equations approach for ordered-categorical items. *Structural Equation Modeling: A Multidisciplinary Journal*, 28(3), 410–422. <https://doi.org/10.1080/10705511.2020.1810047>
- Baron, R. M., & Kenny, D. A. (1986). The moderator–mediator variable distinction in social psychological research: Conceptual, strategic, and statistical considerations. *Journal of Personality and Social Psychology*, 51(6), 1173–1182. <https://doi.org/10.1037/0022-3514.51.6.1173>
- Bartlett, J., Keogh, R., & Bonneville, E. F. (2021). *smcfc: Multiple imputation of covariates by substantive model compatible fully conditional specification*. <https://CRAN.R-project.org/package=smcfc>
- Bauer, D. J. (2017). A more general model for testing measurement invariance and differential item functioning. *Psychological Methods*, 22(3), 507–526. <https://doi.org/10.1037/met0000077>
- Bauer, D. J., & Curran, P. J. (2005). Probing interactions in fixed and multilevel regression: Inferential and graphical techniques. *Multivariate Behavioral Research*, 40(3), 373–400. https://doi.org/10.1207/s15327906mbr4003_5
- Blackwell, M., Honaker, J., & King, G. (2017). A unified approach to measurement error and missing data: Overview and applications. *Sociological Methods & Research*, 46(3), 303–341. <https://doi.org/10.1177/0049124115585360>
- Bohmstedt, G. W., & Goldberger, A. S. (1969). On the exact covariance of products of random variables. *Journal of the American Statistical Association*, 64(328), 1439–1442. <https://doi.org/10.1080/01621459.1969.10501069>
- Bowling, N. A., & Eschleman, K. J. (2010). Employee personality as a moderator of the relationships between work stressors and counterproductive work behavior. *Journal of Occupational Health Psychology*, 15(1), 91–103. <https://doi.org/10.1037/a0017326>
- Cham, H., Reshetnyak, E., Rosenfeld, B., & Breitbart, W. (2017). Full information maximum likelihood estimation for latent variable interactions with incomplete indicators. *Multivariate Behavioral Research*, 52(1), 12–30. <https://doi.org/10.1080/00273171.2016.1245600>
- Cham, H., West, S. G., Ma, Y., & Aiken, L. S. (2012). Estimating latent variable interactions with nonnormal observed data: A comparison of four approaches. *Multivariate Behavioral Research*, 47(6), 840–876. <https://doi.org/10.1080/00273171.2012.732901>
- Cook, R., & Weisberg, S. (1999). *Applied regression including computing and graphics*. John Wiley & Sons.
- Costa, P. T., & McCrae, R. R. (1992). Normal personality assessment in clinical practice: The neo personality inventory. *Psychological Assessment*, 4(1), 5–13. <https://doi.org/10.1037/1040-3590.4.1.5>
- Cowles, M. K. (1996). Accelerating Monte Carlo Markov chain convergence for cumulative-link generalized linear models. *Statistics and Computing*, 6(2), 101–111. <https://doi.org/10.1007/BF00162520>
- Dempster, A. P., Laird, N. M., & Rubin, D. B. (1977). Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)*, 39(1), 1–22. <https://doi.org/10.1111/j.2517-6161.1977.tb01600.x>
- Du, H., Enders, C., Keller, B. T., Bradbury, T. N., & Karney, B. R. (2022). A Bayesian latent variable selection model for nonignorable missingness. *Multivariate Behavioral Research*, 57(2–3), 478–512. <https://doi.org/10.1080/00273171.2021.1874259>
- Enders, C. K. (2022). *Applied missing data analysis*. Guilford Publications.
- Enders, C. K. (2025). Missing data: An update on the state of the art. *Psychological Methods*, 30(2), 322–339. <https://doi.org/10.1037/met0000563>
- Enders, C. K., Baraldi, A. N., & Cham, H. (2014). Estimating interaction effects with incomplete predictor variables. *Psychological Methods*, 19(1), 39–55. <https://doi.org/10.1037/a0035314>
- Enders, C. K., Du, H., & Keller, B. T. (2020). A model-based imputation procedure for multilevel regression models with random coefficients, interaction effects, and other nonlinear terms. *Psychological Methods*, 25(1), 88–112. <https://doi.org/10.1037/met0000228>
- Enders, C. K., Vera, J. D., Keller, B. T., Lenartowicz, A., & Loo, S. K. (2024). Building a simpler moderated nonlinear factor analysis model with Markov chain Monte Carlo estimation. *Psychological Methods*. Advance online publication. <https://doi.org/10.1037/met0000712>
- Erler, N. S., Rizopoulos, D., & Lesaffre, E. M. E. H. (2021). JointAI: Joint analysis and imputation of incomplete data in R. *Journal of Statistical Software*, 100(20), 1–56. <https://doi.org/10.18637/jss.v100.i20>
- Erler, N. S., Rizopoulos, D., Rosmalen, J. v., Jaddoe, V. W., Franco, O. H., & Lesaffre, E. M. (2016). Dealing with missing covariates in epidemiologic studies: A comparison between multiple imputation and a full Bayesian approach. *Statistics in Medicine*, 35(17), 2955–2974. <https://doi.org/10.1002/sim.6944>
- Ernst, A. F., Albers, C. J., & Timmerman, M. E. (2024). A comprehensive model framework for between-individual differences in longitudinal data. *Psychological Methods*, 29(4), 748–766. <https://doi.org/10.1037/met0000585>
- Feng, Y., & Hancock, G. R. (2024). A structural equation modeling approach for modeling variability as a latent variable. *Psychological Methods*, 29(2), 262–286. <https://doi.org/10.1037/met0000477>
- Fox, S., Spector, P. E., & Miles, D. (2001). Counterproductive work behavior (CWB) in response to job stressors and organizational justice: Some mediator and moderator tests for autonomy and emotions. *Journal of Vocational Behavior*, 59(3), 291–309. <https://doi.org/10.1006/jvbe.2001.1803>
- Gelman, A., Carlin, J. B., Stern, H. S., Dunson, D. B., Vehtari, A., & Rubin, D. B. (2013). *Bayesian data analysis*. CRC Press.
- Gelman, A., & Rubin, D. B. (1992). Inference from iterative simulation using multiple sequences. *Statistical Science*, 7(4), 457–472. <https://doi.org/10.1214/ss/1177011136>
- Gelman, A., Stern, H. S., Carlin, J. B., Dunson, D. B., Vehtari, A., & Rubin, D. B. (2013). *Bayesian data analysis*. Chapman & Hall/CRC Press.
- Geman, S., & Geman, D. (1984). Stochastic relaxation, Gibbs distributions, and the Bayesian restoration of images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-6 (6), 721–741. <https://doi.org/10.1109/TPAMI.1984.4767596>
- Geyer, C. J. (2011). Introduction to Markov chain Monte Carlo. In S. Brooks, A. Gelman, G. Jones, & X.-L. Meng (Eds.), *Handbook of Markov chain Monte Carlo* (Vol. 20116022, No. 45, p. 22). CRC Press.
- Goldstein, H., Carpenter, J. R., & Browne, W. J. (2014). Fitting multilevel multivariate models with missing data in responses and covariates that may include interactions and non-linear terms. *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 177(2), 553–564. <https://doi.org/10.1111/rssa.12022>
- Hallquist, M. N., & Wiley, J. F. (2018). MplusAutomation: An R package for facilitating large-scale latent variable analyses in Mplus. *Structural Equation Modeling*, 25(4), 621–638. <https://doi.org/10.1080/10705511.2017.1402334>
- Harring, J. R., & Zou, J. (2023). Nonlinear structural equation models: Advanced methods and applications. In R. H. Hoyle (Ed.), *Handbook of structural equation modeling* (pp. 681–700). Guilford Publications.
- Ibrahim, J. G. (1990). Incomplete data in generalized linear models. *Journal of the American Statistical Association*, 85(411), 765–769. <https://doi.org/10.1080/01621459.1990.10474938>
- Ibrahim, J. G., Chen, M.-H., & Lipsitz, S. R. (1999). Monte Carlo EM for missing covariates in parametric regression models. *Biometrics*, 55(2), 591–596. <https://doi.org/10.1111/j.0006-341X.1999.00591.x>
- Ibrahim, J. G., Chen, M.-H., & Lipsitz, S. R. (2002). Bayesian methods for generalized linear models with covariates missing at random. *Canadian Journal of Statistics*, 30(1), 55–78. <https://doi.org/10.2307/3315865>

- Jenss, R. M., & Bayley, N. (1937). A mathematical method for studying the growth of a child. *Human Biology*, 9(4), Article 556.
- Johnson, P. O., & Neyman, J. (1936). Tests of certain linear hypotheses and their application to some educational problems. *Statistical Research Memoirs*, 1, 57–93.
- Johnson, V. E., & Albert, J. H. (2006). *Ordinal data modeling*. Springer Science & Business Media.
- Kaplan, D., & Depaoli, S. (2012). Bayesian structural equation modeling. In R. H. Hoyle (Ed.), *Handbook of structural equation modeling* (pp. 650–673). The Guilford Press.
- Kelava, A. (2009). *Multikollinearität in nicht-linearen latenten Strukturgleichungsmodellen* [Unpublished doctoral dissertation]. Dissertation, Johann Wolfgang Goethe-Universität.
- Kelava, A., & Brandt, H. (2009). Estimation of nonlinear latent structural equation models using the extended unconstrained approach. *Review of Psychology*, 16(2), 123–132.
- Kelava, A., & Brandt, H. (2023). Latent interaction effects. In R. H. Hoyle (Ed.), *Handbook of structural equation modeling* (pp. 427–446). Guilford Publications.
- Kelava, A., Werner, C. S., Schermelleh-Engel, K., Moosbrugger, H., Zapf, D., Ma, Y., Cham, H., Aiken, L. S., & West, S. G. (2011). Advanced nonlinear latent variable modeling: Distribution analytic LMS and QML estimators of interaction and quadratic effects. *Structural Equation Modeling: A Multidisciplinary Journal*, 18(3), 465–491. <https://doi.org/10.1080/10705511.2011.582408>
- Keller, B. T. (2022). An introduction to factored regression models with Blimp. *Psych*, 4(1), 10–37. <https://doi.org/10.3390/psych4010002>
- Keller, B. T. (2024a). *Automatically derived full conditionals for factored regression models*. Manuscript in preparation.
- Keller, B. T. (2024b). *rblimp: Integration of Blimp software into R* [Computer software manual]. <https://github.com/blimp-stats/rblimp> (R Ppackage Version 0.2.2, commit 9f762667777a5aeb350cb407d0cc3c682fccc5).
- Keller, B. T. (2024c). *Supplemental materials for preprint: A factored regression approach to modeling latent variable interactions and nonlinear effects*. OSF. <https://doi.org/10.17605/OSF.IO/2GHNY>
- Keller, B. T., & Enders, C. K. (2022). *Blimp user's guide* (Version 3.0).
- Keller, B. T., & Enders, C. K. (2023). An investigation of factored regression missing data methods for multilevel models with cross-level interactions. *Multivariate Behavioral Research*, 58(5), 938–963. <https://doi.org/10.1080/00273171.2022.2147049>
- Kenny, D. A., & Judd, C. M. (1984). Estimating the nonlinear and interactive effects of latent variables. *Psychological Bulletin*, 96(1), 201–210. <https://doi.org/10.1037/0033-2909.96.1.201>
- Kim, S., Lee, S., Kim, J., & Whittaker, T. A. (2024). Investigating latent interaction effects in multiple-group analysis in the structural equation modeling framework. *Structural Equation Modeling: A Multidisciplinary Journal*, 31(6), 1043–1055. <https://doi.org/10.1080/10705511.2024.2363827>
- Klein, A., & Moosbrugger, H. (2000). Maximum likelihood estimation of latent interaction effects with the LMS method. *Psychometrika*, 65(4), 457–474. <https://doi.org/10.1007/BF02296338>
- Klein, A., & Muthén, B. O. (2007). Quasi-maximum likelihood estimation of structural equation models with multiple interaction and quadratic effects. *Multivariate Behavioral Research*, 42(4), 647–673. <https://doi.org/10.1080/00273170701710205>
- Klopp, E., & Klößner, S. (2021). The impact of scaling methods on the properties and interpretation of parameter estimates in structural equation models with latent variables. *Structural Equation Modeling: A Multidisciplinary Journal*, 28(2), 182–206. <https://doi.org/10.1080/10705511.2020.1796673>
- Koehler, E., Brown, E., & Haneuse, S. J.-P. (2009). On the assessment of Monte Carlo error in simulation-based statistical analyses. *The American Statistician*, 63(2), 155–162. <https://doi.org/10.1198/tast.2009.0030>
- Lee, S.-Y. (2007). *Structural equation modeling: A Bayesian approach*. John Wiley & Sons.
- Lee, S.-Y., & Shi, J.-Q. (2000). Joint Bayesian analysis of factor scores and structural parameters in the factor analysis model. *Annals of the Institute of Statistical Mathematics*, 52(4), 722–736. <https://doi.org/10.1023/A:1017529427433>
- Lee, S.-Y., & Song, X.-Y. (2003). Model comparison of nonlinear structural equation models with fixed covariates. *Psychometrika*, 68(1), 27–47. <https://doi.org/10.1007/BF02296651>
- Lee, S.-Y., & Song, X.-Y. (2012). *Basic and advanced Bayesian structural equation modeling: With applications in the medical and behavioral sciences*. John Wiley & Sons.
- Lee, S.-Y., Song, X. Y., & Poon, W.-Y. (2004). Comparison of approaches in estimating interaction and quadratic effects of latent variables. *Multivariate Behavioral Research*, 39(1), 37–67. https://doi.org/10.1207/s15327906mbr3901_2
- Lee, S.-Y., Song, X.-Y., & Tang, N.-S. (2007). Bayesian methods for analyzing structural equation models with covariates, interaction, and quadratic latent variables. *Structural Equation Modeling: A Multidisciplinary Journal*, 14(3), 404–434. <https://doi.org/10.1080/10705510701301511>
- Lee, S.-Y., & Zhu, H.-T. (2000). Statistical analysis of nonlinear structural equation models with continuous and polytomous data. *British Journal of Mathematical and Statistical Psychology*, 53(2), 209–232. <https://doi.org/10.1348/000711000159303>
- Lester, H. F., Cullen-Lester, K. L., & Walters, R. W. (2021). From nuisance to novel research questions: Using multilevel models to predict heterogeneous variances. *Organizational Research Methods*, 24(2), 342–388. <https://doi.org/10.1177/1094428119887434>
- Levy, R., & Mislevy, R. J. (2017). *Bayesian psychometric modeling*. CRC Press.
- Lipsitz, S. R., & Ibrahim, J. G. (1996). A conditional model for incomplete covariates in parametric regression models. *Biometrika*, 83(4), 916–922. <https://doi.org/10.1093/biomet/83.4.916>
- Little, R. J., & Rubin, D. B. (2002). *Statistical analysis with missing data*. John Wiley & Sons.
- Liu, J., Gelman, A., Hill, J., Su, Y.-S., & Kropko, J. (2014). On the stationary distribution of iterative imputations. *Biometrika*, 101(1), 155–173. <https://doi.org/10.1093/biomet/ast044>
- Lonati, S., Rönkkö, M., & Antonakis, J. (2024). Normality assumption in latent interaction models. *Psychological Methods*. Advance online publication. <https://doi.org/10.1037/met0000657>
- Lüdtke, O., Robitzsch, A., & West, S. G. (2020a). Analysis of interactions and nonlinear effects with missing data: A factored regression modeling approach using maximum likelihood estimation. *Multivariate Behavioral Research*, 55(3), 361–381. <https://doi.org/10.1080/00273171.2019.1640104>
- Lüdtke, O., Robitzsch, A., & West, S. G. (2020b). Regression models involving nonlinear effects with missing data: A sequential modeling approach using Bayesian estimation. *Psychological Methods*, 25(2), 157–181. <https://doi.org/10.1037/met0000233>
- Lynch, S. M. (2007). *Introduction to applied Bayesian statistics and estimation for social scientists*. Springer Science & Business Media.
- Mansolf, M. (2023). A true score imputation method to account for psychometric measurement error. *Psychological Methods*, 30(3), 636–659. <https://doi.org/10.1037/met0000578>
- Marsh, H. W., Wen, Z., & Hau, K.-T. (2004). Structural equation models of latent interactions: Evaluation of alternative estimation strategies and indicator construction. *Psychological Methods*, 9(3), 275–300. <https://doi.org/10.1037/1082-989X.9.3.275>
- McNeish, D. (2021). Specifying location-scale models for heterogeneous variances as multilevel sems. *Organizational Research Methods*, 24(3), 630–653. <https://doi.org/10.1177/1094428120913083>
- Mealli, F., & Rubin, D. B. (2015). Clarifying missing at random and related definitions, and implications when coupled with exchangeability. *Biometrika*, 102(4), 995–1000. <https://doi.org/10.1093/biomet/asv035>

- Merkle, E. C., Fitzsimmons, E., Uanboro, J., & Goodrich, B. (2020). *Efficient Bayesian structural equation modeling in Stan*. arXiv preprint <https://arxiv.org/abs/2008.07733>
- Merkle, E. C., & Rosseel, Y. (2018). blavaan: Bayesian structural equation models via parameter expansion. *Journal of Statistical Software*, 85(4), 1–30. <https://doi.org/10.18637/jss.v085.i04>
- Muthén, B., & Asparouhov, T. (2009). Growth mixture modeling: Analysis with non-Gaussian random effects. In G. Fitzmaurice, M. Davidian, G. Verbeke, & G. Molenberghs (Eds.), *Longitudinal data analysis* (pp. 143–165). Chapman & Hall/CRC Press.
- Muthén, B., & Asparouhov, T. (2012). Bayesian structural equation modeling: A more flexible representation of substantive theory. *Psychological Methods*, 17(3), 313–335. <https://doi.org/10.1037/a0026802>
- Muthén, L., & Muthén, B. (2017). *Mplus user's guide* (8th ed.).
- Palomo, J., Dunson, D. B., & Bollen, K. (2007). Bayesian structural equation modeling. In S.-Y. Lee (Ed.), *Handbook of latent variable and related models* (pp. 163–188). Elsevier.
- Penney, L. M., & Spector, P. E. (2005). Job stress, incivility, and counterproductive work behavior (CWB): The moderating role of negative affectivity. *Journal of Organizational Behavior: The International Journal of Industrial, Occupational and Organizational Psychology and Behavior*, 26(7), 777–796. <https://doi.org/10.1002/job.336>
- Plummer, M. (2017). *Jags Version 4.3.0 User Manual* [Computer Software].
- Preacher, K. J., Zhang, Z., & Zyphur, M. J. (2016). Multilevel structural equation models for assessing moderation within and across levels of analysis. *Psychological Methods*, 21(2), 189–205. <https://doi.org/10.1037/met0000052>
- R Core Team. (2022). *R: A language and environment for statistical computing* [Computer software manual]. <https://www.R-project.org/>
- Robitzsch, A., & Lüdtke, O. (2021). *mdmb: Model based treatment of missing data*. <https://CRAN.R-project.org/package=mdmb>
- Rubin, D. B. (1976). Inference and missing data. *Biometrika*, 63(3), 581–592. <https://doi.org/10.1093/biomet/63.3.581>
- Rubin, D. B. (1987). *Multiple imputation for nonresponse in surveys*. Wiley.
- Seaman, S. R., Bartlett, J. W., & White, I. R. (2012). Multiple imputation of missing covariates with non-linear effects and interactions: An evaluation of statistical methods. *BMC Medical Research Methodology*, 12(1), Article 46. <https://doi.org/10.1186/1471-2288-12-46>
- Shen, Z., Chen, L., Somer, E., Miočević, M., & Falk, C. F. (2025). Latent variable interactions with categorical indicators: Continuous and categorical latent moderated structural equations approaches. *Structural Equation Modeling: A Multidisciplinary Journal*, 32(3), 460–474. <https://doi.org/10.1080/10705511.2025.2497088>
- Spector, P. E., & Fox, S. (2005). The stressor-emotion model of counterproductive work behavior. In S. Fox & P. E. Spector (Eds.), *Counterproductive work behavior: Investigations of actors and targets* (pp. 151–174). American Psychological Association. <https://doi.org/10.1037/10893-007>
- Stan Development Team. (2023). *The Stan Core Library* [Computer Software]. <https://mc-stan.org>
- Tanner, M. A., & Wong, W. H. (1987). The calculation of posterior distributions by data augmentation. *Journal of the American Statistical Association*, 82(398), 528–540. <https://doi.org/10.1080/01621459.1987.10478458>
- Yeo, I.-K., & Johnson, R. A. (2000). A new family of power transformations to improve normality or symmetry. *Biometrika*, 87(4), 954–959. <https://doi.org/10.1093/biomet/87.4.954>
- Zhang, Q., & Wang, L. (2017). Moderation analysis with missing data in the predictors. *Psychological Methods*, 22(4), 649–666. <https://doi.org/10.1037/met0000104>
- Zhou, Z. E., Meier, L. L., & Spector, P. E. (2014). The role of personality and job stressors in predicting counterproductive work behavior: A three-way interaction. *International Journal of Selection and Assessment*, 22(3), 286–296. <https://doi.org/10.1111/ijsa.12077>

Appendix A

Algorithmically Derived Conditional Distributions

This discussion provides an overview of how Blimp derives the analytical distribution for general products and centering. Further details, including the extension to multilevel and multivariate models, are given by Keller (2024a). Finally, note that if a model does not map onto the models discussed in this section, Blimp can still estimate them using an adaptively tuned random-walk Metropolis–Hastings algorithm.

Derivation of General Underlying Model

Suppose I am interested in deriving the conditional density $f(X | Y, \dots)$, where the ellipsis (“...”) represent any other variables

in the model. I express this density up to proportionality as follows:

$$f(Y | X, \dots) \propto f(X | \dots). \quad (\text{A1})$$

Under normality assumptions for the two densities in the product, if I rearrange the two densities into the following forms:

$$f(Y | X, \dots) \rightarrow y_i \sim \mathcal{N}(A_i + B_i x_i, \sigma_e^2), \quad (\text{A2})$$

$$f(X | \dots) \rightarrow x_i \sim \mathcal{N}(C_i, \sigma_r^2), \quad (\text{A3})$$

where A_i , B_i , and C_i are any computation that does not include

X or Y , then I can solve for $f(X | Y, \dots)$ as follows:

$$\begin{aligned}
 f(X | Y, \dots) &\propto \exp\left\{-\frac{[y_i - (A_i + B_i x_i)]^2}{2\sigma_e^2}\right\} \times \exp\left\{-\frac{(x_i - C_i)^2}{2\sigma_f^2}\right\} \\
 &\propto \exp\left\{-\left\{\frac{[(y_i - A_i) - B_i x_i]^2}{2\sigma_e^2} + \frac{(x_i - C_i)^2}{2\sigma_f^2}\right\}\right\} \\
 &\propto \exp\left\{-\left[\frac{x_i^2 B_i^2 - 2x_i B_i (y_i - A_i) + \text{const}}{2\sigma_e^2} + \frac{x_i^2 - 2x_i C_i + \text{const}}{2\sigma_f^2}\right]\right\} \\
 &\propto \exp\left\{-\left[\frac{x_i^2 B_i^2 - 2x_i B_i (y_i - A_i)}{2\sigma_e^2} + \frac{x_i^2 - 2x_i C_i}{2\sigma_f^2}\right]\right\} \\
 &\propto \exp\left\{-\frac{x_i^2 B_i^2 \sigma_f^2 - 2x_i B_i \sigma_f^2 (y_i - A_i) + x_i^2 \sigma_e^2 - 2x_i \sigma_e^2 C_i}{2\sigma_e^2 \sigma_f^2}\right\} \\
 &\propto \exp\left\{-\frac{x_i^2 (B_i^2 \sigma_f^2 + \sigma_e^2) - 2x_i [B_i \sigma_f^2 (y_i - A_i) + \sigma_e^2 C_i]}{2\sigma_e^2 \sigma_f^2}\right\} \\
 &\propto \exp\left\{-\frac{x_i^2 - \{2x_i [B_i \sigma_f^2 (y_i - A_i) + \sigma_e^2 C_i] / (B_i^2 \sigma_f^2 + \sigma_e^2)\}}{2\sigma_e^2 \sigma_f^2 / (B_i^2 \sigma_f^2 + \sigma_e^2)}\right\} \\
 &\propto \exp\left\{-\frac{x_i^2 - \{2x_i [B_i \sigma_f^2 (y_i - A_i) + \sigma_e^2 C_i] / (B_i^2 \sigma_f^2 + \sigma_e^2)\} + \text{const}}{2\sigma_e^2 \sigma_f^2 / (B_i^2 \sigma_f^2 + \sigma_e^2)}\right\} \\
 &\propto \exp\left\{-\frac{(x_i - \{[B_i \sigma_f^2 (y_i - A_i) + \sigma_e^2 C_i] / (\sigma_f^2 B_i^2 + \sigma_e^2)\})^2}{2[\sigma_e^2 \sigma_f^2 / (B_i^2 \sigma_f^2 + \sigma_e^2)]}\right\} \\
 \therefore x_i &\sim \mathcal{N}\left(\frac{\sigma_f^2 B_i (y_i - A_i) + \sigma_e^2 C_i}{\sigma_f^2 B_i^2 + \sigma_e^2}, \frac{\sigma_e^2 \sigma_f^2}{\sigma_f^2 B_i^2 + \sigma_e^2}\right).
 \end{aligned} \tag{A4}$$

The above result illustrates that if I can rearrange the two densities into some form that maps onto Equations A2 and A3, then the product will result in a distribution that is also normally distributed. Note that the i subscript is both in the mean and the variance; therefore, the conditional distribution in Equation A4 is computed on every observation. However, this maintains a very general result that easily extends to common situations. For example, any moderator variables for the regression of Y on X can be included in the B term. Similarly, anything additive to X can be in A . As another example, if X is centered, then A would include the subtraction of the mean of X times the regression coefficient.

Chaining Products of Multiple Densities

Suppose now I have three variables, X , Y , and Z , with a fully factored distribution as follows:

$$f(Y | X, Z) \times f(Z | X) \times f(X), \tag{A5}$$

with the following three regressions as the models for each density.

$$\begin{aligned}
 f(Y | X, Z) &\rightarrow y_i \sim \mathcal{N}(\beta_0 + \beta_1 x_i + \beta_2 z_i + \beta_3 (x_i \times z_i), \sigma_Y^2), \\
 f(Z | X) &\rightarrow z_i \sim \mathcal{N}(\alpha_0 + \alpha_1 x_i, \sigma_Z^2), \\
 f(X) &\rightarrow x_i \sim \mathcal{N}(\gamma_0, \sigma_X^2).
 \end{aligned} \tag{A6}$$

If the goal is to obtain the distribution for $f(X | Y, Z)$ —for example, as required for a Gibbs sampler to impute X —we first break down this product into two separate steps.

1. Obtain $f(X | Y, Z) \propto f(Y | X, Z) \times f(X | Z)$.
2. Obtain $f(X | Z) \propto f(Z | X) \times f(X)$.

By breaking the problem into products of two distributions, I simplify the problem and perform the steps backwards to obtain the desired distribution. Starting with Step 2, I apply the result from Equation A4 to obtain the conditional distribution.

$$f(X | Z) \rightarrow x_i \sim \mathcal{N}\left(\frac{\sigma_X^2 \alpha_1 (z_i - \alpha_0) + \sigma_Z^2 \gamma_0}{\sigma_X^2 \alpha_1^2 + \sigma_Z^2}, \frac{\sigma_Z^2 \sigma_X^2}{\sigma_X^2 \alpha_1^2 + \sigma_Z^2}\right). \tag{A7}$$

Next, I take the above distribution and apply Equation A4 again in Step 1 to obtain a normal distribution

$$f(X | Y, Z) \rightarrow x_i \sim \mathcal{N}(\mu_{X|Y,Z}, \sigma_{X|Y,Z}^2), \tag{A8}$$

where the mean and variances are as follows:

Equation A9 (see below)

The above distribution may not have a nice compact form; however, using element-wise calculations, it is straightforward to compute for every observation, resulting in mean and variance vectors. These vectors can then be used to sample from the derived conditional distribution for every observation. More importantly, the described algorithm is straightforward to implement in code, thus allowing a wide range of models to have the conditional posterior predictive distributions automatically derived. I include an R implementation of this algorithm and demonstrate how to use functional programming techniques at <https://github.com/blimp-stats/Latent-Interactions>.

$$\begin{aligned}
 \mu_{X|Y,Z} &= \frac{[\sigma_Z^2 \sigma_X^2 / (\sigma_X^2 \alpha_1^2 + \sigma_Z^2)](\beta_1 + \beta_3 z_i)[y_i - (\beta_0 + \beta_2 z_i)] + \sigma_Y^2 \{[\sigma_X^2 \alpha_1 (z_i - \alpha_0) + \sigma_Z^2 \gamma_0] / (\sigma_X^2 \alpha_1^2 + \sigma_Z^2)\}}{[\sigma_Z^2 \sigma_X^2 / (\sigma_X^2 \alpha_1^2 + \sigma_Z^2)](\beta_1 + \beta_3 z_i)^2 + \sigma_Y^2}, \\
 \sigma_{X|Y,Z}^2 &= \frac{\sigma_Y^2 [\sigma_Z^2 \sigma_X^2 / (\sigma_X^2 \alpha_1^2 + \sigma_Z^2)]}{[\sigma_Z^2 \sigma_X^2 / (\sigma_X^2 \alpha_1^2 + \sigma_Z^2)](\beta_1 + \beta_3 z_i)^2 + \sigma_Y^2}.
 \end{aligned} \tag{A9}$$

(Appendices continue)

Appendix B

Full Blimp Syntax for Latent Moderation Example

```
# Data set's file name
DATA: data_cat.txt;

# Variable names in text data set
VARIABLES:
  consci:consci10 emosta1:emosta10
  orgcon1:orgcon10 cwb1:cwb23
  bin_orgcon cat_orgcon;

## Specify that the items are ordinal.
ORDINAL:
  consci:consci10 orgcon1:orgcon10
  cwb1:cwb23;

## Define latent variables
LATENT: consci orgcon cwb;

## Begin Modeling Syntax
MODEL:
  ## Latent Moderation Model
  cwb_model:
    # Explicitly write out the product between
    # the two latent variables
    cwb ~ consci orgcon consci*orgcon;

    # Fix residual variance to 1 for identification
    cwb ~~ cwb@1;

  ## Correlating predictor latent variables
  predictor_model:
    consci ~~ orgcon;
    consci ~~ consci@1;
    orgcon ~~ orgcon@1;

  ## Measurement Models for conscientiousness items
  measurement_consci:
    consci -> consci1@l_con consci2:consci10;

  ## Measurement Models for org constraints items
  measurement_orgcon:
    orgcon -> orgcon1@l_org orgcon2:orgcon10;

  ## Measurement Models for CWB items
  measurement_cwb:
    cwb -> cwb1@l_cwb cwb2:cwb23;

# Specify first loading as always positive
PARAMETERS:
  l_con ~ truncate(0, Inf);
  l_org ~ truncate(0, Inf);
  l_cwb ~ truncate(0, Inf);

## Conditional Effects Analysis with +/- 1 SD of orgcon
# mean = 0 by definition
# SD = 1 by definition
SIMPLE:
  consci | orgcon @ -1.0;
  consci | orgcon @ 0.0;
  consci | orgcon @ 1.0;

# Algorithmic Options
SEED: 298722; # PRNG Seed
CHAINS: 10; # Number of chains
BURN: 20000; # Burn-in iterations
ITER: 50000; # Post burn-in iterations
```

Received March 13, 2024
 Revision received May 5, 2025
 Accepted June 11, 2025 ■