

Approximate Post-Selective Inference for Regression with the Group LASSO

Snigdha Panigrahi

Peter W. MacDonald

Department of Statistics

University of Michigan

Ann Arbor, MI 48109-1107, USA

PSNIGDHA@UMICH.EDU

PWMACDON@UMICH.EDU

Daniel Kessler

Department of Statistics and Department of Psychiatry

University of Michigan

Ann Arbor, MI 48109-1107, USA

KESSLERD@UMICH.EDU

Editor: Daniela Witten

Abstract

After selection with the Group LASSO (or generalized variants such as the overlapping, sparse, or standardized Group LASSO), inference for the selected parameters is unreliable in the absence of adjustments for selection bias. In the penalized Gaussian regression setup, existing approaches provide adjustments for selection events that can be expressed as linear inequalities in the data variables. Such a representation, however, fails to hold for selection with the Group LASSO and substantially obstructs the scope of subsequent post-selective inference. Key questions of inferential interest—for example, inference for the effects of selected variables on the outcome—remain unanswered. In the present paper, we develop a consistent, post-selective, Bayesian method to address the existing gaps by deriving a likelihood adjustment factor and an approximation thereof that eliminates bias from the selection of groups. Experiments on simulated data and data from the Human Connectome Project demonstrate that our method recovers the effects of parameters within the selected groups while paying only a small price for bias adjustment.

Keywords: Conditional inference, Group sparsity, Group LASSO, Laplace approximation, Selective inference

1. Introduction

Modern statistical analysis of complex data does not always fit into the classical inferential framework. Instead, analysis splits into two distinct stages: a *selection* stage, in which we formulate a model and hypotheses of interest; and an *inference* stage, in which we estimate parameters, quantify uncertainties, and test hypotheses under our selected model. However, classical coverage guarantees for credible and confidence intervals fail dramatically when data used for selection is naively re-used for inference; see Berk et al. (2013); Lee et al. (2016); Benjamini (2020) and references therein. Simple procedures like data splitting preserve validity of post-selective inference if two independent subsets of data are used for the selection and inference stages. However, discarding all the data used in the selection stage is inefficient, and there is potential for methodology which can safely reuse a portion

of the information from selection for valid inference. By adopting a conditional approach, recent tools in selective inference reduce this wastefulness when selection algorithms are applied to data prior to statistical modeling and inference. As examples, conditional methods by Suzumura et al. (2017); Zhao and Panigrahi (2019); Gao et al. (2020); Tanizaki et al. (2020) provide adjustments for selection bias in different post-selective inference tasks.

To briefly outline the essence of the conditional approach, consider a variable selection algorithm applied to data Y with p (fixed) predictors X . Suppose the algorithm returns as output $\hat{E}(Y)$, a subset of $\{1, 2, \dots, p\}$ such that each index represents a variable (column of X), and therefore $\hat{E}(Y)$ is associated with a model selected from 2^p possibilities. After selecting a given (non-empty) subset of variables E , our interest lies in inference for a set of post-selective parameters

$$\Theta_E = \left\{ \theta_E^{(j)} \in \mathbb{R}, j \in E \right\}$$

using the observed data $\{Y = y\}$. Post-selective inference for Θ_E proceeds by conditioning on the selection event

$$\left\{ \hat{E}(Y) = E \right\},$$

which is motivated by the fact that conditional coverage implies unconditional coverage under selection. That is, for a set $C\left(\hat{E}(Y), Y\right) \subseteq \mathbb{R}$ that depends on both the output of selection and the data, Lee et al. (2016) note

$$\mathbb{P}\left(\theta_{\hat{E}(Y)}^{(j)} \in C\left(\hat{E}(Y), Y\right) \mid \hat{E}(Y) = E\right) \geq 1 - \alpha \Rightarrow \mathbb{P}\left(\theta_{\hat{E}(Y)}^{(j)} \in C\left(\hat{E}(Y), Y\right)\right) \geq 1 - \alpha.$$

In many problems, it may be more convenient to instead condition on $\mathcal{A}_E$, where $\mathcal{A}_E \subseteq \left\{ \hat{E}(Y) = E \right\}$. Inference remains valid by the argument above even when conditioning on a proper subset of the selection event.

After conditioning on $\mathcal{A}_E$, post-selective inference may be carried out via either a frequentist or Bayesian framework. A Bayesian framework in Yekutieli (2012); Panigrahi et al. (2021) relies on a conditional, *selection-informed* likelihood to facilitate posterior sampling. The Bayesian approach is especially useful for inferring about vector-valued parameters or functions thereof, and permits flexible inference in different models informed by selection, for instance, models with unknown noise variance. In the remainder of this paper, we develop an approximate Bayesian method for post-selective inference with the Group LASSO and several of its variants. The setup for our problem is the following: (i) the covariates act naturally in groups known a priori in the analysis; (ii) only a few of these groups of covariates affect the outcome, captured effectively by a parsimonious model. A well-developed class of algorithms in Yuan and Lin (2005); Jacob et al. (2009); Simon et al. (2013) among others exploits this knowledge about the covariate space in order to select regression models with grouped covariates. Post-selective inference in the resulting selection-informed models is a natural next step that is addressed by our method.

We structure our paper as follows. We begin by situating the contributions of our method in the post-selective literature in Section 2. In Section 3, we present a selection-informed posterior that serves as the methodological centerpiece of our Bayesian framework. In Section 4, we obtain an exact value for a likelihood adjustment factor in our selection-informed posterior to eliminate bias from the selection of groups. We then apply a generalized version

of Laplace-type approximations to obtain feasible sampling updates from an approximate version of the posterior. In Section 5, we generalize our method to models informed by different forms of grouped covariates. We establish large-sample theory for our approximate Bayesian methods in Section 6. We demonstrate the potential of our methods in numerical experiments and in a human neuroimaging application in Section 7. Proofs for our technical results and further supporting information are included in the appendices.

2. Related Work and Contributions

Below, we identify the challenges that preclude the use of existing methods for post-selective inference, and their immediate modifications, for the Group LASSO. Fixing some notation, suppose we observe n independent instances of a scalar response variable Y_i and a p -dimensional vector of covariates X_i for $i = 1, \dots, n$. We denote the response vector by $y = (Y_1 \dots Y_n)^\top \in \mathbb{R}^n$ and the corresponding covariate matrix by $X = (X_1 \dots X_n)^\top \in \mathbb{R}^{n \times p}$. Let $\mathcal{G}$ be a prespecified partition of our p covariates into G groups. We refer to a group in $\mathcal{G}$ by lowercase g , and use $|g| \in \mathbb{N}$ to denote the number of covariates within group g .

For now, we consider non-overlapping groups defined by the partition $\mathcal{G}$. Suppose, we solve the familiar Group LASSO objective in Yuan and Lin (2005):

$$\hat{\beta}^{(\mathcal{G})} \in \underset{\beta}{\operatorname{argmin}} \frac{1}{2} \|y - X\beta\|_2^2 + \sum_{g \in \mathcal{G}} \lambda_g \|\beta_g\|_2. \quad (1)$$

For each group $g \in \mathcal{G}$, $\beta_g \in \mathbb{R}^{|g|}$ is a vector with entries corresponding to the covariates in group g , and $\lambda_g \geq 0$ is a tuning parameter for this group. The solution of (1) returns a subset of the covariates

$$\hat{E}(y) = \operatorname{supp}(\hat{\beta}^{(\mathcal{G})}),$$

where the support of the Group LASSO estimator respects the prespecified groups. Specifically, the selected set of covariates can be written as a union of selected groups in $\mathcal{G}$, which we denote by $\mathcal{G}_{\hat{E}}$ in the paper.

2.1 From Atoms to Groups

In the special case when each covariate forms an atomic group of size 1, the objective in (1) agrees with the widely studied LASSO. Established in Lee et al. (2016), the selection of any subset of atoms is a polyhedral event, which means that the event is expressible as a union of linear inequalities in the response vector y . Existing methods for post-selective inference in Lee et al. (2016); Suzumura et al. (2017); Liu et al. (2018) readily adjust for bias from selection by reducing the polyhedral conditioning event to univariate truncations. However, when we transition from atoms to nontrivial groups, the selection of promising groups no longer results in polyhedral events. We visualize this fact through Figure 1 in a simple example, when the sample size and the number of predictors are both equal to 2.

In Figure 1, we contrast the geometry of the selection event for the LASSO and the Group LASSO. Our covariates are the columns of an identity matrix and the tuning parameters are set to be 1. Under the grouped scenario, the two orthogonal covariates comprise

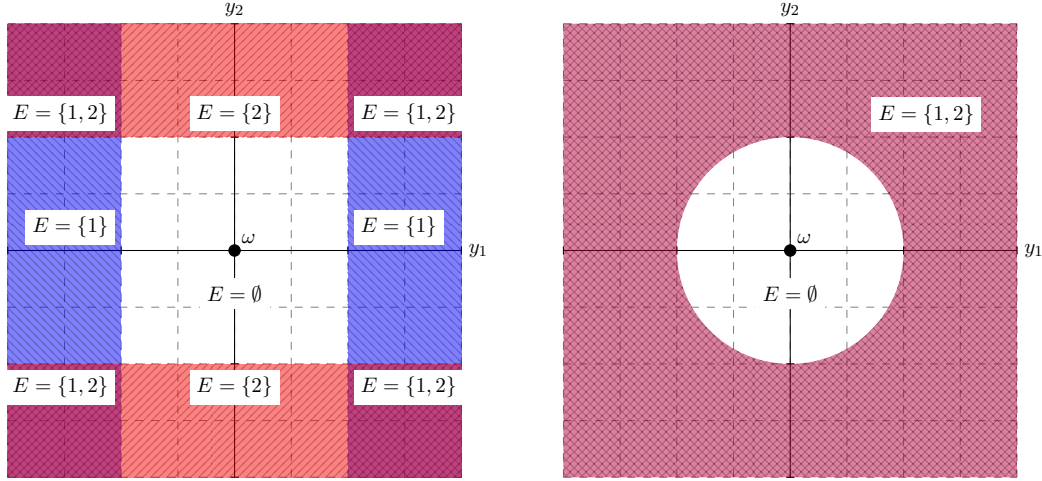


Figure 1: Geometry of selection events for the LASSO (left) and the Group LASSO (right) as a function of (y_1, y_2) . In the case of no randomization, the origin ω is the point $(0, 0)$, but see text surrounding (3) for discussion of how randomization affects the origin. The LASSO can select any of $\emptyset, \{1\}, \{2\}$, or $\{1, 2\}$ as predictors; the Group LASSO can select $\emptyset$ or $\{1, 2\}$.

a single group, whereas in the LASSO each covariate is an atom. For the LASSO, the event leading to the selection of the active set E is a union of rectangular regions in the plane that are highlighted by the same color. A proper subset of this event is obtained by further restricting the signs of selected covariates to match the observed signs. This proper subset leads to one of the rectangular regions in the plane; see left panel. In contrast, the selection of an active group for the Group LASSO is depicted as the complement of a ball in the right panel, which can no longer be characterized as a union of polyhedral events.

2.2 Post-selective Inference for Overall Group Effects

We now turn to recent results by Loftus and Taylor (2015); Yang et al. (2016) which provide post-selective inference for overall effects of groups after solving the Group LASSO. Introducing some more notation, let $\mathcal{W}$ represent an operator that maps the vector v to the unit vector $(\|v\|_2)^{-1} \cdot v$. For a linear subspace $\mathbb{S} \subseteq \mathbb{R}^n$ and its orthogonal complement $\mathbb{S}^\perp$, let $\mathcal{P}_{\mathbb{S}}$ and $\mathcal{P}_{\mathbb{S}^\perp}$ denote the projection operators onto the subspaces $\mathbb{S}$ and $\mathbb{S}^\perp$ respectively.

Consider solving the Group LASSO in (1). Let E be the realized value of $\hat{E}$. Suppose we assume the saturated model: $y \sim \mathcal{N}_n(\mu, \sigma^2 I_n)$ for inference. For $g \in \mathcal{G}_E$ and the subspace $\mathbb{S}_{g,E} = \text{span} \left(\mathcal{P}_{X_{\mathcal{G}_E \setminus g}^\perp} (X_g) \right)$, consider the post-selective parameter

$$\mu_g = \|\mathcal{P}_{\mathbb{S}_{g,E}}(\mu)\|_2 \in \mathbb{R}$$

after selection with the Group LASSO. A significant p -value under the null hypothesis $H_{0,g} : \mu_g = 0$ confirms the presence of the selected group g in the estimated support;

confidence bounds for μ_g measure the overall effect of the selected group g . The main result by Yang et al. (2016) allows post-selective inference for μ_g through a conditional distribution for $\|\mathcal{P}_{\mathbb{S}_{g,\hat{E}}}(y)\|_2$, which we revisit in the following lemma.

Lemma 1 Yang et al. (2016). *Conditional upon the event*

$$\left\{ y : \hat{E}(y) = E, \mathcal{W}\left(\mathcal{P}_{\mathbb{S}_{g,\hat{E}}}(y)\right) = U_g, \mathcal{P}_{\mathbb{S}_{g,\hat{E}}}^\perp(y) = W_g \right\}, \quad (2)$$

the density for $\|\mathcal{P}_{\mathbb{S}_{g,\hat{E}}}(y)\|_2$ at γ_g is proportional to:

$$\gamma_g^{|g|-1} \cdot \exp\left(-\frac{1}{2\sigma^2}(\gamma_g^2 - 2\gamma_g \cdot U_g^\top \mu)\right) \cdot 1_{\mathcal{R}_E}(\gamma_g),$$

where $\mathcal{R}_E = \left\{ \gamma_g \in \mathbb{R}^+ : \hat{E}(U_g \gamma_g + W_g) = E \right\}$.

Applying a probability integral transform to the conditional law in Lemma 1 produces a pivot for

$$\tilde{\mu}_g = U_g^\top \mu$$

conditional upon (2). In particular, $\tilde{\mu}_g$ agrees with μ_g under the null $H_{0,g}$. As a result, a pivot for the former parameter yields a valid p -value for testing $H_{0,g}$ and coincides with the p -value in Loftus and Taylor (2015). The two parameters, however, do not coincide in general. Instead, the following relation holds by Cauchy-Schwarz:

$$\mu_g \geq \tilde{\mu}_g,$$

and inverting the pivot thus provides a conservative, lower confidence bound for μ_g .

As emphasized in the preceding discussion in Section 2, the selection of groups is no longer a polyhedral event. Indeed, the difficulties posed by the non-polyhedral geometry for the Group LASSO continue to persist; we note that the truncating region $\mathcal{R}_E$ in Lemma 1 lacks a closed-form description. The outlined approach overcomes this barrier to some extent by narrowing down the scope of inferential targets to conducting inference on overall group effects, in which case one only needs to explore a positive half-line to approximately compute $\mathcal{R}_E$. Besides lacking an upper confidence bound for overall group effects, the existing approach does not yield interval estimates for the effects of the individual variables in the selected groups, nor does it identify a joint distribution for the individual effects.

2.3 Our Method

Closing existing gaps, we develop a Bayesian method for post-selective inference after conducting a randomized selection of groups. Our method accounts for the non-polyhedral selection of groups via a likelihood adjustment factor and characterizes a *selection-informed* posterior distribution based on the likelihood adjustment. Working with a selection-informed posterior grants us the flexibility to estimate the individual effects within selected groups and functions thereof through credible regions and general posterior expectations. At the same time, a randomized selection of groups permits us a very simple and exact charac-

terization for the truncating region in the conditional likelihood that makes subsequent inference easily feasible.

The randomizing variable, or *randomization*, used for the selection of groups is a Gaussian variable throughout the remainder of the paper and is hereafter termed Gaussian randomization. Our methods based on Gaussian randomization are closely related to data carving proposals in Fithian et al. (2014); Panigrahi et al. (2021); Panigrahi (2018); Schultheiss et al. (2021), wherein selection operates only on a subset of the samples, but subsequent post-selective estimation uses the full data. The variance of the Gaussian randomization is a tuning parameter analogous to the split proportion in data splitting; this provides us control of the relative amount of information used in selecting a group-sparse model versus estimating the post-selective parameters. The information borrowed by our approach from selection yields credible intervals which are shorter than the corresponding interval estimates for data splitting with roughly the same information split.

In the following sections, we develop our method in two steps. First, we account for the selection of groups, a non-polyhedral event, via an exact likelihood adjustment factor. Rather than characterizing the non-polyhedral event in the space of the data and randomization variables, we develop a change of variables that is motivated by ideas in Tian et al. (2016). Our choice of conditioning event is characterized by simple sign constraints in the new variables which yields us a selection-informed posterior distribution. In the next step, we propose a computationally feasible surrogate for this posterior distribution with a (generalized) Laplace approximation. Our Bayesian method delivers statistically consistent estimates using a selection-informed posterior distribution for the group-sparse parameter vector. Continuing with our simple grouped example introduced earlier and depicted in the right panel of Figure 1, Figure 2 serves to preview the distribution of 1400 samples from our surrogate selection-informed posterior by varying the number of observations $n = 25, 50, 100, 250, 500, 1000$. Assuredly, as n increases, the support of the posterior concentrates around the true bivariate parameter, suggesting the statistical consistency of our method which we justify theoretically in Section 6.

3. Framework for Selection-informed Inference

Consistent with a post-selective setting, under a fixed X regression, our problem proceeds in two stages. First, we select promising groups by optimizing an objective inducing grouped sparsity. Then, we specify a group-sparse linear model which is informed by the groups of covariates selected from the previous stage. We begin by describing our methods for non-overlapping groups, based on a prespecified partition $\mathcal{G}$ of p covariates into G groups. In Section 5, we present a larger category of grouped sparsities that our methods successfully encompass.

Using notation defined in Section 2, we consider the Group LASSO objective in (1) with an added randomization term:

$$\hat{\beta}^{(\mathcal{G})} \in \underset{\beta}{\operatorname{argmin}} \left\{ \frac{1}{2} \|y - X\beta\|_2^2 + \sum_{g \in \mathcal{G}} \lambda_g \|\beta_g\|_2 - \omega^\top \beta \right\}. \quad (3)$$

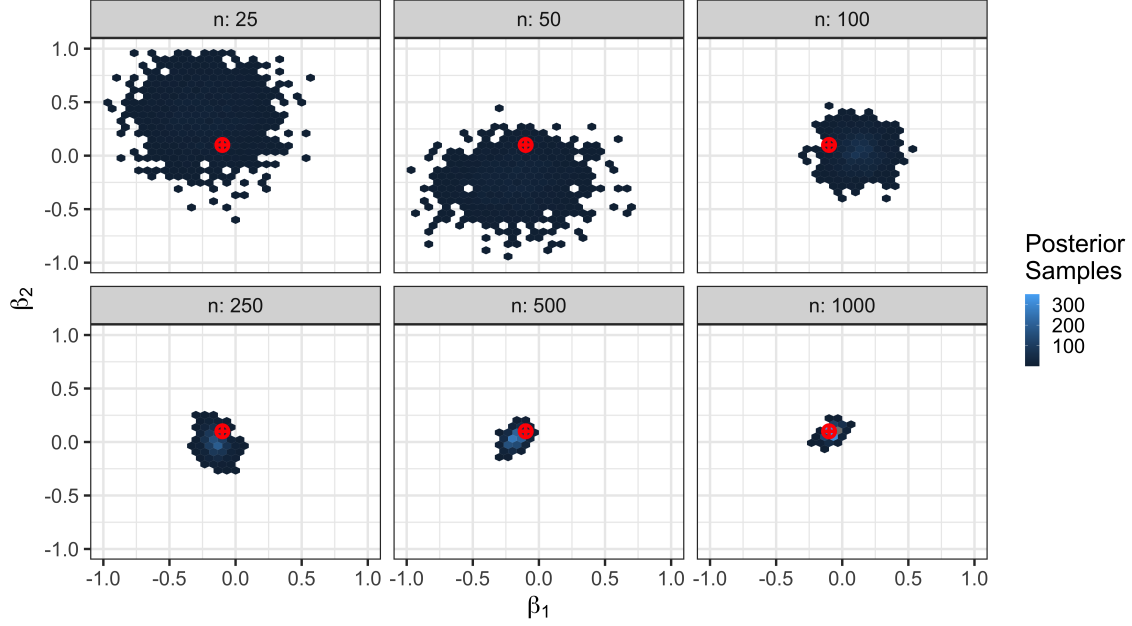


Figure 2: Distribution of 1400 samples from a surrogate of our selection-informed posterior based on a bivariate model with a single group of orthogonal covariates for varying n . The color of each hex depicts the number of posterior samples drawn in that region, while a red crosshair indicates the location of the true value of the parameter $\beta = (-0.1 \ 0.1)^\top$.

In the final term of this objective, $\omega \sim \mathcal{N}_p(0, \Omega)$ is a Gaussian randomization variable independent of the data. As indicated previously, perturbing the optimization problem with a Gaussian randomization variable ω introduces a tradeoff between selection and inference, giving the user the ability to reserve some information from the selection stage to perform inference. Additional discussion of the role of randomization in (3), and the relation of randomization variance with data splitting is given in Section 7. Hereafter, focusing on the solution of (3), we let

$$\hat{E} = \text{supp}(\hat{\beta}^{(\mathcal{G})})$$

be the support of the randomized Group LASSO estimator and let $\mathcal{G}_{\hat{E}}$ be the selected groups of covariates according to the estimated support.

Revisiting the example in Section 2 and the related Figure 1, we note that the selection regions have a similar geometry with the added randomization: the randomization instance ω merely shifts the origin in both panels of the figure. Elaborating on the example, suppose that $n = p = 2$, $X = I_n$, the identity matrix and $\lambda_g = \lambda = 1$. Let $\omega \sim \mathcal{N}_2(0, I_2)$. The stationary mapping for the optimization in (3) is given by:

$$y + \omega = \hat{\beta}^{(\mathcal{G})} + z,$$

where the final term is the subgradient of the Group LASSO penalty evaluated at the solution. In the case that the single group of two covariates is not selected, $\hat{\beta}^{(g)} = 0$ and $\|z\|_2 < 1$. For any fixed ω , the collection of y that leads to no selection is equivalent to a ball centered at ω . The complement of this region characterizes the selection of the group of size 2. Instead, when we have two atoms (groups with size 1 each), i.e., we solve a randomized version of the LASSO in Tian and Taylor (2018), then the selection of a subset of covariates with fixed signs is equivalent to linear inequalities in y and ω . Once again, shifting the origin in the left panel of Figure 1 to ω depicts the polyhedral selection event for the randomized LASSO. Recent work by Panigrahi et al. (2021); Panigrahi and Taylor (2022); Panigrahi et al. (2022) provide a likelihood after the randomized LASSO; but, these methods are not applicable to the present problem, because the selection of groups with size greater than 1 does not admit a polyhedral form.

In the next stage, we specify a model after selection. Letting E be the realized value of $\hat{E}$, we model our response as

$$y \sim \mathcal{N}_n(X_E \beta_E, \sigma^2 I_n). \quad (4)$$

The selected model in (4), using the solution of (3), may indeed be misspecified. Suppose, the true distribution for our response is

$$y \sim \mathcal{N}_n(X \bar{\beta}, \bar{\sigma}^2 I_n),$$

for some $\bar{\beta} \in \mathbb{R}^p$, $\bar{\sigma}^2 \in \mathbb{R}^+$. Under the true distribution, our method always delivers inference for the best linear representation of the response mean using the selected covariates X_E , regardless of model misspecification. We elaborate further on this point when we turn to a selection-informed likelihood based on the model in (4).

By analogy with Yekutieli (2012); Panigrahi and Taylor (2018), we pose a selection-informed prior for our post-selective parameter

$$\beta_E \sim \pi_E \quad (5)$$

to invoke a Bayesian framework after selection. Both the selected model and the selection-informed prior depend on the observed data. However, they do so only through the selection event accounted for by conditioning.

Without loss of generality, we are able to assume that the variance parameter σ is known. Following the lines of Panigrahi et al. (2021), the Bayesian approach we take easily accommodates the case of unknown variance by treating it as a parameter. In this case, we can proceed by posing a joint selection-informed prior on β_E and σ .

3.1 Selection-informed Posterior

In this section we define a selection-informed posterior using the model for y in (4) and the prior in (5). Through the remainder of the paper, we use the notation $p(\mu, \Sigma; b)$ for a multivariate normal density function with mean μ and covariance Σ evaluated at b . To state the selection-informed posterior, we will use the following variables that are involved in selection: (i) the randomization variable ω ; (ii) the least squares estimate based on (X_E, y)

$$\hat{\beta}_E = (X_E^T X_E)^{-1} X_E^T y;$$

(iii) the orthogonal projection $N_E = X^\top (I_n - X_E (X_E^\top X_E)^{-1} X_E^\top) y$, assuming X_E is full rank. Under the selected model, $\hat{\beta}_E$ has mean β_E , and N_E has mean 0. Denote the covariance of $\hat{\beta}_E$ by $\Sigma_E = \sigma^2 (X_E^\top X_E)^{-1}$ and let Ψ_E be the covariance of N_E . Ignoring selection, the usual joint likelihood for these three variables is given by:

$$p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot p(0, \Psi_E; N_E) \cdot p(0, \Omega; \omega).$$

The factorization follows directly from the independence between $\hat{\beta}_E$ and N_E and their independence with the randomization variable ω .

Accounting for the selection-informed nature of our model, the likelihood we work with conditions upon an event:

$$\mathcal{A}_E \subseteq \{(\hat{\beta}_E, N_E, \omega) : \hat{E} = E\}.$$

The conditioning event $\mathcal{A}_E$ for the group-sparse problem is a proper subset of the selection event $\{\hat{E} = E\}$, which we define precisely in Theorem 2. After truncating realizations to the event $\mathcal{A}_E$, the resulting conditional likelihood is proportional to

$$p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot p(0, \Psi_E; N_E) \cdot p(0, \Omega; \omega) \cdot \mathbf{1}_{\mathcal{A}_E}(\hat{\beta}_E, N_E, \omega).$$

Now we state our selection-informed likelihood, derived after conditioning further upon the ancillary statistic N_E and integrating out the randomization variable ω . Up to proportionality in β_E , the expression for this likelihood agrees with

$$\{\mathbb{P}(\mathcal{A}_{E;N_E} \mid \beta_E)\}^{-1} \cdot p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot \int p(0, \Omega; \omega) \cdot \mathbf{1}_{\mathcal{A}_{E;N_E}}(\hat{\beta}_E, \omega) d\omega; \quad (6)$$

$\mathcal{A}_{E;N_E}$ is the set of $\hat{\beta}_E, \omega$ that result in the event $\mathcal{A}_E$ for the fixed instance N_E , and

$$\mathbb{P}(\mathcal{A}_{E;N_E} \mid \beta_E) = \int p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot p(0, \Omega; \omega) \cdot \mathbf{1}_{\mathcal{A}_{E;N_E}}(\hat{\beta}_E, \omega) d\omega d\hat{\beta}_E,$$

where $\mathbb{P}(\cdot \mid \beta_E)$ highlights the dependence of the probability for the event $\mathcal{A}_{E;N_E}$ on the post-selective parameters β_E .

More generally, the selection-informed likelihood in (6) yields us inference for the best linear representation of the response mean in terms of the selected covariates. To note this generality, suppose that our response is generated from the linear model: $y \sim \mathcal{N}_n(X\bar{\beta}, \sigma^2 I_n)$. For any fixed set E with size $|E|$, we have $\hat{\beta}_E \sim \mathcal{N}_{|E|}(\beta_E, \Sigma_E)$ and $N_E \sim \mathcal{N}_p(\xi_E, \Psi_E)$ where

$$\beta_E = \operatorname{argmin}_{b \in \mathbb{R}^{|E|}} \|X\bar{\beta} - X_E b\|_2^2,$$

and $\xi_E = X^\top \mathcal{P}_{\mathbb{S}_E^\perp}(X\bar{\beta})$ for $\mathbb{S}_E = \operatorname{span}(X_E)$. The likelihood for $\hat{\beta}_E, N_E$ and ω factorizes as

$$p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot p(\xi_E, \Psi_E; N_E) \cdot p(0, \Omega; \omega).$$

Treating ξ_E as nuisance parameters post selection, we condition on N_E to obtain a likelihood function of β_E , free from nuisance parameters. It is easy to see that our selection-informed

likelihood assumes the expression in (6) and inference proceeds identically, regardless of model misspecification.

Using (6) in conjunction with our selection-informed prior (5) ultimately yields us our selection-informed posterior distribution for β_E :

$$\pi_E(\beta_E \mid \widehat{\beta}_E, N_E) \propto (\mathbb{P}(\mathcal{A}_{E;N_E} \mid \beta_E))^{-1} \cdot \pi_E(\beta_E) \cdot p(\beta_E, \Sigma_E; \widehat{\beta}_E). \quad (7)$$

Evaluating $\mathbb{P}(\mathcal{A}_{E;N_E} \mid \beta_E)$, called the likelihood *adjustment factor* in Panigrahi et al. (2021), is the prime technical hurdle in carrying out selection-informed Bayesian inference. By carefully choosing the conditioning event $\mathcal{A}_{E;N_E}$, we develop mathematical expressions for the adjustment factor and update estimates in a feasible analytic form for the group-sparse problem. We take this up in the next section.

4. Selection-informed Bayesian Methods

We begin by identifying an exact theoretical value for the likelihood adjustment factor in our selection-informed posterior (7). With a slight abuse of notation, hereon, we denote the event $\mathcal{A}_{E;N_E}$ by $\mathcal{A}_E$ and the related adjustment factor by $\mathbb{P}(\mathcal{A}_E \mid \beta_E)$.

4.1 Likelihood Adjustment Factor

We return to our primary case study of the (non-overlapping) Group LASSO. We express the non-zero solution for a selected group, $g \in \mathcal{G}_E$, as

$$\widehat{\beta}_g^{(\mathcal{G})} = \|\widehat{\beta}_g^{(\mathcal{G})}\|_2 \cdot \frac{\widehat{\beta}_g^{(\mathcal{G})}}{\|\widehat{\beta}_g^{(\mathcal{G})}\|_2} = \gamma_g u_g, \quad (8)$$

where $\gamma_g = \|\widehat{\beta}_g^{(\mathcal{G})}\|_2 > 0$ is a scalar representing the size of the selected group and $u_g = (\|\widehat{\beta}_g^{(\mathcal{G})}\|_2)^{-1} \cdot \widehat{\beta}_g^{(\mathcal{G})}$ is a unit vector in $\mathbb{R}^{|g|}$ satisfying $\|u_g\|_2 = 1$. The stationary mapping for (3) at the solution is given by:

$$\omega = X^\top X \begin{pmatrix} (\gamma_g u_g)_{g \in \mathcal{G}_E} \\ 0 \end{pmatrix} - X^\top X_E \widehat{\beta}_E - N_E + \begin{pmatrix} (\lambda_g u_g)_{g \in \mathcal{G}_E} \\ (\lambda_g z_g)_{g \in -\mathcal{G}_E} \end{pmatrix}.$$

We note that $((\lambda_g u_g)_{g \in \mathcal{G}_E}^\top \ (\lambda_g z_g)_{g \in -\mathcal{G}_E}^\top)^\top$ is the subgradient of the ℓ_2 -norm Group LASSO penalty at the solution, where z_g is a vector in $\mathbb{R}^{|g|}$ satisfying $\|z_g\|_2 < 1$ for each non-selected group $g \in -\mathcal{G}_E$. We collect the following optimization variables:

$$\widehat{\gamma} = (\gamma_g : g \in \mathcal{G}_E)^\top \in \mathbb{R}^{|\mathcal{G}_E|}, \quad \widehat{\mathcal{U}} = \{u_g : g \in \mathcal{G}_E\}, \quad \widehat{\mathcal{Z}} = \{z_g : g \in -\mathcal{G}_E\},$$

calling their respective realizations γ , $\mathcal{U}$ and $\mathcal{Z}$. Letting $\text{diag}(\cdot)$ operate on an ordered collection of matrices and return the corresponding block diagonal matrix, we fix $U = \text{diag}((u_g)_{g \in \mathcal{G}_E})$. Then, based on the stationary mapping from the Group LASSO, define

$$\phi_{\widehat{\beta}_E}(\widehat{\gamma}, \widehat{\mathcal{U}}, \widehat{\mathcal{Z}}) = A \widehat{\beta}_E + B(\widehat{\mathcal{U}}) \widehat{\gamma} + c(\widehat{\mathcal{U}}, \widehat{\mathcal{Z}}),$$

where

$$A = -X^\top X_E, \quad B = X^\top X_E U, \quad c = -N_E + \left((\lambda_g u_g)_{g \in \mathcal{G}_E}^\top \quad (\lambda_g z_g)_{g \in -\mathcal{G}_E}^\top \right)^\top.$$

Theorem 2 gives us the expression for the likelihood adjustment factor after applying the change of variables:

$$\omega \rightarrow (\hat{\gamma}, \hat{\mathcal{U}}, \hat{\mathcal{Z}}), \quad \text{where} \quad (\hat{\gamma}, \hat{\mathcal{U}}, \hat{\mathcal{Z}}) = \phi_{\hat{\beta}_E}^{-1}(\omega). \quad (9)$$

For each $g \in \mathcal{G}_E$, we construct the orthonormal basis completion for u_g , which we denote by $\bar{U}_g \in \mathbb{R}^{|g| \times |g|-1}$. Further, $x > t$ for $x \in \mathbb{R}^k$ and $t \in \mathbb{R}$ simply means that the inequality holds in a coordinate-wise sense.

Theorem 2 *Consider the conditioning event*

$$\mathcal{A}_E = \{\hat{E} = E, \hat{\mathcal{U}} = \mathcal{U}, \hat{\mathcal{Z}} = \mathcal{Z}\}.$$

Define the following matrices:

$$\bar{U} = \text{diag} \left((\bar{U}_g)_{g \in \mathcal{G}_E} \right), \quad \Gamma = \text{diag} \left((\hat{\gamma}_g I_{|g|-1})_{g \in \mathcal{G}_E} \right), \quad \Lambda = \text{diag} \left((\lambda_g I_{|g|})_{g \in \mathcal{G}_E} \right).$$

Then, we have

$$\begin{aligned} \mathbb{P}(\mathcal{A}_E \mid \beta_E) \propto & \int \int p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot \exp \left\{ -\frac{1}{2} (A\hat{\beta}_E + B\hat{\gamma} + c)^\top \Omega^{-1} (A\hat{\beta}_E + B\hat{\gamma} + c) \right\} \\ & \times J_{\phi_{\hat{\beta}_E}}(\hat{\gamma}; \mathcal{U}, \mathcal{Z}) \cdot \mathbf{1}(\hat{\gamma} > 0) d\hat{\beta}_E d\hat{\gamma}, \end{aligned}$$

where

$$J_{\phi_{\hat{\beta}_E}}(\hat{\gamma}; \mathcal{U}, \mathcal{Z}) = \det \left(\Gamma + \bar{U}^\top (X_E^\top X_E)^{-1} \Lambda \bar{U} \right). \quad (10)$$

The non-polyhedral event we set out to analyze is characterized exactly up to proportionality through the adjustment factor in Theorem 2. This exact characterization is possible due to the choice of conditioning event as well as the specific form of randomization. Drawing an analogy to the conditioning event for the LASSO in Lee et al. (2016), conditioning on $\hat{\mathcal{U}} = \mathcal{U}$ is similar to their required conditioning on the sign of each selected coefficient, where we interpret the sign as the univariate special case of vector direction. By conditioning further upon $\hat{\mathcal{Z}} = \mathcal{Z}$, we avoid an integration over $p - |E|$ variables. Furthermore, the specific form of randomization in (3) allows us to describe our conditioning event $\mathcal{A}_E$ only as a set of simple sign constraints on $\hat{\gamma}$, the sizes (ℓ_2 norms) of the selected groups.

In our likelihood adjustment factor, $J_{\phi_{\hat{\beta}_E}}(\hat{\gamma}; \mathcal{U}, \mathcal{Z})$ represents the Jacobian associated with the change of variables $\phi_{\hat{\beta}_E}(\cdot)$. Noticing the dependence of this function on simply $\hat{\gamma}$ and the conditioned-upon $\mathcal{U}$, we call this function $J_\phi(\cdot, \mathcal{U})$. In the special case where the design matrix of the selected model is orthogonal, the Jacobian takes a much simpler form which is given in Corollary 3.

Corollary 3 *Suppose $X_E^\top X_E = I_{|E|}$. Then*

$$J_\phi(\gamma, \mathcal{U}) = \prod_{g \in \mathcal{G}_E} (\gamma_g + \lambda_g)^{(|g|-1)}.$$

The proof of Corollary 3 is a direct calculation based on (10) and is omitted.

In comparison with the adjustment for polyhedral selection events in Panigrahi et al. (2021), the selection of groups leads to a nontrivial Jacobian function $J_\phi(\cdot, \mathcal{U})$ in our likelihood adjustment. It is easy to note that the Jacobian dissolves as a constant when all the selected groups are atoms with sizes exactly equal to 1. Based on the event $\mathcal{A}_E$, our likelihood adjustment factor in this special case (with a constant Jacobian) gives an adjustment for the randomized LASSO.

4.2 Surrogate Selection-informed Posterior

Plugging the adjustment factor from Theorem 2 into (7) gives us the selection-informed posterior. Proposition 4 simplifies the expression for this posterior, and further expresses it in terms of Gaussian densities. We defer the details for matrices $\bar{A}$, $\bar{R}$, $\bar{b}$, $\bar{s}$, $\bar{\Theta}$, and $\bar{\Omega}$, which have closed forms, which do not depend on β_E or $\hat{\gamma}$, to the appendices.

Proposition 4 *Conditioned on the event $\mathcal{A}_E$ in Theorem 2, the selection-informed posterior in (7) agrees with*

$$\left(\int J_\phi(\hat{\gamma}, \mathcal{U}) \cdot p(\bar{R}\beta_E + \bar{s}, \bar{\Theta}; \hat{\beta}_E) \cdot p(\bar{A}\hat{\beta}_E + \bar{b}, \bar{\Omega}; \hat{\gamma}) \cdot \mathbf{1}(\hat{\gamma} > 0) d\hat{\gamma} d\hat{\beta}_E \right)^{-1} \\ \times \pi_E(\beta_E) \cdot p(\bar{R}\beta_E + \bar{s}, \bar{\Theta}; \hat{\beta}_E).$$

Bypassing integrations, we provide deterministic expressions for a surrogate selection-informed posterior and the corresponding gradient in Theorem 5. Let

$$\beta_E^*, \gamma^* = \underset{\tilde{\beta}_E, \tilde{\gamma}}{\operatorname{argmin}} \left\{ \frac{1}{2} (\tilde{\beta}_E - \bar{R}\beta_E - \bar{s})^\top (\bar{\Theta})^{-1} (\tilde{\beta}_E - \bar{R}\beta_E - \bar{s}) \right. \\ \left. + \frac{1}{2} (\tilde{\gamma} - \bar{A}\tilde{\beta}_E - \bar{b})^\top (\bar{\Omega})^{-1} (\tilde{\gamma} - \bar{A}\tilde{\beta}_E - \bar{b}) + \operatorname{Barr}(\tilde{\gamma}) \right\},$$

where $\operatorname{Barr}(\cdot)$ is a barrier penalty (Auslender, 1999) that takes the value ∞ when the support constraints are violated and imposes a smaller penalty for values farther away from the boundary of the positive orthant. We use C to denote a constant free of β_E . We apply a generalized version of the Laplace approximation (Wong, 2001; Inglet and Majerski, 2014) to the normalizing constant in Proposition 4. This approximation, given by

$$\int J_\phi(\hat{\gamma}, \mathcal{U}) \cdot p(\bar{R}\beta_E + \bar{s}, \bar{\Theta}; \hat{\beta}_E) \cdot p(\bar{A}\hat{\beta}_E + \bar{b}, \bar{\Omega}; \hat{\gamma}) \cdot \mathbf{1}(\hat{\gamma} > 0) d\hat{\gamma} d\hat{\beta}_E \\ \approx C \cdot J_\phi(\gamma^*, \mathcal{U}) \cdot \exp \left(-\frac{1}{2} (\beta_E^* - \bar{R}\beta_E - \bar{s})^\top (\bar{\Theta})^{-1} (\beta_E^* - \bar{R}\beta_E - \bar{s}) \right. \\ \left. - \frac{1}{2} (\gamma^* - \bar{A}\beta_E^* - \bar{b})^\top (\bar{\Omega})^{-1} (\gamma^* - \bar{A}\beta_E^* - \bar{b}) - \operatorname{Barr}(\gamma^*) \right), \quad (11)$$

is the basis of our surrogate posterior. Using convex analysis (Rockafellar, 2015), we state the surrogate selection-informed posterior and its gradient in the following theorem.

Theorem 5 *Fixing $\bar{\Sigma} = \bar{\Omega} + \bar{A}\bar{\Theta}(\bar{A})^\top$, $\bar{P} = \bar{A}\bar{R}$ and $\bar{q} = \bar{A}\bar{s} + \bar{b}$, let*

$$\gamma^* = \operatorname{argmin}_{\gamma \in \mathbb{R}^{|\mathcal{G}_E|}} \frac{1}{2}(\gamma - \bar{P}\beta_E - \bar{q})^\top (\bar{\Sigma})^{-1}(\gamma - \bar{P}\beta_E - \bar{q}) + \operatorname{Barr}(\gamma). \quad (12)$$

Letting $\Gamma^ = \operatorname{diag}\left((\gamma_g^* I_{|g|-1})_{g \in \mathcal{G}_E}\right)$ and M_g be the set of $|g| - 1$ diagonal indices of Γ^* for group g , define*

$$J^* = \left(\sum_{i \in M_g} [(\Gamma^* + \bar{U}^\top (X_E^\top X_E)^{-1} \Lambda \bar{U})^{-1}]_{ii} \right)_{g \in \mathcal{G}_E} \in \mathbb{R}^{|\mathcal{G}_E|}.$$

Using the approximation in (11), the logarithm of our surrogate selection-informed posterior is equal to

$$\log \pi_E(\beta_E) + \log p(\bar{R}\beta_E + \bar{s}, \bar{\Theta}; \hat{\beta}_E) - \log p(\bar{P}\beta_E + \bar{q}, \bar{\Sigma}; \gamma^*) + \operatorname{Barr}(\gamma^*) - \log J_\phi(\gamma^*; \mathcal{U}),$$

and has gradient equal to

$$\nabla \log \pi_E(\beta_E) + \bar{R}^\top (\bar{\Theta})^{-1} (\hat{\beta}_E - \bar{R}\beta_E - \bar{s}) + \bar{P}^\top (\bar{\Sigma})^{-1} (\bar{P}\beta_E + \bar{q} - \gamma^* - (\bar{\Sigma}^{-1} + \nabla^2 \operatorname{Barr}(\gamma^*))^{-1} J^*).$$

Recall, $|\mathcal{G}_E|$ is the number of selected groups after solving (3). As is evident from Theorem 5, gradient-based sampling from the surrogate posterior requires us to solve a $\mathbb{R}^{|\mathcal{G}_E|}$ -dimensional optimization problem in every update, without carrying out integrations for the theoretical adjustment. Algorithm 1 outlines a prototype implementation of our methods to generate estimates from the surrogate posterior.

We revisit our simple running example in Section 3 to instantiate Algorithm 1. In the selection stage, we solve (3) with $\omega \sim \mathcal{N}_2(0, \tau^2 I_2)$ where τ^2 is the randomization variance. Before noting the updates from the surrogate posterior, we assess in Figure 3 the relative accuracy of the Laplace approximation in (11) with respect to the exact normalizing constant. Because we have exactly one selected group of covariates, the exact normalizer is a one-dimensional integral that can be computed numerically. In particular, we observe how the relative accuracy of the approximation varies with sample size n and randomization variance τ^2 . For any fixed sample size, the accuracy of approximation for the integral with the mode decreases as the randomization variance increases, or equivalently as the concentration of probability mass in the integrand has a greater spread. As expected, the relative accuracy converges to 0 with growing sample size for all values of randomization variance.

Now, we exemplify Algorithm 1 for the simple example. Applying Theorems 2 and 5, (**Laplace**) in this instance solves the one-dimensional optimization

$$\gamma^{*(k)} = \operatorname{argmin}_{\gamma \in \mathbb{R}} \left\{ \frac{1}{2}(\gamma - \bar{P}\beta_E^{(k)} - \bar{q})^\top (\bar{\Sigma})^{-1}(\gamma - \bar{P}\beta_E^{(k)} - \bar{q}) + \operatorname{Barr}(\gamma) \right\}$$

Algorithm 1 A Prototype Implementation of our Selection-informed Bayesian Method

SELECT: $(y, X, \Omega, \mathcal{G}, \{\lambda_g\}_{g \in \mathcal{G}}) \xrightarrow{\text{Optimize (3)}} \hat{E} = E, \hat{\mathcal{U}} = \mathcal{U}, \hat{\mathcal{Z}} = \mathcal{Z}$

STEPS: Set up parameters for (Laplace)

(Orthonormal completion) Calculate $\bar{U}$ (see Theorem 2)

(Parameters) Calculate $\bar{R}, \bar{s}, \bar{\Theta}, \bar{P}, \bar{q}, \bar{\Sigma}$ (see Theorem 5)

STEPS: Implementation for a generic gradient-based sampler

(Initialize) Sample: $\beta_E^{(1)} = \hat{\beta}_E$, Step Size: η , Proposal Scale: χ , Number of Samples: K

for $k = 1, 2, \dots, K - 1$ **do**

(**Laplace**) Solve $\gamma^{*(k)} = \underset{\gamma \in \mathbb{R}^{|\mathcal{G}_E|}}{\text{argmin}} \left\{ \frac{1}{2}(\gamma - \bar{P}\beta_E^{(k)} - \bar{q})^\top (\bar{\Sigma})^{-1}(\gamma - \bar{P}\beta_E^{(k)} - \bar{q}) + \text{Barr}(\gamma) \right\}$

(Jacobian) Calculate $J^{*(k)}$ (see Theorem 5)

(Gradient) Calculate $\nabla \log \pi_E(\beta_E^{(k)} \mid \hat{\beta}_E, N_E)$ (see Theorem 5)

(Update) $\beta_E^{(k+1)} \leftarrow \beta_E^{(k)} + \eta \chi \nabla \log \pi_E(\beta_E^{(k)} \mid \hat{\beta}_E, N_E) + \sqrt{2\eta} \epsilon^{(k)}, \epsilon^{(k)} \sim \mathcal{N}(0, \chi)$.

end for

where $\bar{P} = U^\top(2 \cdot I_2 - UU^\top)^{-1}$, $\bar{q} = -1$, $\bar{\Sigma} = 1 + U^\top(2 \cdot I_2 - UU^\top)^{-1}U$ and U is the unit vector in $\mathbb{R}^2$ from writing the Group LASSO solution as (8). The Jacobian function is given by $J^{*(k)} = (1 + \gamma^{*(k)})^{-1}$. We generate our prototype update for $\beta_E^{(k+1)}$ using the expression for the gradient of the surrogate in Theorem 5.

At last, we briefly comment on the case when σ , the error variance in the data, is treated as an unknown parameter. Using a joint prior on (β_E, σ) in conjunction with our selection-informed likelihood gives us Bayesian inference in this situation. As prescribed in Panigrahi et al. (2021), one could run a Gibbs sampler that alternates between drawing (i) an update for β_E given σ and (ii) an update for σ given β_E . Note that the updates for β_E , using a gradient-based sampler, will assume the expression provided in Theorem 5.

5. Generalization to Other Grouped Sparsities

We now generalize our selection-informed methods to learning algorithms which target other forms of grouped sparsities and covariate structures.

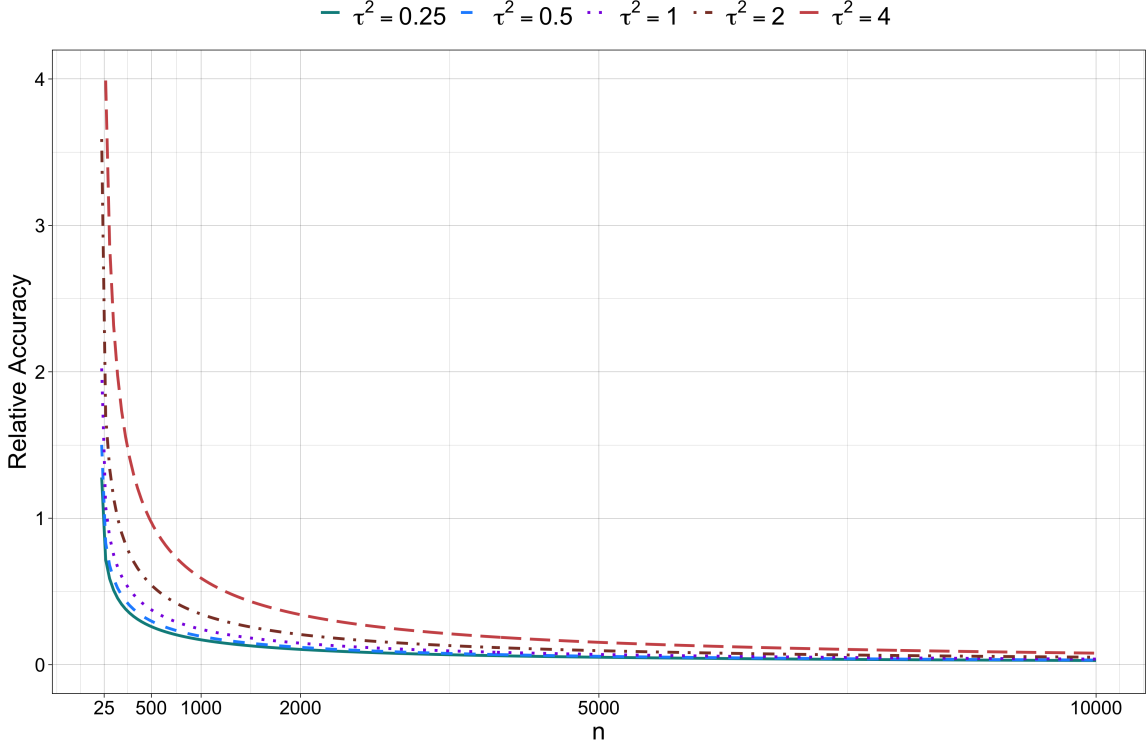


Figure 3: Plot for relative accuracy of the log-Laplace approximation with respect to the exact log-normalizing constant.

5.1 Overlapping Group LASSO

One approach for the overlapping Group LASSO recovers superposed groups by augmenting the covariate space with duplicated predictors (Jacob et al., 2009). We proceed by implementing the solver (3) with a randomization variable $\omega^* \sim \mathcal{N}_{p^*}(0, \Omega)$ such that p^* is the number of covariates after duplication. To formalize the setting, we let $X^* \in \mathbb{R}^{n \times p^*}$ denote the augmented matrix of covariates constructed from $X \in \mathbb{R}^{n \times p}$, given that these overlapping groups are determined before selection. Each set of selected covariates E^* in the augmented space maps to a set of selected variables E in the original space by reversing the duplication. The stationary mapping for the overlapping Group LASSO with augmented matrix X^* and randomization ω^* , which we call ϕ^* , is given by

$$\omega^* = X^{*\top} X^* \begin{pmatrix} (\gamma_g^* u_g^*)_{g \in \mathcal{G}_{E^*}} \\ 0 \end{pmatrix} - X^{*\top} X_E \hat{\beta}_E - N_E + \begin{pmatrix} (\lambda_g u_g^*)_{g \in \mathcal{G}_{E^*}} \\ (\lambda_h z_h^*)_{h \in -\mathcal{G}_{E^*}} \end{pmatrix} \quad (13)$$

where $\hat{\beta}_E$ and $N_E = (X^*)^\top (y - X_E \hat{\beta}_E)$ are the refitted and the ancillary statistics in our selected model (cf. (4)). Following our notation,

$$\hat{\gamma}^* = (\gamma_g^* : g \in \mathcal{G}_{E^*})^\top, \quad \hat{\mathcal{U}}^* = \{u_g^* : g \in \mathcal{G}_{E^*}\},$$

represent a decomposition of the overlapping Group LASSO solution as done in (8). Let $U^* = \text{diag} \left((u_g^*)_{g \in \mathcal{G}_{E^*}} \right)$. Finally,

$$\widehat{\mathcal{Z}}^* = \{z_g^* : g \in -\mathcal{G}_{E^*}\}$$

are the subgradient variables from the Group LASSO penalty for the non-selected groups in the augmented predictor space.

The conditioning event we study for an analytically feasible selection-informed posterior is given by

$$\mathcal{A}_{E^*} = \{\widehat{E}^* = E^*, \widehat{U}^* = U^*, \widehat{\mathcal{Z}}^* = \mathcal{Z}^*\}$$

where E^* , U^* and $\mathcal{Z}^*$ are the corresponding observed instances. We then recover an expression for the adjustment factor along the lines of Theorem 2 using the matrices

$$A = -(X^*)^\top X_E, \quad B = (X^*)^\top X_{E^*}^* U^*, \quad c = -N_E + \left((\lambda_g u_g^*)_{g \in \mathcal{G}_E}^\top \quad (\lambda_g z_g^*)_{g \in -\mathcal{G}_E}^\top \right)^\top \quad (14)$$

from the mapping in (13).

Proposition 6 *Define the following matrices:*

$$\bar{U}^* = \text{diag} \left((\bar{U}_g^*)_{g \in \mathcal{G}_{E^*}} \right), \Gamma^* = \text{diag} \left((\hat{\gamma}_g^* I_{|g|})_{g \in \mathcal{G}_{E^*}} \right), \Lambda = \text{diag} \left((\lambda_g I_{|g|})_{g \in \mathcal{G}_{E^*}} \right),$$

where A , B , c are specified in (14). Then, $\mathbb{P}(\mathcal{A}_{E^*} \mid \beta_E)$ is proportional to

$$\int \int p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot \exp \left\{ -\frac{1}{2} (A \hat{\beta}_E + B \hat{\gamma}^* + c)^\top \Omega^{-1} (A \hat{\beta}_E + B \hat{\gamma}^* + c) \right\} \\ \times J_{\phi^*}(\hat{\gamma}^*; U^*) \cdot \mathbf{1}(\hat{\gamma}^* > 0) d\hat{\beta}_E d\hat{\gamma}^*$$

where

$$J_{\phi^*}(\hat{\gamma}^*; U^*) = \det \left(((X_{E^*}^*)^\top X_{E^*}^* \Gamma^* + \Lambda) \bar{U}^* \quad (X_{E^*}^*)^\top X_{E^*}^* U^* \right). \quad (15)$$

Notice that since X^* contains overlapping groups, the matrix $(X_{E^*}^*)^\top X_{E^*}^*$ may not be invertible. Thus, in comparison to Theorem 2, (15) provides a different expression for the Jacobian function involving the sizes of the selected groups of variables. In solving (3), one may introduce a ridge penalty $\epsilon \|\beta\|_2^2/2$, where ϵ is a small positive number. This will in turn lead to (14) with

$$B = \begin{pmatrix} (X_E^*)^\top X_{E^*}^* + \epsilon \cdot I_{E^*} \\ (X_{-E^*}^*)^\top X_{E^*}^* \end{pmatrix} U^*.$$

This results in the following Jacobian

$$J_{\phi^*}(\hat{\gamma}^*; U^*) = \det \left(\Gamma^* + (\bar{U}^*)^\top ((X_{E^*}^*)^\top X_{E^*}^* + \epsilon \cdot I_{E^*})^{-1} \Lambda \bar{U}^* \right)$$

where the ridge parameter can be used to counter the collinearity in the augmented predictor matrix.

5.2 Standardized Group LASSO

An alternate treatment to the Group LASSO objective is popularly applied in problems with correlated covariates when a within-group orthonormality is desired. The learning algorithm proposed by Simon and Tibshirani (2012) addresses the selection of groups in the presence of such correlations via a modification to the Group LASSO penalty. Equivalently, the canonical objective (3) is reparameterized in the standardized formulation; for each submatrix X_g containing the predictors in group $g \in \mathcal{G}$, the quadratic loss function is now given by

$$\|y - \sum_g X_g \beta_g\|_2^2 = \|y - \sum_g W_g \theta_g\|_2^2, \quad (16)$$

$X_g = W_g R_g$ for W_g , an orthonormal matrix, and R_g , an invertible matrix, and $\theta_g = R_g \beta_g$, the reparameterized vector. The standardized Group LASSO optimizes the Group LASSO objective in terms of θ rather than β :

$$\hat{\theta}^{(\mathcal{G})} = \underset{\theta}{\operatorname{argmin}} \left\{ \frac{1}{2} \|y - \sum_{g \in \mathcal{G}} W_g \theta_g\|_2^2 + \sum_{g \in \mathcal{G}} \lambda_g \|\theta_g\|_2 - \omega^\top \theta \right\}. \quad (17)$$

Lastly, the original parameters of interest are estimated by $R_g^{-1} \hat{\theta}_g^{(\mathcal{G})}$. The selected set of groups E comprises groups with non-zero coordinates in $\hat{\theta}^{(\mathcal{G})}$ after solving (17).

Letting $W \in \mathbb{R}^{n \times p}$ be the column-wise concatenation of the standardized groups of covariates W_g , the stationary mapping for the standardized Group LASSO is given by

$$\omega = W^\top W \begin{pmatrix} (\gamma_g u_g)_{g \in \mathcal{G}_E} \\ 0 \end{pmatrix} - W^\top X_E \hat{\beta}_E - N_E + \begin{pmatrix} (\lambda_g u_g)_{g \in \mathcal{G}_E} \\ (\lambda_h z_h)_{h \in -\mathcal{G}_E} \end{pmatrix}; \quad (18)$$

$\hat{\beta}_E$ is our usual refitted statistic and $N_E = W^\top (Y - X_E \hat{\beta}_E)$. Consistent with our approach, the (non-zero) standardized Group LASSO solution $\hat{\theta}^{(\mathcal{G})}$ can be decomposed in terms of

$$\gamma = (\gamma_g : g \in \mathcal{G}_E), \quad \mathcal{U} = \{u_g : g \in \mathcal{G}_E\}.$$

Recall that

$$\mathcal{Z} = \{z_g : g \in -\mathcal{G}_E\}$$

are the subgradient variables for the non-selected groups. Setting

$$\mathcal{A}_E = \{\hat{E} = E, \hat{\mathcal{U}} = \mathcal{U}, \hat{\mathcal{Z}} = \mathcal{Z}\},$$

$$A = -W^\top X_E, \quad B = W^\top W_E U, \quad c = -N_E + \left((\lambda_g u_g)_{g \in \mathcal{G}_E}^\top \quad (\lambda_g z_g)_{g \in -\mathcal{G}_E}^\top \right)^\top, \quad (19)$$

we present the adjustment factor in line with Theorem 2 after a change of variables from inverting the stationary mapping of (17).

Proposition 7 *Consider the following matrices:*

$$\bar{U} = \operatorname{diag} \left((\bar{U}_g)_{g \in \mathcal{G}_E} \right), \Gamma = \operatorname{diag} \left((\hat{\gamma}_g I_{|g|-1})_{g \in \mathcal{G}_E} \right), \Lambda = \operatorname{diag} \left((\lambda_g I_{|g|})_{g \in \mathcal{G}_E} \right).$$

We then have

$$\mathbb{P}(\mathcal{A}_E \mid \beta_E) \propto \int \int p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot \exp \left\{ -\frac{1}{2} (A\hat{\beta}_E + B\hat{\gamma} + c)^\top \Omega^{-1} (A\hat{\beta}_E + B\hat{\gamma} + c) \right\} \\ \times J_\phi(\hat{\gamma}; \mathcal{U}) \cdot \mathbf{1}(\hat{\gamma} > 0) d\hat{\beta}_E d\hat{\gamma}$$

where

$$J_\phi(\hat{\gamma}; \mathcal{U}) = \det(\Gamma + \bar{U}^\top (W_E^\top W_E)^{-1} \Lambda \bar{U}).$$

5.3 Sparse Group LASSO

The sparse Group LASSO (Simon et al., 2013) produces solutions that are sparse at both the group level and the individual level within selected groups by deploying the Group LASSO penalty along with the usual ℓ_1 penalty. A randomized formulation of the sparse Group LASSO is given by

$$\operatorname{argmin}_{\beta} \left\{ \frac{1}{2} \|y - X\beta\|_2^2 + \sum_g \lambda_g \|\beta_g\|_2 + \lambda_0 \|\beta\|_1 - \omega^\top \beta \right\};$$

the sum over g forms a non-overlapping partition of the predictors. Notice that this criterion may be viewed as a special case of the overlapping Group LASSO where each predictor i appears in a group $g \in \{1, \dots, G\}$, as well as in its own individual group.

Setting up notations, recall that $\mathcal{G}_E$ denotes the set of selected groups and $-\mathcal{G}_E$ its complement; let T_g denote the selected predictors in group g , and $-T_g$ the corresponding complement. Finally, we define

$$\check{E} = \bigcup_{g \in \mathcal{G}_E} T_g,$$

the set of selected predictors in the selected groups which parameterize our model. We write the stationary mapping for the sparse Group LASSO below:

$$\omega = X^\top X \begin{pmatrix} (\gamma_g u_g)_{g \in \mathcal{G}_E} \\ 0 \end{pmatrix} - X^\top X_{\check{E}} \hat{\beta}_{\check{E}} - N_{\check{E}} + \begin{pmatrix} (\lambda_g u_g)_{g \in \mathcal{G}_E} \\ (\lambda_h z_h)_{h \in -\mathcal{G}_E} \end{pmatrix} + \lambda_0 \left(\begin{pmatrix} (s_j)_{j \in T_g} \\ (s_j)_{j \in -T_g} \end{pmatrix}_{g \in \mathcal{G}} \right) \quad (20)$$

where $\hat{\beta}_{\check{E}}$ and $N_{\check{E}}$ are defined as per projections according to the selected model; we denote this mapping by $\check{\phi}$. Recall that γ , $\mathcal{U}$, and $\mathcal{Z}$ are consistent in their definition in terms of the groups we select after solving (3). In addition, the subgradient variables from the ℓ_1 penalty are represented by s_j , with $s_j = \operatorname{sign}(\hat{\beta}_j^{(\mathcal{G})})$ for $j \in \check{E}$, and $|s_j| < 1$ for $j \in -\check{E}$. We collect these scalar variables into the two sets $\mathcal{S}_{\mathcal{G}_E}$ and $\mathcal{S}_{-\mathcal{G}_E}$, depending on whether the predictor is in a selected group. For non-selected predictors in selected groups, the corresponding entry of u_g is zero. That is, we can interpret u_g as a unit vector with the same dimension as the selected part of group g , denoting it by $\check{u}_g = (u_{g,j} : j \in T_g)^\top \in \mathbb{R}^{|T_g|}$. Set $\check{U} = \operatorname{diag}((\check{u}_g)_{g \in \mathcal{G}_E})$. Define the selection event

$$\mathcal{A}_{\check{E}} = \{\hat{E} = \check{E}, \hat{\mathcal{U}} = \mathcal{U}, \hat{\mathcal{Z}} = \mathcal{Z}, \hat{\mathcal{S}}_{\mathcal{G}_E} = \mathcal{S}_{\mathcal{G}_E}, \hat{\mathcal{S}}_{-\mathcal{G}_E} = \mathcal{S}_{-\mathcal{G}_E}\},$$

$$A = -X^\top X_{\check{E}}, \quad B = X^\top X_{\check{E}} \check{U}, \quad c = -N_E + \left((\lambda_g u_g)_{g \in \mathcal{G}_E}^\top \quad (\lambda_g z_g)_{g \in -\mathcal{G}_E}^\top \right)^\top + \lambda_0 \left(((s_j)_{j \in g})_{g \in \mathcal{G}} \right)^\top, \quad (21)$$

using the stationary mapping for the learning algorithm under scrutiny. We then recover the following theoretical expression for the adjustment factor after selection via the sparse Group LASSO.

Proposition 8 *For each selected group $g \in \mathcal{G}_E$, construct $\check{U}_g \in \mathbb{R}^{|T_g| \times (|T_g|-1)}$ as the orthonormal basis completion of $\check{u}_g$. Define the following matrices:*

$$\check{U} = \text{diag} \left(\left(\check{U}_g \right)_{g \in \mathcal{G}_E} \right), \Gamma = \text{diag} \left((\gamma_g I_{|T_g|-1})_{g \in \mathcal{G}_E} \right), \Lambda = \text{diag} \left((\lambda_g I_{|T_g|})_{g \in \mathcal{G}_E} \right),$$

based upon (21). Then, we have

$$\begin{aligned} \mathbb{P}(\mathcal{A}_E \mid \beta_{\check{E}}) &\propto \int \int \mathbb{P}(\beta_{\check{E}}, \Sigma_{\check{E}}; \hat{\beta}_{\check{E}}) \cdot \exp \left\{ -\frac{1}{2} (A \hat{\beta}_{\check{E}} + B \hat{\gamma} + c)^\top \Omega^{-1} (A \hat{\beta}_{\check{E}} + B \hat{\gamma} + c) \right\} \\ &\quad \times J_{\check{\phi}_{\hat{\beta}_{\check{E}}}}(\hat{\gamma}; \mathcal{U}) \cdot \mathbf{1}(\hat{\gamma} > 0) d\hat{\beta}_{\check{E}} d\hat{\gamma} \end{aligned}$$

where

$$J_{\check{\phi}}(\hat{\gamma}; \mathcal{U}) = \det(\Gamma + \check{U}^\top (X_E^\top X_{\check{E}})^{-1} \Lambda \check{U}).$$

6. Large Sample Theory

We establish statistical credibility of inference done using our surrogate selection-informed posterior under a frequentist regime with fixed p and growing n . We fix $\beta_{n,E}$ to be the sequence of parameters governing our generating model such that $\sqrt{n}\beta_{n,E} = b_n \bar{\beta}_E$, where $n^{-1/2}b_n = O(1)$, $b_n \rightarrow \infty$ as $n \rightarrow \infty$. Introducing the dependence on the sample size, let $\ell_{n,E}(\beta_{n,E}; \hat{\beta}_{n,E} \mid N_{n,E})$ represent the surrogate (log) selection-informed likelihood given in Theorem 5. After ignoring constants and the prior, the surrogate (log) likelihood takes the value

$$\begin{aligned} \log \mathbb{P}(\bar{R}\sqrt{n}\beta_{n,E} + \bar{s}, \bar{\Theta}; \sqrt{n}\hat{\beta}_{n,E}) &+ \frac{n}{2}(\gamma_n^* - \bar{P}\beta_{n,E} - n^{-1/2}\bar{q})^\top (\bar{\Sigma})^{-1}(\gamma_n^* - \bar{P}\beta_{n,E} - n^{-1/2}\bar{q}) \\ &+ \text{Barr}(\sqrt{n}\gamma_n^*) - \log J_{\phi}(\sqrt{n}\gamma_n^*; \mathcal{U}), \end{aligned}$$

which uses the optimizer

$$\gamma_n^* = \underset{\gamma}{\text{argmin}} \quad \frac{n}{2}(\gamma - \bar{P}\beta_{n,E} - n^{-1/2}\bar{q})^\top (\bar{\Sigma})^{-1}(\gamma - \bar{P}\beta_{n,E} - n^{-1/2}\bar{q}) + \text{Barr}(\sqrt{n}\gamma).$$

Recall that appending our surrogate selection-informed likelihood to a selection-informed prior $\pi_E(\cdot)$ gives us our selection-informed posterior. We denote the measure of a set $\mathcal{K}$ with respect to this posterior distribution as follows:

$$\Pi_{n,E}(\mathcal{K} \mid \hat{\beta}_{n,E}; N_{n,E}) = \frac{\int_{\mathcal{K}} \pi_E(z_n) \cdot \exp(\ell_{n,E}(z_n; \hat{\beta}_{n,E} \mid N_{n,E})) \, dz_n}{\int \pi_E(z_n) \cdot \exp(\ell_{n,E}(z_n; \hat{\beta}_{n,E} \mid N_{n,E})) \, dz_n}.$$

We let

$$\mathcal{B}(\beta_{n,E}, \delta) = \{z : \|z - \beta_{n,E}\|_2^2 < \delta\}$$

denote a ball of radius δ around our parameter of interest, $\beta_{n,E}$, and we let $\mathcal{B}^c(\beta_{n,E}, \delta)$ denote its complement. Lastly, we use $\mathbb{P}_{n,E}(\cdot)$ to represent the selection-informed probability after conditioning upon our selection event and the ancillary statistic under the generating parameter, $\beta_{n,E}$.

Our main theoretical result, Theorem 12, proves that, subject to mild conditions on the selection-informed prior distribution, our surrogate version of the selection-informed posterior concentrates around the true parameter as the sample size grows infinitely large; this gives us the rate of contraction. We begin with two supporting propositions: (i) Proposition 9 proves the convergence of the approximate normalizing constant based on (11) to the exact counterpart when the support constraints for the sizes of the selected groups are restricted to a compact subset; (ii) Proposition 11 bounds the curvature of the surrogate (log) selection-informed likelihood around its maximizer. Proofs of these main results are in Appendix B. To support the claim in Proposition 11, Lemma 13 and Lemma 14 provide supplementary theory to control the asymptotic orders of the gradient and Hessian of the (log) Jacobian in our surrogate selection-informed posterior. We include both these results in Appendix C.

Proposition 9 *Suppose that*

$$\lim_{n \rightarrow \infty} (b_n)^{-2} \{ \log \mathbb{P}(\sqrt{n}\gamma_n > 0) - \log \mathbb{P}(\sqrt{n}\gamma_n > \bar{q}) \} = 0. \quad (22)$$

Let $\bar{\gamma}_n^$ be given as*

$$\operatorname{argmin}_{\bar{\gamma} < \bar{Q} \cdot 1_{|E|}} \frac{b_n^2}{2} (\bar{\gamma} - \bar{P}\bar{\beta}_E - (b_n)^{-1}\bar{q})^\top \bar{\Sigma}^{-1} (\bar{\gamma} - \bar{P}\bar{\beta}_E - (b_n)^{-1}\bar{q}) + \operatorname{Barr}(b_n \bar{\gamma}),$$

where $\bar{Q} > 0$. Then we have

$$\begin{aligned} \lim_{n \rightarrow \infty} (b_n)^{-2} \log \mathbb{P}(0 < \sqrt{n}\gamma_n < b_n \bar{Q} \cdot 1_{|E|} + \bar{q}) + (b_n)^{-2} \operatorname{Barr}(b_n \bar{\gamma}_n^*) - (b_n)^{-2} \log J_\phi(b_n \bar{\gamma}_n^*; \mathcal{U}) \\ + \frac{1}{2} (\bar{\gamma}_n^* - \bar{P}\bar{\beta}_E - (b_n)^{-1}\bar{q})^\top \bar{\Sigma}^{-1} (\bar{\gamma}_n^* - \bar{P}\bar{\beta}_E - (b_n)^{-1}\bar{q}) = 0. \end{aligned}$$

Remark 10 *For a fixed selection-informed prior $\pi_E(\cdot)$, we may choose $\bar{Q}$ to be a positive constant in order to consider a sufficiently large compact subset of our selection region that would work for all $\bar{\beta}_E$ in a bounded set of probability close to 1 under our prior. Proposition 9 now implies that our surrogate (log) selection-informed likelihood converges to its exact counterpart, obtained by plugging in the exact probability of selection under the parameter sequence $\beta_{n,E}$, as the sample size grows to ∞ .*

Proposition 11 *Fix $\mathcal{C} \in \mathbb{R}^{|E|}$, a compact set. Define $\hat{\beta}_{n,E}^{\max}$ to be the maximizer of the selection-informed likelihood sequence, $\ell_{n,E}(\cdot; \hat{\beta}_{n,E} \mid N_{n,E})$. Then there exist positive con-*

stants $C_0 \leq C_1$ and $N \in \mathbb{N}$ such that for any $0 < \epsilon_0 < C_0$,

$$\begin{aligned} -\frac{n}{2}(C_1 + \epsilon_0)\|z_n - \hat{\beta}_{n,E}^{max}\|_2^2 &\leq \ell_{n,E}(z_n; \hat{\beta}_{n,E} \mid N_{n,E}) - \ell_{n,E}(\hat{\beta}_{n,E}^{max}; \hat{\beta}_{n,E} \mid N_{n,E}) \\ &\leq -\frac{n}{2}(C_0 - \epsilon_0)\|z_n - \hat{\beta}_{n,E}^{max}\|_2^2 \end{aligned}$$

for all $n \geq N$ and $z_n \in \mathcal{C}$.

We are now ready to state and prove our main theoretical result on the concentration properties of our selection-informed posterior.

Theorem 12 *Suppose a selection-informed prior $\pi_E(\cdot)$ with compact support $\mathcal{C}$ assigns non-zero probability to $\mathcal{B}(\beta_{n,E}, \delta') \subset \mathcal{C}$ for any $\delta' > 0$ such that the inclusion is satisfied. Further, assume for the associated prior measure $\Pi_E(\cdot)$ that*

$$\lim_{n \rightarrow \infty} \exp(-b_n^2 \delta^2 K/2) / \Pi_E(\mathcal{B}(\beta_{n,E}, \kappa \delta_n)) = 0$$

for any $K > 0$, $\kappa \in (0, 1)$, and $\delta > 0$, where δ_n is defined by $\sqrt{n}\delta_n = b_n\delta$. Then, the following convergence must hold for any $\epsilon > 0$:

$$\mathbb{P}_{n,E} \left(\Pi_{n,E} \left(\mathcal{B}^c(\beta_{n,E}, \delta_n) \mid \hat{\beta}_{n,E}; N_{n,E} \right) \leq \epsilon \right) \rightarrow 1 \text{ as } n \rightarrow \infty.$$

We visualize in Figure 2 the concentration theory presented in Theorem 12 for our simple example with a single group of two covariates.

7. Empirical Investigations

7.1 Experimental Design

In all of our experiments with synthetic data, we construct the design X by drawing $n = 500$ rows independently according to $\mathcal{N}_p(0, \Sigma)$; Σ follows an autoregressive structure with the (i, j) -th entry of the covariance matrix $\Sigma_{(i,j)} = 0.2^{|i-j|}$. We fix the support and vary the values of β according to a variety of schemes described below; in each case, we have a “Low,” “Medium,” and “High” signal-to-noise ratio (SNR) regime as detailed in Appendix D. Finally, we draw $Y \sim \mathcal{N}_n(X\beta, \sigma^2 I_n)$, a Gaussian group-sparse linear model; we fix $\sigma = 3$ in our experiments. We consider the following settings for our grouped covariates:

- **Balanced:** In the balanced setting, we partition (i) $p = 100$ covariates into 25 disjoint groups each of cardinality four when we solve the canonical Group LASSO and the standardized Group LASSO to learn a group-sparse model; (ii) $p = 103$ covariates into 34 groups of four predictors each, but the last feature of the first group is also the first feature of the second group, and so on when we solve the overlapping Group LASSO. In the latter case, the first and last groups each have three features in no other groups, and all other groups have two features in no other groups. We randomly select three of these candidate groups to be active, and let each coefficient assume a random sign with the same magnitudes that in turn depend on the SNR regime.

- **Heterogeneous:** In the heterogeneous setting, we allow the disjoint groups of covariates to differ in their sizes. Further, our groups of covariates now display heterogeneity in the signal amplitudes both within and between the active groups. We have 3 groups with three predictors each, 4 groups with four predictors each, 5 groups with five predictors each, and 5 groups with ten predictors each. We set one of each of the three-, four-, and five-predictor groups to be active with linearly increasing signal magnitudes and each active coefficient is assigned a random sign.

For each realization of the data and each setting under study, we apply three methods: (i) “Selection-informed”, the selection-informed implementation summarized in Algorithm 1; drawing each sample remarkably solves only a $|\mathcal{G}_E|$ -dimensional optimization problem as can be appreciated by reviewing Step **(Laplace)** in Algorithm 1; (ii) “Naive”, the standard inferential tool that first fits the usual Group LASSO (3) with no randomization to identify the active set E , and then fits $\hat{\beta}_E$ using ordinary least squares restricted to X_E from which we obtain credible intervals ignoring the effects of selection; (iii) “Split”, the sample splitting method follows the same procedure as “Naive” except that this method partitions the data at a prespecified ratio r , that is, “Split” applies the usual Group LASSO to $[rn]$ randomly chosen subsamples without replacement to obtain E and then uses the remaining (holdout) samples to fit a linear model restricted to E for interval estimation. The nominal level for the interval estimates is set at 90%. Following Algorithm 1, a (gradient-based) Langevin sampler takes a noisy step along the gradient of our posterior (5) at each draw:

$$\beta_E^{(K+1)} = \beta_E^{(K)} + \eta \chi \nabla \log \pi_E(\beta_E^{(k)} \mid \hat{\beta}_E, N_E) + \sqrt{2\eta} \epsilon^{(K)},$$

where $\eta > 0$ is a predetermined step size, $\pi_E(\beta_E \mid \hat{\beta}_E, N_E)$ denotes our surrogate selection-informed posterior at the parameter vector β_E , $\epsilon^{(K)} \sim N(0, \chi)$, our Gaussian proposal (Shang et al., 2015), and we plug in the expression for the gradient of the surrogate posterior from Theorem 5. In practice, we set $\eta = 1$ and determine χ from the inverse of the Hessian of our (negative-log) posterior. It bears emphasis that this sampler serves as a representative execution of our methods; more generally, other sampling schemes to deliver inference based upon our selection-informed posterior (and its gradient) are clearly possible. We construct credible intervals from the appropriate quantiles in the posterior sample of the parameter for “Selection-informed” under a diffuse Gaussian (selection-informed) prior, i.e., $\pi_E = \mathcal{N}_{|E|}(0, r_0 \sigma^2)$ with $r_0 = 100$. Our credible intervals for “Split” and “Naive” are reported for the same prior.

In our experimental findings, we draw attention to comparisons between “Split” based on an allocation of r fraction of the samples for selecting a group-sparse linear model and “Selection-informed” based on the solution of the randomized Group LASSO (3) with a p -dimensional isotropic Gaussian randomization variable independent of our response; that is, $\Omega = \tau^2 \cdot I_p$. To (approximately) match the amount of information used during selection by “Split” for an honest assessment of inference for the randomized methods, we fix the ratio of randomization variation to the noise level in our response as follows:

$$\frac{\tau^2}{\sigma^2} = \frac{(1-r)}{r}, \quad (23)$$

where r is the proportion of data reserved by “Split” for solving the Group LASSO. The value for randomization variation we set for comparisons is motivated from an asymptotic equivalence between data splitting and a Gaussian randomization scheme proved in Panigrahi et al. (2021, Proposition 4.1), which notes that a regularized regression objective using $[rn]$ subsamples, under an i.i.d. generative process for the response and covariates, can be formulated as (3) with the randomization covariance $\Omega = \frac{(1-r)}{r} \cdot \mathbb{E}[\frac{1}{n} X^\top X]$. In this sense, our choice of τ^2 allows us to mimic data carving during post-selective inference for the group-sparse parameters. Clearly, there will be a tradeoff between the information used for selection and inference when the same data is used for learning a group-sparse model and inferring for these selection-informed parameters. We remark that “Naive”, deploying all the samples for selection, does not require an analogous tradeoff between these two intertwined goals.

7.2 Inferential Findings for Different Grouped Sparsities

Based on the design of experiment in the preceding section, we undertake 100 rounds of numerical simulations for each level of randomization variation and a category of the three SNR regimes, “Low,” “Medium,” and “High.” Our experimental findings after solving the canonical Group LASSO (3) are summarized for Balanced and Heterogeneous groups in Figures 4, 6, and 8. These results are presented as box plots with the outliers suppressed. On the x-axis of these figures, we vary the level of randomization with decreasing levels from left to right. The randomization level is determined by the ratio of data r allocated for model selection and reserved for inference in the case of “Split,” and the level of variation for the corresponding Gaussian randomization scheme τ^2 is set according to (23) in the case of “Selection-informed”; this value is denoted by the label $x : y$ such that $(x + y)^{-1}x = r$ on the x-axis. Note that because “Naive” deploys no randomization, we found it instructive to assign it a label “0” for the randomization level on the x-axis. We then highlight how our method can be adapted for extensions to the overlapping and standardized Group LASSO through Figures 5, 7, and 9. To implement the overlapping Group LASSO, we take the approach of Jacob et al. (2009) that duplicates overlapping features to obtain an augmented design matrix $X^\star$ with no overlaps. Because columns are duplicated, the expanded design matrix is rank deficient; we therefore incorporate a (small) ridge term as per the prescription in Section 5. After selecting the active set $E^\star$, we map back to the set of selected variables E in the original space to define our group-sparse model and perform inference for β_E .

In terms of our findings, Figures 4 and 5 first depict an assessment of the model selection accuracy which we measure in terms of the F1 score:

$$\text{F1 score} = \frac{\text{True Positives}}{\text{True Positives} + \frac{1}{2}(\text{False Positives} + \text{False Negatives})}.$$

These plots corroborate the approximate correspondence in the amount of information used for model selection by the two randomized methods and subsequently in the quality of models selected by them. We note that the selection accuracy for the randomized methods increases with decreased levels of randomization and is bounded above by the “Naive” selection based on all the data. Figures 6 and 7 plot the distribution of coverages of

the credible intervals for all three methods grouped by the level of randomization. We emphasize that the selected model is likely misspecified (with a very high probability) in our settings which is evident from Figures 4 and 5. Regardless of model misspecification, we deliver inference for a well-defined target—the best linear representation of the response mean using X_E —as was described in Section 3. Consistent with expectations, “Naive” does not yield honest interval estimates with the shortfall in coverage understandably more severe for the lower SNR regimes. “Selection-informed” and “Split” discard information from model selection to counteract the bias in uncertainty estimation. Figures 8 and 9 illustrate our motivation to borrow residual information from selection in order to construct more efficient inferential procedures than the benchmark offered by sample splitting. The gains in efficiency for “Selection-informed” are noticeable from the clear separation in the distributions of the lengths of the interval estimates produced by the randomized approaches at a fixed level of randomization in diverse grouped settings, for example, we may compare the third quartile for the “Selection-informed” distribution with the first quartile for the corresponding “Split.” The “Naive” intervals in the figure for lengths highlight the price paid in terms of efficiency for constructing honest estimates of uncertainty post selection. We note a marginal loss in inferential accuracy for the interval estimates based on our method as the level of randomization increases. This is attributed to the Laplace-type approximation we apply to replace the probability of selection with the mode of the associated integrand. That is, the quality of approximation under a fixed sample size deteriorates as the concentration of probability mass is more spread out, for instance for randomization level “1:2,” which has four times the level of randomization variation compared to “2:1.” The above intuitive explanation is corroborated by Figure 3 for the simple, running example in the paper.

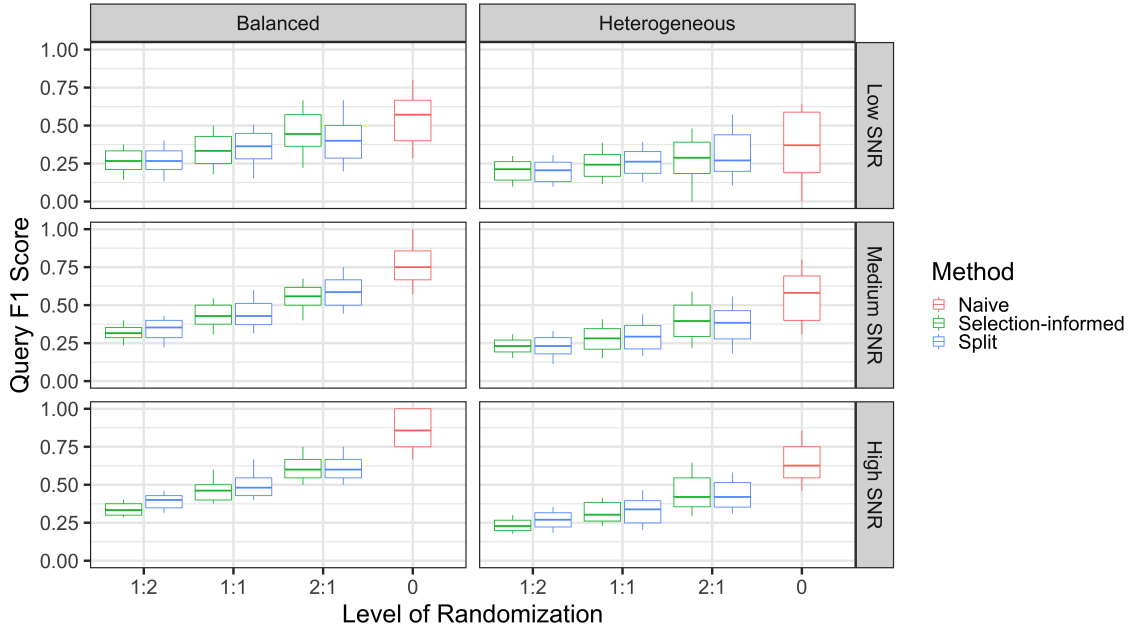


Figure 4: Box plots for model selection accuracy under the Group LASSO for the balanced and heterogeneous groups.

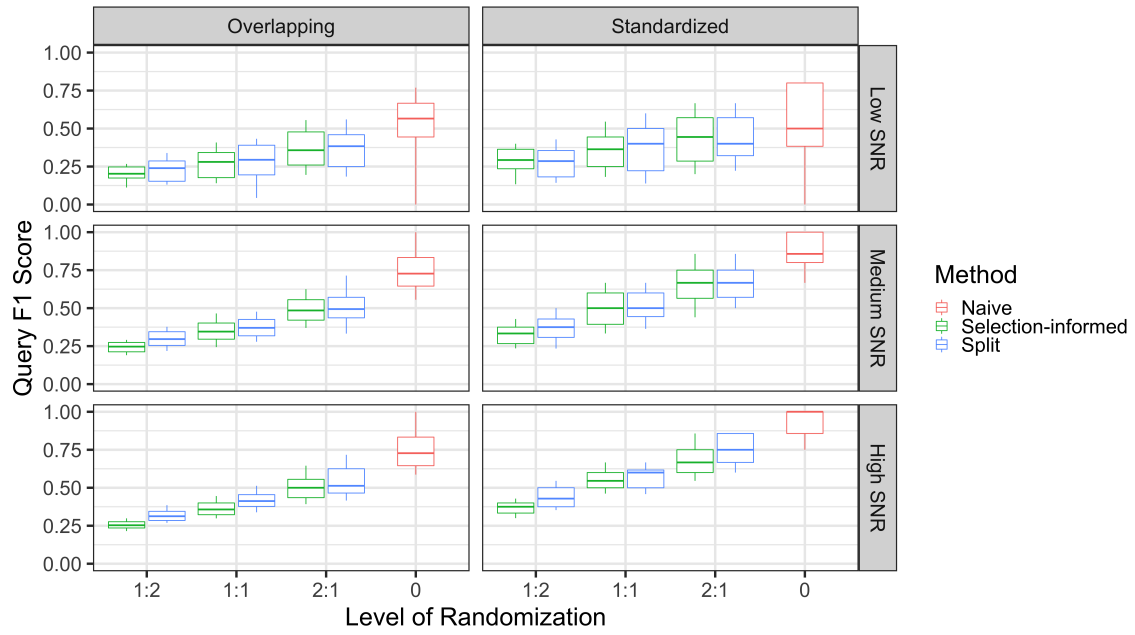


Figure 5: Box plots for model selection accuracy under the two extensions of the Group LASSO.

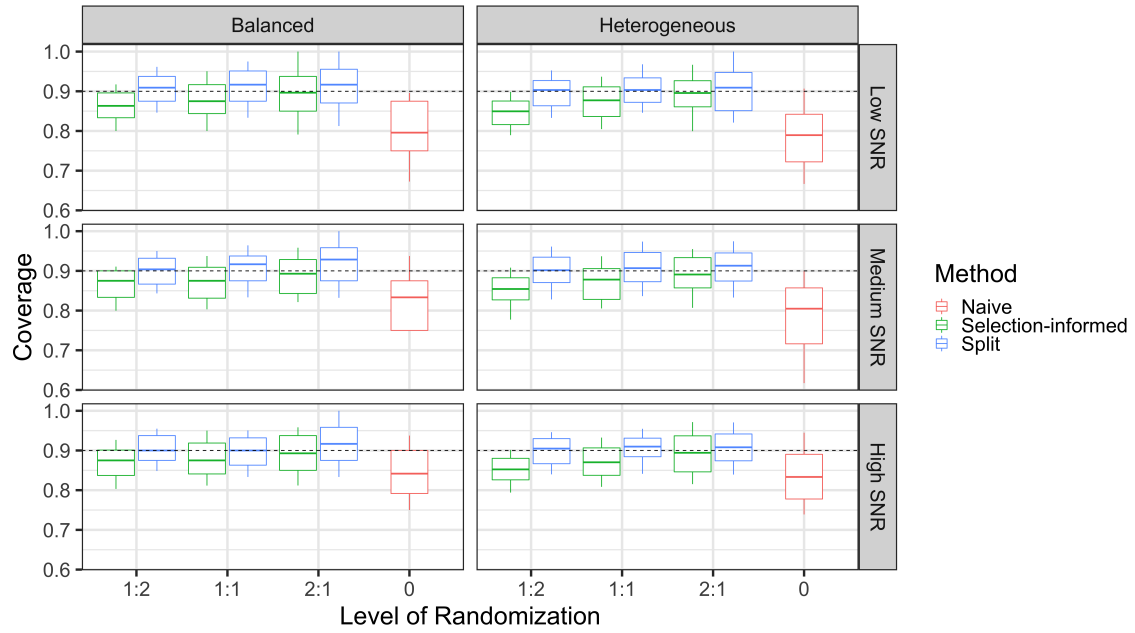


Figure 6: Box plots for coverage of credible intervals post the Group LASSO for the balanced and heterogeneous groups.

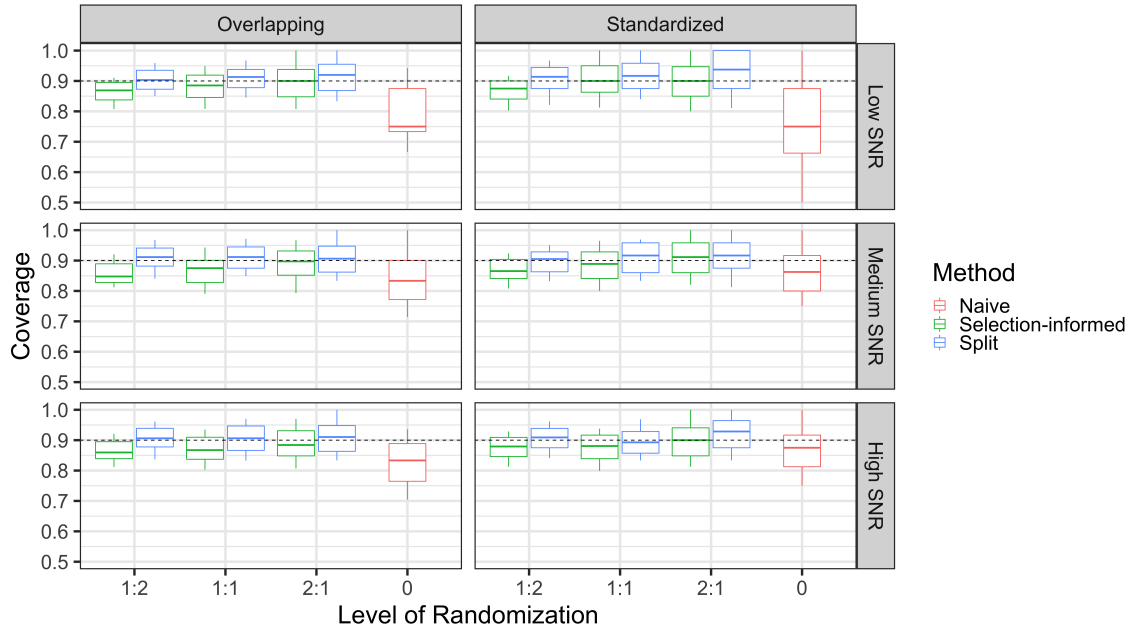


Figure 7: Box plots for coverage of credible intervals post the two extensions of the Group LASSO.

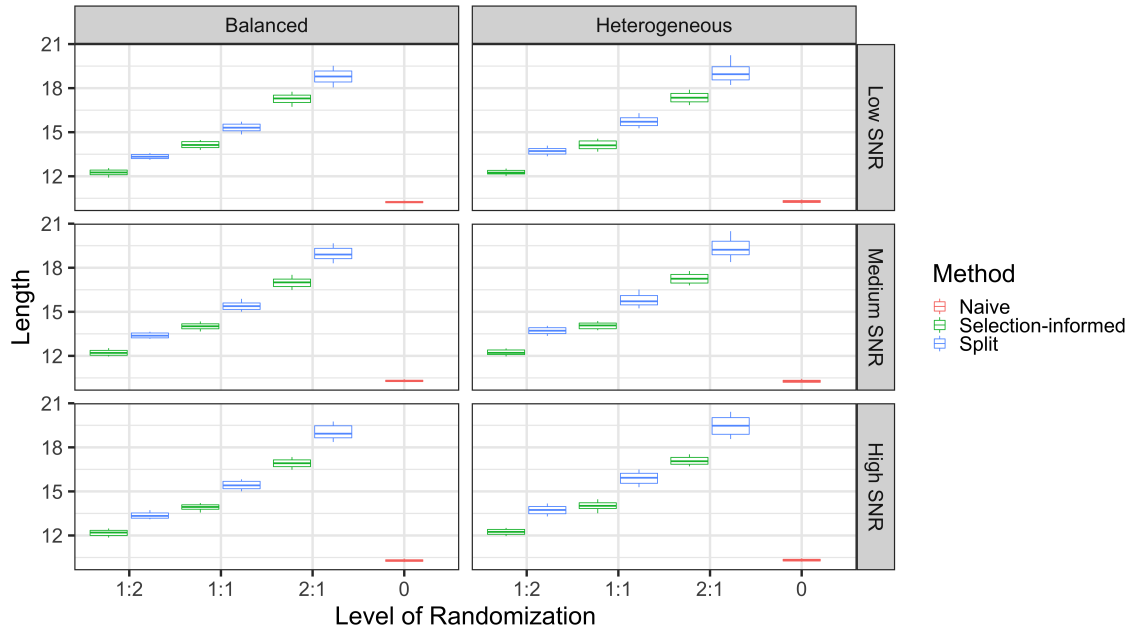


Figure 8: Box plots for lengths of credible intervals post the canonical Group LASSO for the balanced and heterogeneous groups.

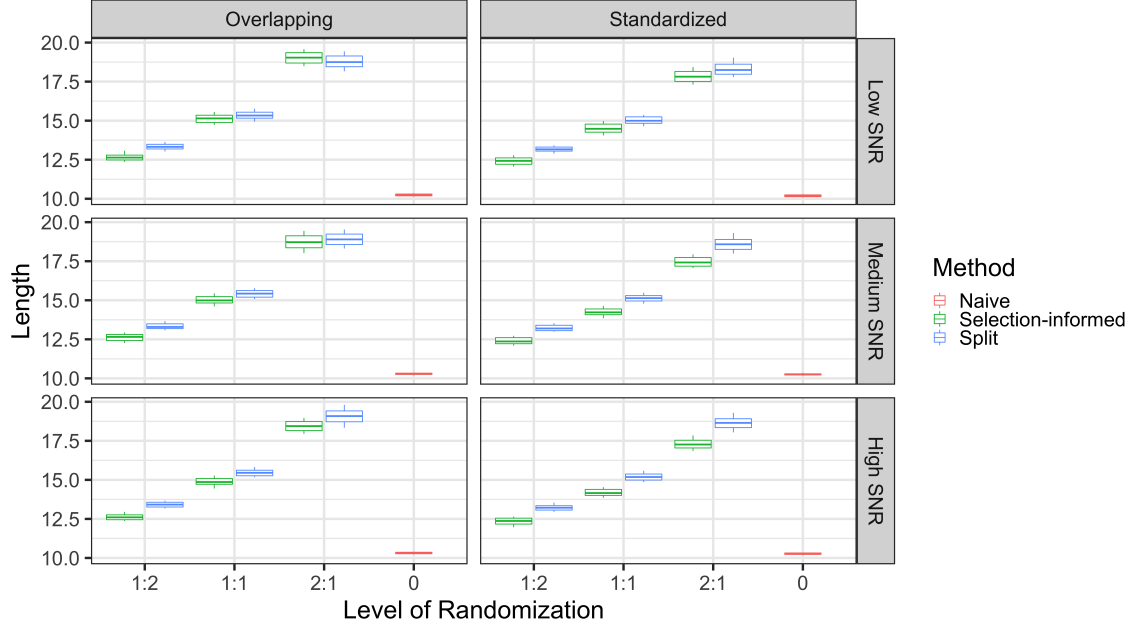


Figure 9: Box plots for lengths of credible intervals post the two extensions of the Group LASSO.

7.3 Application to Neuroimaging Data

We apply our method to a subset ($n = 785$) of human neuroimaging data from the Human Connectome Project (HCP) (Van Essen et al., 2013), a landmark study undertaken by a consortium involving Washington University, the University of Minnesota, and Oxford University. The HCP has led to a substantial advancement of human neuroimaging methodology and included the collection of several corpora of data which are available to researchers interested in studying brain function and connectivity.

We consider below a linear model to understand participant accuracy on a working memory task using brain activity measured at $p = 236$ locations in the brain during performance of the task, using data graciously processed by the lab of our collaborator (see Acknowledgements). Using both behavioral and functional magnetic resonance imaging (fMRI) measurements recorded from a cognitive task, a standardized measure of accuracy for each participant during this task will be our response y and contrasts relying upon brain activation records during the task form our covariates X . We provide a summary of these details in Appendix D.1, accompanied by a description for the preprocessing steps and parameter settings. Our analysis here groups the covariates by brain system and applies our selection-informed method to calibrate interval estimates for coefficients within any selected systems.

We apply both the randomized Group LASSO and the Group LASSO to a random split of this data set. We consider the level of isotropic Gaussian randomization to be “1:1,” “2:1,” and “9:1” by setting the variance parameter τ^2 according to (23) after replacing σ^2

with

$$\hat{\sigma}^2 = (n - p)^{-1} \left\| y - X (X^\top X)^{-1} X^\top y \right\|_2^2,$$

for $r = 1/2, 2/3, 9/10$, respectively. In all cases, we select one group: the “Fronto-parietal Task Control” (FP) system. To restore inferential validity, we reuse data via our “Selection-informed” method to draw samples from the surrogate selection-adjusted posterior. Given the overlap between data sets employed, it is not surprising that Sripada et al. (2020) also found that activation in the Fronto-parietal Task Control system could be used to predict “General Cognitive Ability” (GCA): this general pattern that included activation in the fronto-parietal system and deactivation in the default mode system under general cognitive demands (including working memory) is discussed and reviewed in Sripada et al. (2020). We depict 90% interval estimates for each of the locations in the brain within the selected FP system under both the “Selection-informed” and “Split” methods in Figure 10 at varying levels of randomization. Corroborating the general pattern from our numerical experiments, the intervals from our “Selection-informed” methods are roughly 8% shorter than those from “Split” on an average in the application.

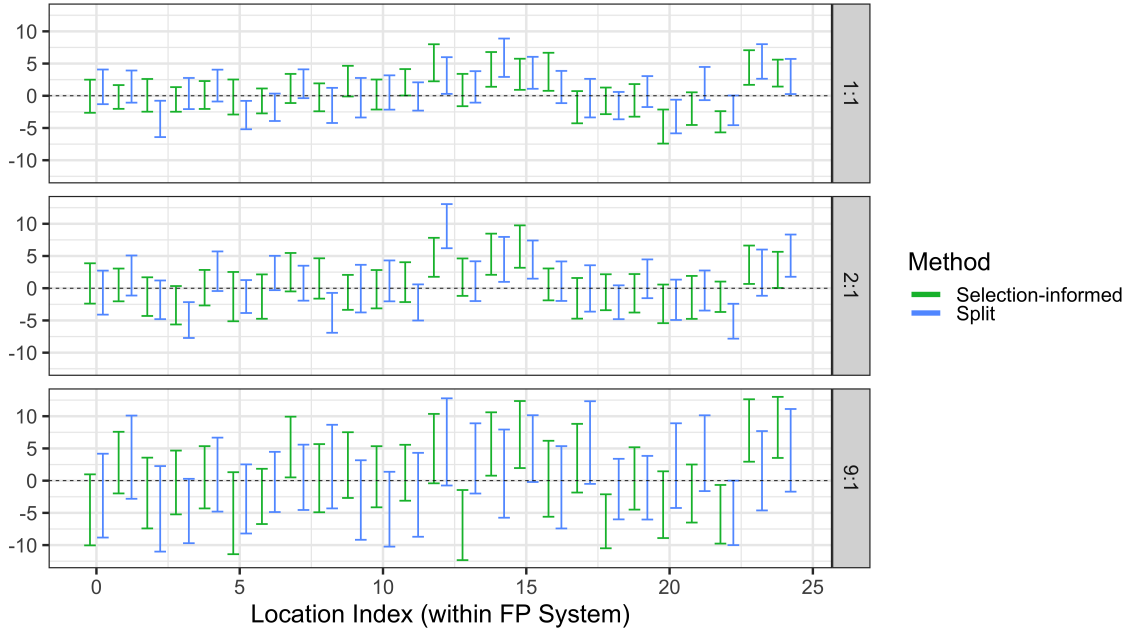


Figure 10: Interval estimates for coefficients at each location within the selected FP brain system for human neuroimaging application. Intervals are provided for both our “Selection-informed” method and “Split” at a variety of randomization levels. At the levels of randomization “1:1,” “2:1,” and “9:1,” average interval widths for (“Selection-informed”, “Split”) are (4.72, 5.10), (5.99, 6.17), (10.03, 11.77).

8. Conclusion

In this paper, we provide methods to account for the selection-informed nature of models after solving group-sparse learning algorithms. Deriving conditional inference in these settings is particularly challenging due to the unavailability of a polyhedral representation for the selection event. Formally cast into a Bayesian framework, we successfully characterize an exact adjustment factor to account for selections of grouped variables and importantly, provide a computationally feasible solution to bridge the gap between theory and practice for a general class of group-sparse models. Appealingly amenable to a large class of grouped sparsities and context-relevant targets, the efficiency of our methods is evident from the minimal price we pay to correct for selection in comparison to naive inference.

Our work leaves room for promising directions of research which we hope to take on as future investigations. The performance of our methods serve as an encouraging direction to tackle non-affine geometries in general, seen often with penalties disparate in behavior from ℓ_1 -sparsity imposing algorithms. These developments do not preclude an asymptotic framework for valid inference after the Group LASSO, when we deviate from Gaussian distributions. Lastly, grouped selections form the first step of many hierarchical exploratory pipelines (for example, in genomic studies) to select predictors with better meaning and accuracy. Our solutions for reusing data after a class of grouped selection rules hold the potential to infer using automated models from these complicated selection pipelines.

Acknowledgments

S.P. was supported by NSF-DMS 1951980 and NSF-DMS 2113342. D.K. was supported by NSF-DMS 1646108 and a Rackham Predoctoral Fellowship from the University of Michigan. P.W.M. was supported by NSF-DMS 1916222 and a Rackham Predoctoral Fellowship from the University of Michigan. Data were provided in part by the Human Connectome Project, WU-Minn Consortium (Principal Investigators: David Van Essen and Kamil Ugurbil; 1U54MH091657) funded by the 16 NIH Institutes and Centers that support the NIH Blueprint for Neuroscience Research; and by the McDonnell Center for Systems Neuroscience at Washington University. We thank Dr. Chandra Sripada and his research group (with particular thanks to Saige Rutherford) for providing a processed version of this data as well as helpful comments. This research was supported in part through computational resources and services provided by Advanced Research Computing (ARC), a division of Information and Technology Services (ITS) at the University of Michigan, Ann Arbor. In addition, this work used the Extreme Science and Engineering Discovery Environment (XSEDE), which is supported by National Science Foundation grant number ACI-1548562.

Appendix A. Proofs for Technical Developments (Section 4)

Proof [Proof of Theorem 2.] We write our adjustment factor as follows:

$$\mathbb{P}(\mathcal{A}_E \mid \beta_E) = \int \int p(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot p(0, \Omega; \omega) \cdot \mathbf{1}_{\mathcal{A}_E}(\hat{\beta}_E, \omega) \, d\omega d\hat{\beta}_E.$$

Next, the change of variables in (9) together with the conditioning upon $\widehat{\mathcal{U}} = \mathcal{U}$ and $\widehat{\mathcal{Z}} = \mathcal{Z}$ results in the below simplification

$$\begin{aligned} \mathbb{P}(\mathcal{A}_E \mid \beta_E) &= \int \int J_{\phi_{\widehat{\beta}_E}}(\widehat{\gamma}; \mathcal{U}, \mathcal{Z}) \cdot p(\beta_E, \Sigma_E; \widehat{\beta}_E) \cdot p(0, \Omega; \phi_{\widehat{\beta}_E}(\widehat{\gamma}, \mathcal{U}, \mathcal{Z})) \cdot \mathbf{1}(\widehat{\gamma} > 0) d\widehat{\gamma} d\widehat{\beta}_E \\ &\propto \int \int J_{\phi_{\widehat{\beta}_E}}(\widehat{\gamma}; \mathcal{U}, \mathcal{Z}) \cdot p(\beta_E, \Sigma_E; \widehat{\beta}_E) \cdot \exp \left\{ -\frac{1}{2} (A\widehat{\beta}_E + B(\mathcal{U})\widehat{\gamma} + c(\mathcal{U}, \mathcal{Z}))^\top \Omega^{-1} \right. \\ &\quad \left. (A\widehat{\beta}_E + B(\mathcal{U})\widehat{\gamma} + c(\mathcal{U}, \mathcal{Z})) \right\} \cdot \mathbf{1}(\widehat{\gamma} > 0) d\widehat{\gamma} d\widehat{\beta}_E \end{aligned} \quad (24)$$

where $J_{\phi_{\widehat{\beta}_E}}(\gamma; \mathcal{U}, \mathcal{Z})$ is the Jacobian associated with the change of variables $\phi_{\widehat{\beta}_E}(\cdot)$. To compute the nontrivial Jacobian, we write the change of variables map in (9) as follows:

$$(\widehat{\gamma}, \mathcal{U}, \mathcal{Z}) \xrightarrow{\phi} (\phi_1, \phi_2);$$

$$\begin{aligned} \phi_1(\widehat{\beta}_E, \widehat{\gamma}, \mathcal{U}, \mathcal{Z}) &= QU\widehat{\gamma} - Q\widehat{\beta}_E - ((N_{E,j})_{j \in E}) + ((\lambda_g u_g)_{g \in \mathcal{G}_E}), \\ \phi_2(\widehat{\beta}_E, \widehat{\gamma}, \mathcal{U}, \mathcal{Z}) &= X_{-E}^\top X_E (U\widehat{\gamma} - \widehat{\beta}_E) - ((N_{E,j})_{j \in -E}) + ((\lambda_g z_g)_{g \in -\mathcal{G}_E}), \end{aligned}$$

such that $Q = X_E^\top X_E$. Then, the derivative matrix is given by:

$$D_{\phi_{\widehat{\beta}_E}} = \begin{pmatrix} \frac{\partial \phi_1}{\partial \mathcal{U}} & \frac{\partial \phi_1}{\partial \widehat{\gamma}} & \frac{\partial \phi_1}{\partial \mathcal{Z}} \\ \frac{\partial \phi_2}{\partial \mathcal{U}} & \frac{\partial \phi_2}{\partial \widehat{\gamma}} & \frac{\partial \phi_2}{\partial \mathcal{Z}} \end{pmatrix},$$

where $\frac{\partial}{\partial \mathcal{U}}(\cdot)$ refers to differentiation with respect to each u_g in the coordinates of its tangent space. Note that the block $\frac{\partial \phi_1}{\partial \mathcal{Z}}$ above the diagonal is zero and $\det \left(\frac{\partial \phi_2}{\partial \mathcal{Z}} \right) \propto 1$. Thus, it follows that

$$J_{\phi_{\widehat{\beta}_E}}(\widehat{\gamma}; \mathcal{U}, \mathcal{Z}) = \det(D_{\phi_{\widehat{\beta}_E}}) \propto \det \left(\begin{pmatrix} \frac{\partial \phi_1}{\partial \mathcal{U}} & \frac{\partial \phi_1}{\partial \widehat{\gamma}} \end{pmatrix} \right).$$

First, it is easy to see $\frac{\partial \phi_1}{\partial \widehat{\gamma}} = QU$. For computing the other block $\frac{\partial \phi_1}{\partial \mathcal{U}}$, let $u_g \in \mathcal{S}^{|g|-1}$ be associated with the tangent space: $T_{u_g} \mathcal{S}^{|g|-1} = \{v : v^\top u_g = 0\}$, the orthogonal complement of $\text{span}\{u_g\}$; $\bar{U}_g$ is a fixed orthonormal basis for this tangent space. For a vector of coordinates $y_g \in \mathbb{R}^{|g|-1}$ and a general function h ,

$$\frac{\partial h(u_g)}{\partial u_g} := \frac{\partial h(u_g + \bar{U}_g y_g)}{\partial y_g}.$$

Writing this more compactly with a stacked vector $y = (y_1, \dots, y_{|\mathcal{G}_E|})^\top$, it follows that for fixed g ,

$$\begin{aligned} \frac{\partial \phi_1}{\partial u_g} &= \frac{\partial}{\partial y_g} \{Q(U + \bar{U}y)\widehat{\gamma} + ((\lambda_g(u_g + \bar{U}_g y_g))_{g \in E})\} \\ &= Q \begin{pmatrix} 0 & \cdots & (\widehat{\gamma}_g \bar{U}_g)^\top & \cdots & 0 \end{pmatrix}^\top + \begin{pmatrix} 0 & \cdots & (\lambda_g \bar{U}_g)^\top & \cdots & 0 \end{pmatrix}^\top, \end{aligned}$$

and combining these column-wise we obtain the full derivative matrix

$$\begin{aligned}\frac{\partial \phi_1}{\partial \mathcal{U}} &= \begin{pmatrix} \frac{\partial \phi_1}{\partial u_1} & \cdots & \frac{\partial \phi_1}{\partial u_{|\mathcal{G}_E|}} \end{pmatrix} = Q \cdot \text{diag}((\hat{\gamma}_g \bar{U}_g)_{g \in E}) + \text{diag}((\lambda_g \bar{U}_g)_{g \in E}) \\ &= Q \bar{\Gamma} \bar{U} + \Lambda \bar{U}; \quad \text{where } \bar{\Gamma} = \text{diag}((\hat{\gamma}_g I_{|g|})_{g \in \mathcal{G}_E}), \\ &= (Q \bar{\Gamma} + \Lambda) \bar{U}.\end{aligned}$$

This gives us

$$\det(D_{\phi_{\hat{\beta}_E}}) \propto \det((Q \bar{\Gamma} + \Lambda) \bar{U} \quad QU). \quad (25)$$

Simplifying this expression further

$$\begin{aligned}\det(D_{\phi_{\hat{\beta}_E}}) &\propto \det\left(\begin{pmatrix} \bar{U}^\top \\ U^\top \end{pmatrix} (\bar{\Gamma} \bar{U} + Q^{-1} \Lambda \bar{U} \quad U)\right) \\ &= \det\left(\begin{pmatrix} \Gamma + \bar{U}^\top Q^{-1} \Lambda \bar{U} & 0 \\ \cdots & I_{|E|} \end{pmatrix}\right) \\ &= \det(\Gamma + \bar{U}^\top Q^{-1} \Lambda \bar{U}),\end{aligned}$$

This follows since $\bar{U}^\top \bar{\Gamma} \bar{U} = \Gamma$ by block orthogonality, and the final equality is deduced using block triangularity. This proves our claim in the Theorem. $\blacksquare$

Proof [Proof of Proposition 4.] Ignoring the selection-informed prior for now, our likelihood after conditioning upon the event $\mathcal{A}_E$ is given by

$$\frac{\text{p}(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot \int J_\phi(\hat{\gamma}, \mathcal{U}) \exp\left\{-\frac{1}{2}(A\hat{\beta}_E + B\hat{\gamma} + c)^\top \Omega^{-1}(A\hat{\beta}_E + B\hat{\gamma} + c)\right\} \cdot \mathbf{1}(\hat{\gamma} > 0) d\hat{\gamma}}{\int J_\phi(\hat{\gamma}, \mathcal{U}) \cdot \text{p}(\beta_E, \Sigma_E; \hat{\beta}_E) \cdot \exp\left\{-\frac{1}{2}(A\hat{\beta}_E + B\hat{\gamma} + c)^\top \Omega^{-1}(A\hat{\beta}_E + B\hat{\gamma} + c)\right\} \cdot \mathbf{1}(\hat{\gamma} > 0) d\hat{\gamma} d\hat{\beta}_E}.$$

We note that this expression is proportional to:

$$\begin{aligned}\left(\int J_\phi(\hat{\gamma}, \mathcal{U}) \text{p}(\bar{R}\beta_E + \bar{s}, \bar{\Theta}; \hat{\beta}_E) \cdot \text{p}(\bar{A}\hat{\beta}_E + \bar{b}, \bar{\Omega}; \hat{\gamma}) \cdot \mathbf{1}(\hat{\gamma} > 0) d\hat{\gamma} d\hat{\beta}_E\right)^{-1} \\ \times \text{p}(\bar{R}\beta_E + \bar{s}, \bar{\Theta}; \hat{\beta}_E)\end{aligned}$$

leaving out the constants in β_E , where

$$\bar{\Omega} = (B^\top \Omega^{-1} B)^{-1}, \quad \bar{A} = -\bar{\Omega} B^\top \Omega^{-1} A, \quad \bar{b} = -\bar{\Omega} B^\top \Omega^{-1} c,$$

$$\bar{\Theta} = (\Sigma_E^{-1} - (\bar{A})^\top (\bar{\Omega})^{-1} \bar{A} + A^\top \Omega^{-1} A)^{-1}, \quad \bar{R} = \bar{\Theta} \Sigma_E^{-1}, \quad \bar{s} = \bar{\Theta} ((\bar{A})^\top (\bar{\Omega})^{-1} \bar{b} - A^\top \Omega^{-1} c).$$

This display relies on the observation that

$$\begin{aligned}\text{p}(\beta_E, \Sigma_E; \hat{\beta}_E) \exp\left\{-\frac{1}{2}(A\hat{\beta}_E + B\hat{\gamma} + c)^\top \Omega^{-1}(A\hat{\beta}_E + B\hat{\gamma} + c)\right\} \\ = K(\beta_E) \cdot \text{p}(\bar{R}\beta_E + \bar{s}, \bar{\Theta}; \hat{\beta}_E) \cdot \text{p}(\bar{A}\hat{\beta}_E + \bar{b}, \bar{\Omega}; \hat{\gamma}),\end{aligned}$$

such that $K(\beta_E)$ involves β_E alone. ■

Proof [Proof of Theorem 5.] To derive the expression in (5), we note that optimizing over $\tilde{\beta}_E$ in the problem:

$$\begin{aligned} \text{minimize}_{\tilde{\beta}_E, \tilde{\gamma}} \left\{ \frac{1}{2}(\tilde{\beta}_E - \bar{R}\beta_E - \bar{s})^\top \bar{\Theta}^{-1}(\tilde{\beta}_E - \bar{R}\beta_E - \bar{s}) \right. \\ \left. + \frac{1}{2}(\tilde{\gamma} - \bar{A}\tilde{\beta}_E - \bar{b})^\top (\bar{\Omega})^{-1}(\tilde{\gamma} - \bar{A}\tilde{\beta}_E - \bar{b}) + \text{Barr}(\tilde{\gamma}) \right\} \end{aligned}$$

gives us

$$\text{minimize}_{\gamma} \frac{1}{2}(\gamma - \bar{P}\beta_E - \bar{q})^\top (\bar{\Sigma})^{-1}(\gamma - \bar{P}\beta_E - \bar{q}) + \text{Barr}(\gamma).$$

Plugging the value of this optimization into (11), the logarithm of the surrogate posterior is given by the expression

$$\begin{aligned} \log \pi_E(\beta_E) + \log p(\bar{R}\beta_E + \bar{s}, \bar{\Theta}; \hat{\beta}_E) + \frac{1}{2}(\gamma^* - \bar{P}\beta_E - \bar{q})^\top (\bar{\Sigma})^{-1}(\gamma^* - \bar{P}\beta_E - \bar{q}) \\ + \text{Barr}(\gamma^*) - \log J_\phi(\gamma^*; \mathcal{U}), \end{aligned}$$

ignoring additive constants. To compute the gradient of the (log) surrogate posterior, define $\zeta^*(\cdot)$ to be the convex conjugate for the function

$$\zeta(\gamma) = \frac{1}{2}\gamma^\top (\bar{\Sigma})^{-1}\gamma + \text{Barr}(\gamma).$$

This allows us to write (5) as

$$\begin{aligned} \log \pi_E(\beta_E) - \frac{1}{2}(\hat{\beta}_E - \bar{R}\beta_E - \bar{s})^\top (\bar{\Theta})^{-1}(\hat{\beta}_E - \bar{R}\beta_E - \bar{s}) \\ + \frac{1}{2}(\bar{P}\beta_E + \bar{q})^\top (\bar{\Sigma})^{-1}(\bar{P}\beta_E + \bar{q}) - \zeta^*((\bar{\Sigma})^{-1}(\bar{P}\beta_E + \bar{q})) - \log J_\phi(\gamma^*; \mathcal{U}). \end{aligned} \tag{26}$$

Denoting $L(\beta_E) = (\bar{\Sigma})^{-1}(\bar{P}\beta_E + \bar{q})$ and taking the derivative of (26) with respect to β_E gives us

$$\begin{aligned} \nabla \log \pi_E(\beta_E) + (\bar{R})^\top (\bar{\Theta})^{-1}(\hat{\beta}_E - \bar{R}\beta_E - \bar{s}) + \bar{P}^\top (\bar{\Sigma})^{-1}(\bar{P}\beta_E + \bar{q}) \\ - \bar{P}^\top (\bar{\Sigma})^{-1} \nabla_{L(\beta_E)} \zeta^*((\bar{\Sigma})^{-1}(\bar{P}\beta_E + \bar{q})) - \bar{P}^\top (\bar{\Sigma})^{-1} \nabla_{L(\beta_E)} \gamma^*((\bar{\Sigma})^{-1}(\bar{P}\beta_E + \bar{q})) \nabla_{\gamma^*} \log J_\phi(\gamma^*; \mathcal{U}) \\ = \nabla \log \pi_E(\beta_E) + (\bar{R})^\top (\bar{\Theta})^{-1}(\hat{\beta}_E - \bar{R}\beta_E - \bar{s}) + \bar{P}^\top (\bar{\Sigma})^{-1}(\bar{P}\beta_E + \bar{q}) \\ - \bar{P}^\top (\bar{\Sigma})^{-1} (\nabla_{L(\beta_E)} \zeta)^{-1}(L(\beta_E)) - \bar{P}^\top (\bar{\Sigma})^{-1} \nabla_{L(\beta_E)} \gamma^*(L(\beta_E)) \nabla_{\gamma^*} \log J_\phi(\gamma^*; \mathcal{U}) \\ = \nabla \log \pi_E(\beta_E) + (\bar{R})^\top (\bar{\Theta})^{-1}(\hat{\beta}_E - \bar{R}\beta_E - \bar{s}) + \bar{P}^\top (\bar{\Sigma})^{-1}(\bar{P}\beta_E + \bar{q}) \\ - \bar{P}^\top (\bar{\Sigma})^{-1} \gamma^*(L(\beta_E)) - \bar{P}^\top (\bar{\Sigma})^{-1} \nabla_{L(\beta_E)} \gamma^*(L(\beta_E)) \nabla_{\gamma^*} \log J_\phi(\gamma^*; \mathcal{U}). \end{aligned}$$

Note that in the second display we use the fact that $\nabla\zeta^*(\cdot) = (\nabla\zeta)^{-1}(\cdot)$, and the third display follows by observing $(\nabla\zeta)^{-1}(L(\beta_E)) = \gamma^*(L(\beta_E))$, where

$$\gamma^* = \operatorname{argmax}_{\gamma} \gamma^\top L(\beta_E) - \zeta(\gamma),$$

which is equal to the optimizer defined in (12). Computing the two pieces in the final term, $\nabla_{L(\beta_E)}\gamma^*(L(\beta_E))$ and $\nabla_{\gamma^*} \log J_\phi(\gamma^*; \mathcal{U})$, we first have

$$\nabla\gamma^*(\cdot) = \nabla^2\zeta^*(\cdot) = [\nabla^2\zeta(\gamma^*(\cdot))]^{-1}.$$

Since $\nabla^2\zeta(\gamma) = \bar{\Sigma}^{-1} + \nabla^2 \text{Barr}(\gamma)$, we have

$$\nabla_{L(\beta_E)}\gamma^*(L(\beta_E)) = (\bar{\Sigma}^{-1} + \nabla^2 \text{Barr}(\gamma^*(L(\beta_E))))^{-1}.$$

As for $\nabla \log J_\phi(\gamma^*; \mathcal{U})$, recall from Theorem 2

$$J_\phi(\gamma; \mathcal{U}) = \det(\Gamma + \bar{U}^\top (X_E^\top X_E)^{-1} \Lambda \bar{U}).$$

We have the following matrix derivative identities for a square matrix X (Petersen and Pedersen, 2012):

$$\frac{\partial \log \det(X)}{\partial X_{ij}} = (X^{-1})_{ji}, \quad \frac{\partial (X^{-1})_{kl}}{\partial X_{ij}} = -(X^{-1})_{ki} (X^{-1})_{jl}. \quad (27)$$

Using the chain rule,

$$\gamma \xrightarrow{J_1} (\Gamma + \bar{U}^\top (X_E^\top X_E)^{-1} \Lambda \bar{U}) \xrightarrow{J_2} \log \det(\Gamma + \bar{U}^\top (X_E^\top X_E)^{-1} \Lambda \bar{U}),$$

we partition the indices into $\{M_g\}_{g \in E}$ such that M_g is the set of $|g| - 1$ indices along the diagonal of Γ corresponding to group g . Then

$$\left[\frac{\partial J_1}{\partial \gamma_g} \right]_{ij} = \begin{cases} 1, & i = j, i \in M_g, \\ 0, & \text{otherwise.} \end{cases}$$

By (27), the partial derivatives of J_2 are the entries of $(\Gamma + \bar{U}^\top (X_E^\top X_E)^{-1} \Lambda \bar{U})^{-1}$. Putting these together,

$$\frac{\partial \log J_\phi(\cdot; \mathcal{U})}{\partial \gamma_g} = \sum_{i \in M_g} [(\Gamma + \bar{U}^\top (X_E^\top X_E)^{-1} \Lambda \bar{U})^{-1}]_{ii},$$

which gives us the expression for $\nabla_{\gamma^*} \log J_\phi(\gamma^*; \mathcal{U})$. ■

We conclude with a remark highlighting the distinction from the usual Laplace-type approximation, which we adopt for tractable calculations of the adjustment factor. An

alternate approximation for (11) is given by

$$C \exp \left(-\frac{1}{2}(\beta_E^* - \beta_E)^\top \Sigma_E^{-1}(\beta_E^* - \beta_E) - \frac{1}{2}(\gamma^* - A^* \beta_E^* - b^*)^\top (\Sigma^*)^{-1}(\gamma^* - A^* \beta_E^* - b^*) - \text{Barr}(\gamma^*) + \log J_\phi(\gamma^*; \mathcal{U}) \right),$$

the usual Laplace approximation where γ^* and β_E^* are obtained by solving

$$\begin{aligned} \text{minimize}_{\tilde{\beta}_E, \tilde{\gamma}} \left\{ \frac{1}{2}(\tilde{\beta}_E - \beta_E)^\top \Sigma_E^{-1}(\tilde{\beta}_E - \beta_E) + \frac{1}{2}(\tilde{\gamma} - A^* \tilde{\beta}_E - b^*)^\top (\Sigma^*)^{-1}(\tilde{\gamma} - A^* \tilde{\beta}_E - b^*) \right. \\ \left. + \text{Barr}(\tilde{\gamma}) - \log J_\phi(\tilde{\gamma}; \mathcal{U}) \right\}. \end{aligned} \quad (28)$$

Problem (28) deviates from our current formulation (12) in terms of the part the (log) Jacobian term plays in determining the mode of the optimization. We opt specifically for a generalized formulation of the Laplace approximation to compute the Jacobian only once at the mode of (12) for increased computational efficiency, obtaining our selection-informed posterior and the gradient associated with it.

Proof [Proof of Proposition 6.] The proof of this Proposition follows by applying the change of variables map:

$$\omega^* \rightarrow (\hat{\gamma}^*, \hat{\mathcal{U}}^*, \hat{\mathcal{Z}}^*) \text{ where } (\hat{\gamma}^*, \hat{\mathcal{U}}^*, \hat{\mathcal{Z}}^*) = (\phi^*)^{-1}(\omega^*),$$

and conditioning upon $\hat{\mathcal{U}}^* = \mathcal{U}^*$ and $\hat{\mathcal{Z}}^* = \mathcal{Z}^*$. This leads to the below adjustment factor, the probability of the selection event under consideration,

$$\begin{aligned} \mathbb{P}(\mathcal{A}_{E^*} \mid \beta_E) &= \int \int J_{\phi^*}(\hat{\gamma}^*; \mathcal{U}^*) \cdot p(\beta_E, \Sigma_E; \hat{\beta}_E) \\ &\times \exp \left\{ -\frac{1}{2}(A\hat{\beta}_E + B(\mathcal{U})\hat{\gamma}^* + c(\mathcal{U}, \mathcal{Z}))^\top \Omega^{-1}(A\hat{\beta}_E + B(\mathcal{U})\hat{\gamma}^* + c(\mathcal{U}, \mathcal{Z})) \right\} \cdot \mathbf{1}(\hat{\gamma}^* > 0) d\hat{\gamma}^* d\hat{\beta}_E, \end{aligned} \quad (29)$$

$J_{\phi^*}(\hat{\gamma}^*; \mathcal{U}^*)$ is the Jacobian associated with the change of variables derived from $\phi^*(\cdot)$. To complete the proof, we note that the value for $J_{\phi^*}(\hat{\gamma}^*; \mathcal{U}^*)$ is obtained from (25) in the derivation of the adjustment factor when there are no overlaps in the groups, where we simply replace the original design with the augmented version. $\blacksquare$

Proof [Proof of Proposition 7.] The proof is direct from using the change of variables map from inverting the stationary mapping we identify for the solver (17). Based upon the matrices we identify in (19), the argument follows similar lines as Theorem 2 yielding the expression for the adjustment factor. $\blacksquare$

Proof [Proof of Proposition 8.] Modifying the proof of Theorem 2 by replacing the stationary mapping with $\check{\phi}(\cdot)$ defined in (20) results in the claim in this Proposition. We thus omit further details of the proof here. $\blacksquare$

Appendix B. Proofs for Large Sample Theory (Section 6)

We provide in this section the proofs of the large sample claims for our selection-informed posterior in Section 6. Supporting results for this theory, Lemma 13 and 14, are included in Section C.

Proof [Proof of Proposition 9.] First, observe that we write our probability of selection as follows:

$$(b_n)^{-2} \log \mathbb{P} (0 < \sqrt{n}\gamma_n < b_n \bar{Q} \cdot 1_{|E|} + \bar{q}) = (b_n)^{-2} \log \mathbb{E} \left[\exp(\log J_\phi(\sqrt{n}\bar{Z}_n + b_n \bar{P}\bar{\beta}_E + \bar{q}; \mathcal{U})) \times \mathbf{1}(-b_n \bar{P}\bar{\beta}_E - \bar{q} < \sqrt{n}\bar{Z}_n < b_n \bar{Q} \cdot 1_{|E|} - b_n \bar{P}\bar{\beta}_E) \right],$$

(up to an additive constant), where $\sqrt{n}\bar{Z}_n$ is a centered Gaussian random variable with covariance $\bar{\Sigma}$. Define $\mathcal{C}_0 = \{z : -\bar{P}\bar{\beta}_E < z < \bar{Q} \cdot 1_{|E|} - \bar{P}\bar{\beta}_E\}$. Using assumption (22), we deduce

$$\begin{aligned} & \limsup_{n \rightarrow \infty} (b_n)^{-2} \log \mathbb{P} (0 < \sqrt{n}\gamma_n < b_n \bar{Q} \cdot 1_{|E|} + \bar{q}) \\ &= \limsup_{n \rightarrow \infty} (b_n)^{-2} \log \mathbb{E} \left[\exp(\log J_\phi(\sqrt{n}\bar{Z}_n + b_n \bar{P}\bar{\beta}_E + \bar{q}; \mathcal{U})) \cdot 1_{\mathcal{C}_0}(\sqrt{n}\bar{Z}_n/b_n) \right] \\ &\leq \lim_{n \rightarrow \infty} \sup_{z \in \mathcal{C}_0} (b_n)^{-2} |\log J_\phi(b_n z + b_n \bar{P}\bar{\beta}_E + \bar{q}; \mathcal{U})| + \lim_{n \rightarrow \infty} (b_n)^{-2} \log \mathbb{P}(\sqrt{n}\bar{Z}_n/b_n \in \mathcal{C}_0) \\ &= \lim_{n \rightarrow \infty} (b_n)^{-2} \log \mathbb{P}(-\bar{P}\bar{\beta}_E < \sqrt{n}\bar{Z}_n/b_n < \bar{Q} \cdot 1_{|E|} - \bar{P}\bar{\beta}_E). \end{aligned}$$

To justify that the limit of the term involving the Jacobian vanishes, note that for all $z \in \mathcal{C}_0$, we have

$$\bar{q} < b_n z + b_n \bar{P}\bar{\beta}_E + \bar{q} < b_n \bar{Q} 1_{|E|} + \bar{q}.$$

Thus, $J_\phi(b_n z + b_n \bar{P}\bar{\beta}_E + \bar{q}; \mathcal{U})$ is the determinant of a matrix with entries uniformly bounded by $2b_n \bar{Q} > 0$ for sufficiently large n . Then

$$\sup_{z \in \mathcal{C}_0} |J_\phi(b_n z + b_n \bar{P}\bar{\beta}_E + \bar{q}; \mathcal{U})| \leq |E|! (2b_n \bar{Q})^{|E|}.$$

The Jacobian is also bounded away from zero, thus it follows that

$$\sup_{z \in \mathcal{C}_0} (b_n)^{-2} |\log J_\phi(b_n z + b_n \bar{P}\bar{\beta}_E + \bar{q}; \mathcal{U})| \leq (b_n)^{-2} \log \left(|E|! (2b_n \bar{Q})^{|E|} \right)$$

which goes to 0 as $n \rightarrow \infty$. Using an argument along the same line,

$$\begin{aligned} & \liminf_{n \rightarrow \infty} (b_n)^{-2} \log \mathbb{P} (0 < \sqrt{n}\gamma_n < b_n \bar{Q} \cdot 1_{|E|} + \bar{q}) \\ &\geq - \lim_{n \rightarrow \infty} \sup_{z \in \mathcal{C}_0} (b_n)^{-2} |\log J_\phi(b_n z + b_n \bar{P}\bar{\beta}_E + \bar{q}; \mathcal{U})| + \lim_{n \rightarrow \infty} (b_n)^{-2} \log \mathbb{P}(\sqrt{n}\bar{Z}_n/b_n \in \mathcal{C}_0). \end{aligned}$$

From the above limits, we have

$$\lim_{n \rightarrow \infty} (b_n)^{-2} \left(\log \mathbb{P} (0 < \sqrt{n}\gamma_n < b_n \bar{Q} \cdot 1_{|E|} + \bar{q}) - \log \mathbb{P}(\sqrt{n}\bar{Z}_n/b_n \in \mathcal{C}_0) \right) = 0.$$

Using a moderate (large)-deviation type result (De Acosta, 1992) for the limiting value of the probability in the second term, we have

$$\lim_{n \rightarrow \infty} (b_n)^{-2} \log \mathbb{P} (0 < \sqrt{n} \gamma_n < b_n \bar{Q} \cdot 1_{|E|} + \bar{q}) + \inf_{z \in \mathcal{C}_0} z^\top \bar{\Sigma}^{-1} z / 2 = 0.$$

Lastly, we let $z + \bar{P} \bar{\beta}_E + (b_n)^{-1} \bar{q} = \bar{\gamma}$. Using the observation that the optimization $\inf_{z \in \mathcal{C}_0} z^\top \bar{\Sigma}^{-1} z$ has a unique minimum, and relying on the convexity of the objectives in the sequence of optimization problems defined below

$$\inf_{\bar{\gamma} < \bar{Q} \cdot 1_{|E|}} \left\{ \frac{1}{2} (\bar{\gamma} - \bar{P} \bar{\beta}_E - (b_n)^{-1} \bar{q})^\top \bar{\Sigma}^{-1} (\bar{\gamma} - \bar{P} \bar{\beta}_E - (b_n)^{-1} \bar{q}) + (b_n)^{-2} \text{Barr}(b_n \bar{\gamma}) \right\},$$

our claim in the Proposition is complete. $\blacksquare$

Notice that we work with the below approximation for the adjustment factor in Theorem 2:

$$\exp \left(- \frac{b_n^2}{2} (\bar{\gamma}_n^* - \bar{P} \bar{\beta}_E - (b_n)^{-1} \bar{q})^\top \bar{\Sigma}^{-1} (\bar{\gamma}_n^* - \bar{P} \bar{\beta}_E - (b_n)^{-1} \bar{q}) - \text{Barr}(b_n \bar{\gamma}_n^*) + \log J_\phi(b_n \bar{\gamma}_n^*; \mathcal{U}) \right),$$

motivated by Proposition 9 where

$$\bar{\gamma}_n^* = \operatorname{argmin} \frac{1}{2} (\bar{\gamma} - \bar{P} \bar{\beta}_E - (b_n)^{-1} \bar{q})^\top \bar{\Sigma}^{-1} (\bar{\gamma} - \bar{P} \bar{\beta}_E - (b_n)^{-1} \bar{q}) + (b_n)^{-2} \text{Barr}(b_n \bar{\gamma}). \quad (30)$$

Proof [Proof of Proposition 11.] Set $\sqrt{n} z_n = b_n \bar{z}$, and let $\bar{\gamma}_n^*$ be the optimizer defined in (30). Then, the surrogate selection-informed (log) likelihood assumes the form

$$\ell_{n,E}(z_n; \hat{\beta}_{n,E} \mid N_{n,E}) = (\sqrt{n} \hat{\beta}_{n,E})^\top \bar{\Theta}^{-1} \bar{R}(b_n \bar{z}) - b_n^2 C_n(\bar{z}). \quad (31)$$

In the above representation, $C_n(\bar{z})$ equals

$$\begin{aligned} & (\bar{\gamma}_n^*)^\top (\bar{\Sigma})^{-1} (\bar{P} \bar{z} + (b_n)^{-1} \bar{q}) - \frac{1}{2} (\bar{\gamma}_n^*)^\top (\bar{\Sigma})^{-1} \bar{\gamma}_n^* - (b_n)^{-2} \text{Barr}(b_n \bar{\gamma}_n^*) \\ & + (b_n)^{-2} \log J_\phi(b_n \bar{\gamma}_n^*; \mathcal{U}) + \frac{1}{2} \bar{z}^\top \bar{R}^\top (\bar{\Theta} + \bar{A}^\top \bar{\Omega}^{-1} \bar{A})^{-1} \bar{R} \bar{z} - (b_n)^{-1} \bar{z}^\top \bar{P}^\top \bar{\Sigma}^{-1} \bar{q} - \frac{1}{2} (b_n)^{-2} \bar{q}^\top \bar{\Sigma}^{-1} \bar{q}, \end{aligned}$$

which we derive after plugging in the associated (log) approximation.

Next, we define several constants: C_1 is the largest eigenvalue of $\bar{R}^\top \bar{\Theta}^{-1} \bar{R}$ and C_0 is the smallest eigenvalue of $\bar{R}^\top (\bar{\Theta} + \bar{A}^\top \bar{\Omega}^{-1} \bar{A})^{-1} \bar{R}$. Consistent with our parameterization, we denote $b_n \bar{\beta}_E^{\max} = \sqrt{n} \hat{\beta}_{n,E}^{\max}$. It follows then from a Taylor series expansion of $C_n(\bar{z})$ around $\bar{\beta}_E^{\max}$ that the difference of log-likelihoods

$$\ell_{n,E}(z_n; \hat{\beta}_{n,E} \mid N_{n,E}) - \ell_{n,E}(\hat{\beta}_{n,E}^{\max}; \hat{\beta}_{n,E} \mid N_{n,E})$$

equals

$$\begin{aligned}
 & \sqrt{n}(\widehat{\beta}_{n,E})^\top \bar{\Theta}^{-1} \bar{R} b_n(\bar{z} - \bar{\beta}_E^{\max}) - b_n^2 \{C_n(\bar{z}) - C_n(\bar{\beta}_E^{\max})\} \\
 &= b_n(\bar{z} - \bar{\beta}_E^{\max})^\top \bar{R}^\top \bar{\Theta}^{-1} \sqrt{n} \widehat{\beta}_{n,E} - b_n^2 (\bar{z} - \bar{\beta}_E^{\max})^\top \nabla C_n(\bar{\beta}_E^{\max}) \\
 &\quad - \frac{b_n^2}{2} (\bar{z} - \bar{\beta}_E^{\max})^\top \nabla^2 C_n(R(\bar{\beta}_E^{\max}; \bar{z})) (\bar{z} - \bar{\beta}_E^{\max}) \\
 &= -\frac{n}{2} (z_n - \widehat{\beta}_{n,E}^{\max})^\top \nabla^2 C_n(R(\bar{\beta}_E^{\max}; \bar{z})) (z_n - \widehat{\beta}_{n,E}^{\max}).
 \end{aligned}$$

By Lemma 14, there exists $N \in \mathbb{N}$ such that the contribution of the Jacobian term towards the Hessian $(\nabla^2 C_n(\cdot))$,

$$\sup_{\bar{z} \in \mathcal{C}} \|\nabla_{\bar{z}}^2 (b_n)^{-2} \log J_\phi(b_n \bar{\gamma}_n^*(\bar{z}); \mathcal{U})\|_{\text{op}},$$

is uniformly bounded in operator norm by ϵ_0 for all $n \geq N$. Together with the observation that

$$(\bar{\gamma}_n^*)^\top (\bar{\Sigma})^{-1} (\bar{P} \bar{z} + (b_n)^{-1} \bar{q}) - \frac{1}{2} (\bar{\gamma}_n^*)^\top (\bar{\Sigma})^{-1} \bar{\gamma}_n^* - (b_n)^{-2} \text{Barr}(b_n \bar{\gamma}_n^*)$$

is a convex conjugate of the function $\frac{1}{2} (\bar{\gamma})^\top (\bar{\Sigma})^{-1} \bar{\gamma} + (b_n)^{-2} \text{Barr}(b_n \bar{\gamma})$ evaluated at $\bar{\Sigma}^{-1} (\bar{P} \bar{z} + (b_n)^{-1} \bar{q})$, we conclude

$$(C_0 - \epsilon_0) \cdot I \preceq \nabla^2 C_n(\bar{z}) \preceq (C_1 + \epsilon_0) \cdot I$$

for all $\bar{z} \in \mathcal{C}$, where I is the identity matrix of appropriate dimensions. This directly leads to our claim in the Proposition. $\blacksquare$

Proof [Proof of Theorem 12.] Fix $0 < a < 1$ such that

$$4a^2 \cdot (C_1 + C_0/2) - (1-a)^2 \cdot C_0/2 < 0,$$

where C_0 and C_1 are defined in Proposition 11. This follows by noting that the quadratic expression on the left-hand side has a root between $(0, 1)$. Denoting $\mathcal{C} \cap \mathcal{B}^c(\beta_{n,E}, \delta_n) =$

$\mathcal{B}'^c(\beta_{n,E}, \delta_n)$, we observe that there exists N such that for all $n \geq N$ such that

$$\begin{aligned}
 & \mathbb{P}_{n,E} \left(\Pi_{n,E} \left(\mathcal{B}^c(\beta_{n,E}, \delta_n) \mid \widehat{\beta}_{n,E}; N_{n,E} \right) \leq \epsilon \right) \\
 &= \mathbb{P}_{n,E} \left(\int_{\mathcal{B}'^c(\beta_{n,E}, \delta_n)} \pi_E(z_n) \cdot \exp(\ell_{n,E}(z_n; \widehat{\beta}_{n,E} | N_{n,E})) \, dz_n \leq \epsilon \int \pi_E(z_n) \cdot \exp(\ell_{n,E}(z_n; \widehat{\beta}_{n,E} | N_{n,E})) \, dz_n \right) \\
 &\geq \mathbb{P}_{n,E} \left(\int_{\mathcal{B}'^c(\beta_{n,E}, \delta_n)} \pi_E(z_n) \cdot \exp(\ell_{n,E}(z_n; \widehat{\beta}_{n,E} | N_{n,E}) - \ell_{n,E}(\widehat{\beta}_{n,E}^{\max}; \widehat{\beta}_{n,E} | N_{n,E})) \, dz_n \right. \\
 &\quad \left. \leq \epsilon \int_{\mathcal{B}(\beta_{n,E}, a\delta_n)} \pi_E(z_n) \cdot \exp(\ell_{n,E}(z_n; \widehat{\beta}_{n,E} | N_{n,E}) - \ell_{n,E}(\widehat{\beta}_{n,E}^{\max}; \widehat{\beta}_{n,E} | N_{n,E})) \, dz_n \right) \\
 &\geq \mathbb{P}_{n,E} \left(\int_{\mathcal{B}'^c(\beta_{n,E}, \delta_n)} \pi_E(z_n) \cdot \exp(-nC_0 \cdot \|\widehat{\beta}_{n,E}^{\max} - z_n\|^2/4) \, dz_n \right. \\
 &\quad \left. \leq \epsilon \int_{\mathcal{B}(\beta_{n,E}, a\delta_n)} \pi_E(z_n) \cdot \exp(-n(C_1 + C_0/2) \cdot \|\widehat{\beta}_{n,E}^{\max} - z_n\|^2/2) \, dz_n \right).
 \end{aligned}$$

The ultimate display follows by using the bounds in Proposition 11 where we set $\epsilon_0 = C_0/2$. This yields us the bound

$$\begin{aligned}
 & \mathbb{P}_{n,E} \left(\Pi_{n,E} \left(\mathcal{B}^c(\beta_{n,E}, \delta_n) \mid \widehat{\beta}_{n,E}; N_{n,E} \right) \leq \epsilon \right) \\
 &\geq \mathbb{P}_{n,E} \left(\int_{\mathcal{B}'^c(\beta_{n,E}, \delta_n)} \pi_E(z_n) \cdot \exp(-nC_0 \cdot \|\widehat{\beta}_{n,E}^{\max} - z_n\|^2/4) \, dz_n \right. \\
 &\quad \leq \epsilon \int_{\mathcal{B}(\beta_{n,E}, a\delta_n)} \pi_E(z_n) \cdot \exp(-n(C_1 + C_0/2) \cdot \|\widehat{\beta}_{n,E}^{\max} - z_n\|^2/2) \, dz_n, \\
 &\quad \|\widehat{\beta}_{n,E}^{\max} - z_n\| \geq (1-a)\delta_n \text{ for all } z_n \in \mathcal{B}'^c(\beta_{n,E}, \delta_n), \\
 &\quad \left. \|\widehat{\beta}_{n,E}^{\max} - z_n\| \leq 2a\delta_n \text{ for all } z_n \in \mathcal{B}(\beta_{n,E}, a\delta_n) \right) \\
 &\geq \mathbb{P}_{n,E} \left(\exp(-C_0 \cdot (1-a)^2 n \delta_n^2/4) \leq \epsilon \cdot \Pi_{n,E}(\mathcal{B}(\beta_{n,E}, a\delta_n)) \exp(-(C_1 + C_0/2) \cdot 4a^2 \delta_n^2/2) \right. \\
 &\quad \|\widehat{\beta}_{n,E}^{\max} - z_n\| \geq (1-a)\delta_n \text{ for all } z_n \in \mathcal{B}'^c(\beta_{n,E}, \delta_n), \\
 &\quad \left. \|\widehat{\beta}_{n,E}^{\max} - z_n\| \leq 2a\delta_n \text{ for all } z_n \in \mathcal{B}(\beta_{n,E}, a\delta_n) \right) \\
 &\geq \mathbb{P}_{n,E}(\|\widehat{\beta}_{n,E}^{\max} - \beta_{n,E}\| \leq a\delta_n).
 \end{aligned}$$

The argument in the last display follows from our assumptions on the selection-informed prior for sufficiently large n , coupled with the choice of $a \in (0, 1)$. We complete our proof by showing

$$\lim_{n \rightarrow \infty} \mathbb{P}_{n,E}(\|\widehat{\beta}_{n,E}^{\max} - \beta_{n,E}\| > a\delta_n) = 0.$$

To this end, we note that the MLE estimating equation is given by

$$\sqrt{n}\bar{R}^\top\bar{\Theta}^{-1}\hat{\beta}_{n,E} = b_n\nabla C_n(\bar{\beta}_E^{\max})$$

from the surrogate selection-informed (log) likelihood in (31) (Proposition 11) under the assumed parameters. Further, observing that $C_n(\cdot)$ is strongly convex for sufficiently large n , we have

$$(L)^{-2}\|\sqrt{n}\Sigma_E^{-1}\hat{\beta}_{n,E} - b_n\nabla C_n(\bar{\beta}_E)\|^2 \geq \|\sqrt{n}\hat{\beta}_{n,E}^{\max} - \sqrt{n}\beta_{n,E}\|^2;$$

$L = C_0/2$. Denoting the exact counterpart of $C_n(\cdot)$ (obtained upon using the exact probability of selection) by $\bar{C}_n(\cdot)$, we conclude

$$\begin{aligned} \mathbb{P}_{n,E}((b_n)^{-1}\sqrt{n}\|\hat{\beta}_{n,E}^{\max} - \beta_{n,E}\| > a\delta) &\leq (b_na\delta L)^{-2} \cdot \mathbb{E}_{n,E}(\|\sqrt{n}\bar{R}^\top\bar{\Theta}^{-1}\hat{\beta}_{n,E} - b_n\nabla C_n(\bar{\beta}_E)\|^2) \\ &\leq (b_na\delta L)^{-2} \cdot \mathbb{E}_{n,E}(\|\sqrt{n}\bar{R}^\top\bar{\Theta}^{-1}\hat{\beta}_{n,E} - b_n\nabla\bar{C}_n(\bar{\beta}_E)\|^2) \\ &\quad + (a\delta L)^{-2} \cdot \|\nabla C_n(\bar{\beta}_E) - \nabla\bar{C}_n(\bar{\beta}_E)\|^2. \end{aligned}$$

The first term in the final display clearly converges to 0 as $n \rightarrow \infty$. The second term converges to 0, using the result in Proposition 9 combined with the convexity and smoothness of the sequence $C_n(\cdot)$ for large enough n . $\blacksquare$

Appendix C. Supporting Theory (Section 6)

Below, we prove a result on the asymptotic orders of the gradient and Hessian of the (log) Jacobian; this in turn allows us to bound the contribution of the Jacobian term in the Hessian of the (log) likelihood in Proposition 11.

Lemma 13 *For $\eta > 0$, denote*

$$\mathcal{K}_\eta = \{x \in \mathbb{R}^{|E|} : \min_j x_j > \eta\}. \quad (32)$$

We then have the following uniform bounds on the derivatives of the (log)Jacobian:

$$\begin{aligned} \sup_{x \in \mathcal{K}_\eta} \|\nabla_\gamma \log J_\phi(\gamma; \mathcal{U})|_{b_n x}\|_\infty &= O(b_n^{-1}), \\ \sup_{x \in \mathcal{K}_\eta} \|\nabla_\gamma^2 \log J_\phi(\gamma; \mathcal{U})|_{b_n x}\|_{op} &= O(b_n^{-2}). \end{aligned}$$

Proof We begin by deriving an expression for the Hessian $\nabla_\gamma^2 \log J_\phi(\gamma; \mathcal{U})$. Recall from Theorem 5, for $g = 1, \dots, |E|$,

$$\frac{\partial}{\partial \gamma_g} \log J_\phi(\gamma; \mathcal{U}) = \sum_{i \in M_g} [(\Gamma + C)^{-1}]_{ii}, \quad (33)$$

where for simplicity we denote $C = \bar{U}^\top(X_E^\top X_E)^{-1}\Lambda\bar{U}$, $\Gamma = \text{diag}((\gamma_g I_{|g|-1})_{g \in \mathcal{G}_E})$, and M_g denotes the set of indices along the diagonal of Γ corresponding to group g . Using matrix

derivative identities similar to those applied in the proof of Theorem 5, we derive for $g, h = 1, \dots, |E|$,

$$\frac{\partial^2}{\partial \gamma_g \partial \gamma_h} \log J_\phi(\gamma; \mathcal{U}) = - \sum_{i \in M_g} \sum_{j \in M_h} [(\Gamma + C)^{-1}]_{ij} \cdot [(\Gamma + C)^{-1}]^\top_{ij}. \quad (34)$$

Both our claims rely on a uniform bound on the entries of $(\Gamma + C)^{-1}$. Let the operators $s_{\max}$, $\lambda_{\max}$, and $\lambda_{\min}$ denote the largest singular value, largest eigenvalue, and smallest eigenvalue of a matrix, respectively. Fix $x \in \mathcal{K}_\eta$; let $\Gamma(b_n x) = \text{diag}((b_n x_g I_{|g|-1})_{g \in \mathcal{G}_E})$. Denote its dimension by $q = \sum_{g \in E} (|g| - 1)$. Then

$$\begin{aligned} \max_{1 \leq i, j \leq q} |[(\Gamma(b_n x) + C)^{-1}]_{ij}| &\leq s_{\max}((\Gamma(b_n x) + C)^{-1}) \\ &= \lambda_{\max}^{1/2}((\Gamma(b_n x) + C)^{-1} [(\Gamma(b_n x) + C)^{-1}]^\top) \\ &= \lambda_{\min}^{-1/2}((\Gamma(b_n x) + C)^\top (\Gamma(b_n x) + C)). \end{aligned}$$

Now note the following:

$$\begin{aligned} \lambda_{\min}((\Gamma(b_n x) + C)^\top (\Gamma(b_n x) + C)) &= \inf_{\|v\|_2=1} v^\top ((\Gamma(b_n x) + C)^\top (\Gamma(b_n x) + C)) v \\ &\geq b_n^2 \min_{1 \leq j \leq |E|} x_j^2 + \lambda_{\min}(C^\top C) \\ &\geq (b_n \eta)^2. \end{aligned}$$

Combining the previous two displays uniformly over $\mathcal{K}_\eta$, we have that

$$\sup_{x \in \mathcal{K}_\eta} \left(\max_{1 \leq i, j \leq q} |[(\Gamma(b_n x) + C)^{-1}]_{ij}| \right) \leq \frac{1}{b_n \eta} = O(b_n^{-1}). \quad (35)$$

By (33), each entry of the gradient is the sum of up to $p - 1$ entries of $(\Gamma(b_n x) + C)^{-1}$, and by (34), each entry of the Hessian is a sum of up to p^2 products of entries of the same matrix. Lastly, a bound on the ℓ_∞ norm of the gradient follows directly from this element-wise bound. Further, a bound for the operator norm of the Hessian follows after noting that for an $r \times r$ square matrix M , $\|M\|_2 \leq r \max_{ij} |M|_{ij}$. $\blacksquare$

With the previous result in hand, we prove the next Lemma used in Proposition 11.

Lemma 14 *Under the assumptions of Proposition 11, we have*

$$\lim_{n \rightarrow \infty} \left\{ \sup_{\bar{z} \in \mathcal{C}} \|\nabla_{\bar{z}}^2(b_n)^{-2} \log J_\phi(b_n \bar{\gamma}_n^*(\bar{z}); \mathcal{U})\|_{op} \right\} = 0.$$

Proof To proceed with the proof, we derive an expression for the Hessian of the (log) Jacobian. Recall, $\bar{\gamma}_n^*(\bar{z})$ is the optimizer of the convex conjugate of the function

$$(\bar{\gamma})^\top (\bar{\Sigma})^{-1} \bar{\gamma} / 2 + (b_n)^{-2} \text{Barr}(b_n \bar{\gamma})$$

evaluated at $\bar{\Sigma}^{-1}(\bar{P}\bar{z} + (b_n)^{-1}\bar{q})$. By properties of the convex conjugate function and the chain rule,

$$\nabla_{\bar{z}}\bar{\gamma}_n^*(\bar{z}) = \bar{P}^\top \bar{\Sigma}^{-1} \mathcal{H}_n(\bar{z}),$$

where $\mathcal{H}_n(\bar{z})$ denotes the inverse Hessian matrix $(\bar{\Sigma}^{-1} + \nabla^2 \text{Barr}(b_n \bar{\gamma}_n^*(\bar{z})))^{-1}$. Here, we note that the (diagonal) Hessian of the barrier function is positive definite for all $\bar{z}$, which implies $\|\mathcal{H}_n(\bar{z})\|_2 \leq \|\bar{\Sigma}\|_2$ uniformly over all n and $\bar{z} \in \mathcal{C}$. That is,

$$\nabla_{\bar{z}}(b_n)^{-2} \log J_\phi(b_n \bar{\gamma}_n^*(\bar{z}); \mathcal{U}) = (b_n)^{-1} \bar{P}^\top \bar{\Sigma}^{-1} \mathcal{H}_n(\bar{z}) \left(\nabla_x \log J_\phi(x; \mathcal{U}) \Big|_{b_n \bar{\gamma}_n^*(\bar{z})} \right).$$

To compute the Hessian, we will use the following identity

$$\frac{\partial A(x)b(x)}{\partial x} = \frac{\partial A(x)}{\partial x} \times_2 b(x)^\top + A(x) \frac{\partial b(x)}{\partial x}, \quad (36)$$

where $A(x)b(x)$ denotes a matrix-vector product; the 3-dimensional tensor $\frac{\partial A(x)}{\partial x}$ is summed across its second dimension in the first term of (36). Observe that the element-wise derivatives of $\mathcal{H}_n$ with respect to the entries of $\bar{\gamma}_n^*(\bar{z})$ are given by

$$\frac{\partial}{\partial (\bar{\gamma}_n^*)_\ell} [\mathcal{H}_n(\bar{z})]_{ij} = -b_n \nabla_{\ell\ell}^3 \text{Barr}(b_n \bar{\gamma}_n^*(\bar{z})) [\mathcal{H}_n(\bar{z})]_{i\ell} [\mathcal{H}_n(\bar{z})]_{\ell j}, \quad (37)$$

involving the third derivatives of the barrier function. Then, plugging into the first term of (36), we get

$$\begin{aligned} & -b_n \sum_j \nabla_{\ell\ell}^3 \text{Barr}(b_n \bar{\gamma}_n^*(\bar{z})) [\mathcal{H}_n(\bar{z})]_{i\ell} [\mathcal{H}_n(\bar{z})]_{\ell j} [\nabla J_\phi(b_n \bar{\gamma}_n^*(\bar{z}); \mathcal{U})]_j \\ & = -b_n \nabla_{\ell\ell}^3 \text{Barr}(b_n \bar{\gamma}_n^*(\bar{z})) [\mathcal{H}_n(\bar{z})]_{i\ell} \sum_j \left([\mathcal{H}_n(\bar{z})]_{\ell j} [\nabla J_\phi(b_n \bar{\gamma}_n^*(\bar{z}); \mathcal{U})]_j \right). \end{aligned}$$

Elementwise (row i and column ℓ), this matrix has the same entries as $\mathcal{H}_n(\bar{z})$, but with the ℓ th column scaled by

$$-b_n \nabla_{\ell\ell}^3 \text{Barr}(b_n \bar{\gamma}_n^*(\bar{z})) [\mathcal{H}_n(\bar{z})]_{i\ell} [\nabla J_\phi(b_n \bar{\gamma}_n^*(\bar{z}); \mathcal{U})]_\ell.$$

Let $b_n \mathcal{D}_n(\bar{z})$ be a diagonal matrix with these entries on its main diagonal. Then the matrix in the first term of (36) is given by $\mathcal{H}_n(\bar{z}) \mathcal{D}_n(\bar{z})$. The second term of (36) equals

$$b_n \mathcal{H}_n(\bar{z}) \left(\nabla_x^2 \log J_\phi(x; \mathcal{U}) \Big|_{b_n \bar{\gamma}_n^*(\bar{z})} \right).$$

Thus, $\nabla_{\bar{z}}(b_n)^{-2} \log J_\phi(b_n \bar{\gamma}_n^*(\bar{z}); \mathcal{U})$ equals

$$\bar{P}^\top \bar{\Sigma}^{-1} \mathcal{H}_n(\bar{z}) \left(\mathcal{D}_n(\bar{z}) + \nabla_x^2 \log J_\phi(x; \mathcal{U}) \Big|_{b_n \bar{\gamma}_n^*(\bar{z})} \right) \mathcal{H}_n(\bar{z}) \bar{\Sigma}^{-1} \bar{P}.$$

Noting $\bar{P}^\top \bar{\Sigma}^{-1} \mathcal{H}_n(\bar{z}) = O(1)$ uniformly over n and $\bar{z}$, bounding (C) in operator norm follows by uniformly bounding the largest diagonal element of $\mathcal{D}_n(\bar{z})$, and the operator norm of $\nabla_x^2 \log J_\phi(x; \mathcal{U}) \Big|_{b_n \bar{\gamma}_n^*(\bar{z})}$.

Bounds for both terms follow from an application of Lemma 13, along with the use of the observation that the third derivatives of the barrier function are decreasing. To complete the proof, it therefore suffices to show uniformly over $\bar{z} \in \mathcal{C}$, and for sufficiently large n , all the entries of $\bar{\gamma}_n^*(\bar{z})$ are bounded below by a constant η . To this end, we define the limit of $\bar{\gamma}_n^*(\bar{z})$ as $n \rightarrow \infty$:

$$\bar{\gamma}_\infty^*(\bar{z}) = \operatorname{argmin}_{\gamma > 0} \left\{ \frac{1}{2} \gamma^\top \bar{\Sigma}^{-1} \gamma - \gamma^\top \bar{\Sigma}^{-1} \bar{P} \bar{z} \right\}. \quad (38)$$

The image of a compact set $\mathcal{C}$ under the continuous map $\bar{\gamma}_\infty^*(\cdot)$ is compact and a subset of the positive orthant, which we call $\mathcal{C}'$. Define

$$\eta = \frac{1}{2} \min\{|x_j| : x \in \mathcal{C}'\} > 0.$$

Uniform convergence of $\bar{\gamma}_n^*(\bar{z})$ to $\bar{\gamma}_\infty^*(\bar{z})$ on a compact domain leads us to conclude

$$\min_j |[\bar{\gamma}_n^*(\bar{z})]_j| > \eta > 0,$$

for all $\bar{z} \in \mathcal{C}$ and sufficiently large n . ■

Appendix D. Supplementary Details (Section 7)

We outline additional details involving the parameters in our numerical experiments below. For the simulation instances we generate in the atomic and balanced case analyses, each active coefficient has a random sign with magnitude $\sqrt{2n^{-1}t \log p}$. We let $t = 0.2$ for the low SNR setting, $t = 0.5$ for the moderate SNR setting, and $t = 1.5$ for the high SNR setting. In the heterogeneous scenario, the first predictor in the smallest active group has magnitude $\sqrt{2n^{-1}t \log p / |T|}$ and the last predictor in the largest group has magnitude $\sqrt{2n^{-1}t \log p}$ with magnitudes linearly interpolated for intermediate active coefficients; T is the number of signal variables in the instance. Each coefficient assumes a random sign and we set t for the low, medium, and high SNR as in our previous cases.

For the selection step, we set the grouped penalty weights to be

$$\lambda_g = \lambda \rho \sigma \sqrt{2 \log p \frac{|g|}{\bar{g}}}$$

for solving the Group LASSO in both the randomized (“Selection-informed”) and non-randomized formulations (“Naive” and “Split”); $|g|$ is the number of features in group g , and $\bar{g}$ is the floor of the average group size. In the choice of the penalty weights, when solving the Group LASSO and overlapping Group LASSO, for “Split” we set $\rho = r$, the proportion of data used for the query, while we set $\rho = 1$ for “Selection-informed” and “Naive.” When solving the standardized Group LASSO, for “Split” we set $\rho = \sqrt{r}$ while for “Selection-informed” and “Naive” we again set $\rho = 1$. Clearly, in our balanced settings,

we impose a uniform penalty across all groups, while the penalties for the heterogeneous settings scale with the size of our groups.

Addressing selection-informed inference after the Group LASSO, the barrier function used in our optimization problem (see Theorem 5) is given by

$$\text{Barr}(\gamma) = \sum_{g \in \mathcal{G}_E} \log(1 + (\gamma_g)^{-1}).$$

Observe that this choice of penalty assigns higher preference to optimizing variables away from the boundary of the selection region $[0, \infty)^{\mathbb{R}^{|\mathcal{G}_E|}}$. For executing the sampler, we set the initial draw as

$$\beta^{(0)} = \hat{\beta}_E,$$

the refitted least squares estimate in our setup. Completing our specifications, the inferential results we report in Section 7 are based upon 1500 draws of the Langevin sampler. We discard the first 100 samples as burn-in and retain the remainder for uncertainty estimation. The code for our experiments in the paper is available here: https://github.com/snigdhagit/selective-inference/tree/group_LASSO/selection/randomized.

D.1 Supplementary Details for HCP Analysis

The “preprocessed” version of the data set used in our analysis, which had undergone the processing stream described in Glasser et al. (2013), was downloaded from the HCP’s ConnectomeDB platform (Marcus et al., 2011). The fMRI data comprises time courses at many “voxels” throughout the brain that are typically each a few millimeters cubed in volume. This data was preprocessed as described in Sripada et al. (2019), excluding the steps that are specific to resting state processing. While the HCP data includes a variety of imaging modalities, we use both behavioral and functional magnetic resonance imaging (fMRI) measurements recorded from a cognitive task, namely the “N-back” task (Barch et al., 2013).

In the the “N-back” task, participants are presented with a sequence of pictures about which they make judgments, and their accuracy and brain activity is recorded while they perform the task. There are two different conditions of principle interest, each of which are presented in blocks. In the 0-back condition, participants simply judge whether each item is the same as the item presented at the beginning of the block. In the 2-back condition, participants judge whether each item is the same as the item presented two trials previous. As may be intuitively clear, the 2-back condition is appreciably more demanding with respect to working memory. A common approach for analyzing fMRI data involves the construction of “contrasts.” Measuring activity during the 2-back condition would likely indicate activity related to working memory, but it would also include activity indicating many other phenomena such as visual processing, motor activation in order to press buttons to indicate judgments, etc. These phenomena are not of primary interest, so we consider a contrast formed by subtracting the activation during the 0-back condition from the activation during the 2-back condition. This 2-back minus 0-back contrast is standard for the N-back task (Barch et al., 2013). Contrasts were obtained using in-house processing scripts that use SPM12. The standardized accuracy of each participant during this task will be our

target of prediction y and we will use the contrast as the predictor X . As a preprocessing step, columns of the design matrix X are adjusted to have mean 0 and unit norm.

Using the contrast value from each voxel results in very high dimensional data, and analysis is sometimes instead performed at the level of “regions of interest” (ROIs). This provides a means of effectively downsampling the data by aggregating information at each ROI, which is a spatially contiguous group of voxels. These ROIs can be defined *a priori* according to one of a variety of atlases, and this aids interpretability and enables comparisons of findings across studies that use the same atlas. We use ROIs as defined by the “Power Parcellation” (Power et al., 2011). In addition to being a broadly popular atlas, the Power Parcellation is also noteworthy in that it assigns each of its 264 ROIs to a “brain system.” The spatial coordinates of the ROIs, as well as their assignment to brain systems, are described in Power et al. (2011). The MarsBar utility (Brett et al., 2002) was used to extract contrast values for each of these ROIs. Of the 264 ROIs, 236 are assigned to one of 13 distinct, named brain systems while the remainder are simply labeled “unknown” and in our analysis we use only these 236 positively labeled ROIs as predictors in our regression. Because each of these brain systems is putatively believed to underlie a discrete set of functions (e.g., because they typically coactivate for a given type of task), we partition our predictors into groups by brain system label, and then use the Group LASSO to predict accuracy on the N-back task using data from these 236 ROIs. While inference may be performed at the level of individual ROIs, it is also useful to interrogate effects at a system-wide level. Further averaging all of the ROIs within a single system may be too coarse and obscure useful signal, so the Group LASSO provides a means of allowing each ROI to make a distinct predictive contribution while still performing selection at the interpretable level of entire brain systems. Because in this application $n > p$, we estimate $\hat{\sigma}^2 = (n - p)^{-1} \left\| y - X(X^\top X)^{-1} X^\top y \right\|_2^2$. We use the same value for $\hat{\sigma}^2$ for the intervals obtained via data splitting. We set λ_g for each group as described in D and set the randomization level τ to satisfy (23) at varying levels of r . Choosing $\lambda = 1$ (as we did for the simulation studies) yields a fully dense model, so we increase to $\lambda = 10$ which selects just a single group.

D.2 Supplementary Numerical Comparison

We conduct an additional numerical experiment to compare the methods of Yang et al. (2016) and also Loftus and Taylor (2015). Specifically, we consider one instance of the simulation settings considered in Yang et al. (2016) where we draw $X \in \mathbb{R}^{500 \times 500}$ with entries independently and identically distributed as $\mathcal{N}(0, \frac{1}{500})$. The $p = 500$ features are arranged into 50 contiguous groups of 10 features each. The first 10 groups (i.e., first 50 features) are all active with associated coefficient 1.5 and the remainder are inactive with associated coefficients 0, i.e., $\beta = (1.5 \cdot \mathbf{1}_{50}^\top \quad \mathbf{0}_{450}^\top)^\top$. The response y is then generated as $\mathcal{N}(\mu, 1)$, where $\mu = X\beta$. We generate a single realization of the data in this setting and then apply the methods of Yang et al. (2016), Loftus and Taylor (2015), and our method conducted with 5000 posterior samples (with 100 samples discarded as burn-in). Findings in this instance gives us an opportunity to note the extent of agreement between the three methods.

For all methods, the first stage is automatically selecting groups using: (i) the Group LASSO (for Yang et al. (2016)), (ii) the randomized Group LASSO (for our method) in (3), or (iii) forward stage-wise selection (for Loftus and Taylor (2015)). We use $\lambda = 4$ for the approach of Yang et al. (2016) which yields the selection of 11 active groups (i.e., 110 features), and we then tune parameters for the other two methods to select the same number of active features. Once the model has been selected, we proceed to inference with $\alpha = 0.1$.

For the methods by Yang et al. (2016) and Loftus and Taylor (2015), the inferential target for each selected group g in E is an overall group effect μ_g , which we review in more detail under Section 2.2. We apply our methods to construct credible intervals for the individual components of the coefficient vector for each group; inference for individual effects in the selected groups is not addressed by the previous two methods. As described in Section 7, sampling from the selection-informed posterior with a diffuse (non-informative) prior yields credible intervals for the individual effects with “good” frequentist properties. Note, we do not pursue inference for the overall group effect—a (non-linear) function of the selection-informed parameters β_E —using our Bayesian methods. This is because our focus is on the extent of agreement between all three methods in terms of their frequentist properties. Specially, “good frequentist properties” for μ_g will also depend on a choice of prior for this parameter; in this case, a non-informative prior for β_E might not be non-informative for μ_g for $g \in \mathcal{G}_E$.

We summarize our results in Table 1. Each of the three methods selects all of the 5 active groups and 6 additional inactive groups, although the identities of the selected inactive groups differ slightly across the methods due to differences in the query (i.e., Group LASSO vs randomized Group LASSO vs forward stage-wise selection). The method of Loftus and Taylor (2015) yields no significant p-values at $\alpha = 0.10$: it makes no Type I errors, but 5 Type II errors. The method of Yang et al. (2016) correctly rejects the null for 4 of 5 active groups and only makes a Type I error for 1 of 6 six inactive groups. For our method, we report the component-wise coverage of our marginal credible intervals in each group. Empirically, we appear to have coverage that does not appreciably deviate from nominal (i.e., 90%). Coverage for active coefficients is slightly better at 88% as opposed to coverage for inactive coefficients at 80%, although these may just be chance fluctuations. In summary, the test by Loftus and Taylor (2015) seems more conservative than the remaining two methods. Coverage for the individual variable effects in each active group by our method seem to be consistent with the lower bounds for the overall group effect by Yang et al. (2016).

Group #	μ_g	Loftus p -value	Yang LCB	Yang p -value	Ours (Coverage)
1	4.49	0.34	2.39	0.00	1.0
2	4.01	0.14	1.11	0.03	0.9
3	4.19	0.42	1.01	0.03	0.9
4	4.07	0.32	-2.68	0.33	0.8
5	4.20	0.23	3.05	0.00	0.8
8	0	0.55			
11	0				0.7
14	0				0.7
17	0	0.77			
18	0		-9.99	0.81	
20	0	0.54	-0.95	0.21	0.9
28	0		-9.54	0.66	1.0
33	0	0.78	1.74	0.02	
36	0	0.83	-4.08	0.45	0.7
39	0		-3.05	0.50	0.8
46	0	0.83			

Table 1: Comparison of inferential results on single realization of synthetic data using methods of Loftus and Taylor (2015), Yang et al. (2016), and the proposed method. Blanks indicate that the associated method did not select the variable group depicted in the corresponding row. LCB signifies lower confidence bound. Coverage (where applicable) was assessed using 90% credible intervals. “Ours” refers to the Selection-informed method discussed in the manuscript.

References

- Alfred Auslender. Penalty and barrier methods: a unified framework. *SIAM Journal on Optimization*, 10(1):211–230, 1999.
- Deanna M. Barch, Gregory C. Burgess, Michael P. Harms, Steven E. Petersen, Bradley L. Schlaggar, Maurizio Corbetta, Matthew F. Glasser, Sandra Curtiss, Sachin Dixit, Cindy Feldt, Dan Nolan, Edward Bryant, Tucker Hartley, Owen Footer, James M. Bjork, Russ Poldrack, Steve Smith, Heidi Johansen-Berg, Abraham Z. Snyder, and David C. Van Essen. Function in the human connectome: Task-fMRI and individual differences in behavior. *NeuroImage*, 80:169–189, October 2013. ISSN 1053-8119. doi: 10.1016/j.neuroimage.2013.05.033.
- Yoav Benjamini. Selective inference: The silent killer of replicability. *Harvard Data Science Review*, 2(4), 2020.
- Richard Berk, Lawrence Brown, Andreas Buja, Kai Zhang, Linda Zhao, et al. Valid post-selection inference. *The Annals of Statistics*, 41(2):802–837, 2013.
- Matthew Brett, Jean-Luc Anton, Romain Valabregue, and Jean-Baptiste Poline. Region of interest analysis using an SPM toolbox. In *Presented at the 8th International Con-*

- ference on Functional Mapping of the Human Brain*, June 2002. Abstract Available in *NeuroImage*, Vol 16, No 2.
- A De Acosta. Moderate deviations and associated laplace approximations for sums of independent random vectors. *Transactions of the American Mathematical Society*, 329(1):357–375, 1992.
- William Fithian, Dennis Sun, and Jonathan Taylor. Optimal inference after model selection. *arXiv preprint arXiv:1410.2597*, 2014.
- Lucy L Gao, Jacob Bien, and Daniela Witten. Selective inference for hierarchical clustering. *arXiv preprint arXiv:2012.02936*, 2020.
- Matthew F. Glasser, Stamatios N. Sotiropoulos, J. Anthony Wilson, Timothy S. Coalson, Bruce Fischl, Jesper L. Andersson, Junqian Xu, Saad Jbabdi, Matthew Webster, Jonathan R. Polimeni, David C. Van Essen, and Mark Jenkinson. The minimal preprocessing pipelines for the Human Connectome Project. *NeuroImage*, 80:105–124, October 2013. ISSN 1053-8119. doi: 10.1016/j.neuroimage.2013.04.127.
- Tadeusz Inglot and Piotr Majerski. Simple upper and lower bounds for the multivariate laplace approximation. *Journal of Approximation Theory*, 186:1–11, 2014.
- Laurent Jacob, Guillaume Obozinski, and Jean-Philippe Vert. Group Lasso with Overlap and Graph Lasso. In *Proceedings of the 26th Annual International Conference on Machine Learning*, ICML ’09, pages 433–440, New York, NY, USA, 2009. ACM. ISBN 978-1-60558-516-1. doi: 10.1145/1553374.1553431.
- Jason D. Lee, Dennis L. Sun, Yuekai Sun, and Jonathan E. Taylor. Exact post-selection inference with the LASSO. *The Annals of Statistics*, 44(3):907–927, November 2016.
- Keli Liu, Jelena Markovic, and Robert Tibshirani. More powerful post-selection inference, with application to the LASSO. *arXiv preprint arXiv:1801.09037*, 2018.
- Joshua R Loftus and Jonathan E Taylor. Selective inference in regression models with groups of variables. *arXiv preprint arXiv:1511.01478*, 2015.
- Daniel Marcus, John Harwell, Timothy Olsen, Michael Hodge, Matthew Glasser, Fred Prior, Mark Jenkinson, Timothy Laumann, Sandra Curtiss, and David Van Essen. Informatics and data mining tools and strategies for the Human Connectome Project. *Frontiers in Neuroinformatics*, 5, 2011. ISSN 1662-5196. doi: 10.3389/fninf.2011.00004.
- Snigdha Panigrahi. Carving model-free inference. *arXiv preprint arXiv:1811.03142*, 2018.
- Snigdha Panigrahi and Jonathan Taylor. Scalable methods for Bayesian selective inference. *Electronic Journal of Statistics*, 12(2):2355–2400, 2018.
- Snigdha Panigrahi and Jonathan Taylor. Approximate selective inference via maximum likelihood. *Journal of the American Statistical Association*, pages 1–11, 2022.
- Snigdha Panigrahi, Jonathan Taylor, and Asaf Weinstein. Integrative methods for post-selection inference under convex constraints. *Annals of Statistics*, 49(5):2803–2824, 2021.

- Snigdha Panigrahi, Shariq Mohammed, Arvind Rao, and Veerabhadran Baladandayuthapani. Integrative Bayesian models using post-selective inference: A case study in radio-genomics. *Biometrics*, 2022.
- Kaare B Petersen and Michael S Pedersen. *The Matrix Cookbook*. Technical University of Denmark, 2012.
- Jonathan D. Power, Alexander L. Cohen, Steven M. Nelson, Gagan S. Wig, Kelly Anne Barnes, Jessica A. Church, Alecia C. Vogel, Timothy O. Laumann, Fran M. Miezín, Bradley L. Schlaggar, and Steven E. Petersen. Functional network organization of the human brain. *Neuron*, 72(4):665–678, November 2011. ISSN 0896-6273. doi: 10.1016/j.neuron.2011.09.006.
- Ralph Tyrell Rockafellar. *Convex Analysis*. Princeton University Press, 2015.
- Christoph Schultheiss, Claude Renaux, and Peter Bühlmann. Multicarving for high-dimensional post-selection inference. *Electronic Journal of Statistics*, 15(1):1695–1742, 2021.
- Xiaocheng Shang, Zhanxing Zhu, Benedict Leimkuhler, and Amos J Storkey. Covariance-controlled adaptive Langevin thermostat for large-scale Bayesian sampling. *Advances in Neural Information Processing Systems*, 28:37–45, 2015.
- Noah Simon and Robert Tibshirani. Standardization and the group LASSO penalty. *Statistica Sinica*, 22(3):983, 2012.
- Noah Simon, Jerome Friedman, Trevor Hastie, and Robert Tibshirani. A sparse-group LASSO. *Journal of Computational and Graphical Statistics*, 22(2):231–245, 2013.
- Chandra Sripada, Mike Angstadt, Saige Rutherford, Daniel Kessler, Yura Kim, Mike Yee, and Elizaveta Levina. Basic units of inter-individual variation in resting state connectomes. *Scientific Reports*, 9(1):1900, February 2019. ISSN 2045-2322. doi: 10.1038/s41598-018-38406-5.
- Chandra Sripada, Mike Angstadt, Saige Rutherford, Aman Taxali, and Kerby Shedden. Toward a “treadmill test” for cognition: Improved prediction of general cognitive ability from the task activated brain. *Human Brain Mapping*, 41(12):3186–3197, 2020. ISSN 1097-0193. doi: 10.1002/hbm.25007.
- Shinya Suzumura, Kazuya Nakagawa, Yuta Umezū, Koji Tsuda, and Ichiro Takeuchi. Selective inference for sparse high-order interaction models. In *International Conference on Machine Learning*, pages 3338–3347. PMLR, 2017.
- Kosuke Tanizaki, Noriaki Hashimoto, Yu Inatsu, Hidekata Hontani, and Ichiro Takeuchi. Computing valid p-values for image segmentation by selective inference. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9553–9562, 2020.
- Xiaoying Tian and Jonathan Taylor. Selective inference with a randomized response. *The Annals of Statistics*, 46(2):679–710, 2018.

- Xiaoying Tian, Snigdha Panigrahi, Jelena Markovic, Nan Bi, and Jonathan Taylor. Selective sampling after solving a convex problem. *arXiv preprint arXiv:1609.05609*, 2016.
- David C. Van Essen, Stephen M. Smith, Deanna M. Barch, Timothy E. J. Behrens, Essa Yacoub, Kamil Ugurbil, and the WU-Minn HCP Consortium. The WU-Minn Human Connectome Project: An overview. *NeuroImage*, 80:62–79, October 2013. ISSN 1053-8119. doi: 10.1016/j.neuroimage.2013.05.041.
- Roderick Wong. *Asymptotic Approximations of Integrals*. SIAM, 2001.
- Fan Yang, Rina Foygel Barber, Prateek Jain, and John Lafferty. Selective inference for group-sparse linear models. In *Advances in Neural Information Processing Systems*, pages 2469–2477, 2016.
- Daniel Yekutieli. Adjusted Bayesian inference for selected parameters. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 74(3):515–541, 2012.
- Ming Yuan and Yi Lin. Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society: Series B*, 68(1):49–67, 2005.
- Qingyuan Zhao and Snigdha Panigrahi. Selective inference for effect modification: An empirical investigation. *Observational Studies*, 5(2):131–140, 2019.