

Supplementary Material for “Bayesian Group Regularization in Generalized Linear Models with a Continuous Spike-and-Slab Prior”

Ray Bai

Department of Statistics, George Mason University

September 21, 2025

Abstract

Section A and Section B give the proofs for Section 3 and Section 4 of the main manuscript respectively. Section C gives additional details for the expectation-maximization (EM) algorithm introduced in Section 5.1 of the main manuscript. Section D gives the complete Gibbs sampler for the MCMC algorithms introduced in Section 5.3 of the main manuscript. Finally, Section E presents simulation studies for our model in negative binomial regression with a log link.

A Proofs for Section 3

Before proving Theorem 1, it will be necessary to define some terminology. Letting $\pi(\boldsymbol{\beta})$ denote the spike-and-slab group lasso (SSGL) prior (10), we can express the SSGL log prior as

$$\log \pi(\boldsymbol{\beta}) = \sum_{g=1}^G \text{pen}(\boldsymbol{\beta}_g),$$

where

$$\text{pen}(\boldsymbol{\beta}_g) = \log [(1 - \theta) \boldsymbol{\Psi}(\boldsymbol{\beta}_g \mid \lambda_0) + \theta \boldsymbol{\Psi}(\boldsymbol{\beta}_g \mid \lambda_1)]. \quad (\text{A.1})$$

Recall that

$$p_{\theta}^*(\boldsymbol{\beta}_g) = \frac{\theta \boldsymbol{\Psi}(\boldsymbol{\beta}_g \mid \lambda_1)}{\theta \boldsymbol{\Psi}(\boldsymbol{\beta}_g \mid \lambda_1) + (1 - \theta) \boldsymbol{\Psi}(\boldsymbol{\beta}_g \mid \lambda_0)}. \quad (\text{A.2})$$

The derivative of $\text{pen}(\boldsymbol{\beta}_g)$ with respect to $\|\boldsymbol{\beta}_g\|_2$ is

$$\frac{\partial \text{pen}(\boldsymbol{\beta}_g)}{\partial \|\boldsymbol{\beta}_g\|_2} = -\lambda_{\theta}^*(\boldsymbol{\beta}_g), \quad (\text{A.3})$$

where

$$\lambda_{\theta}^*(\boldsymbol{\beta}_g) = \lambda_1 p_{\theta}^*(\boldsymbol{\beta}_g) + \lambda_0 [1 - p_{\theta}^*(\boldsymbol{\beta}_g)]. \quad (\text{A.4})$$

Under the objective function (13), it can then be seen from (5) and (A.3) that any local maximizer $\hat{\beta}$ of (13) must satisfy

$$\mathbf{X}_g^\top \text{diag}\{\xi'(\mathbf{X}\hat{\beta})\} \left(\mathbf{Y} - b'(\xi(\mathbf{X}\hat{\beta})) \right) - \lambda_\theta^*(\hat{\beta}_g) \partial \|\beta_g\|_2 = \mathbf{0}_{m_g} \text{ for all } g \in \{1, \dots, G\}, \quad (\text{A.5})$$

where ∂f denotes the subdifferential of f . From (A.5), we can obtain the necessary first-order Karush-Kuhn-Tucker (KKT) conditions for $\hat{\beta}$ to be a local maximizer of (13),

$$\begin{aligned} \mathbf{X}_g^\top \text{diag}\{\xi'(\mathbf{X}\hat{\beta})\} (\mathbf{Y} - b'(\xi(\mathbf{X}\hat{\beta}))) - \lambda_\theta^*(\hat{\beta}_g) \frac{\hat{\beta}_g}{\|\hat{\beta}_g\|_2} &= \mathbf{0}_{m_g}, & \text{for } \hat{\beta}_g \neq \mathbf{0}_{m_g}, \\ \|\mathbf{X}_g^\top \text{diag}\{\xi'(\mathbf{X}\hat{\beta})\} (\mathbf{Y} - b'(\xi(\mathbf{X}\hat{\beta})))\|_2 &\leq \lambda_\theta^*(\mathbf{0}_{m_g}) & \text{for } \hat{\beta}_g = \mathbf{0}_{m_g}. \end{aligned} \quad (\text{A.6})$$

Proof of Theorem 1 Recall that $\beta_0 = (\beta_{0S_0}^\top, \mathbf{0}^\top)^\top$. To prove Theorem 1, we do the following two steps. First, we consider the subspace of vectors $\mathcal{T} = \{\beta \in \mathbb{R}^p : \beta_{S_0^c} = \mathbf{0}\}$. We show that on $\mathcal{T}$, there exists a local maximizer $\hat{\beta} = (\hat{\beta}_{S_0}^\top, \mathbf{0}^\top)^\top$ of (13) under the SSGL prior (10) such that $\|\hat{\beta} - \beta_0\|_2 = O_p(\epsilon_n)$. Next, we show that $\hat{\beta}$ is a strict local maximizer of (13) under the SSGL prior (10) on the *entire* parameter space $\mathbb{R}^p$.

Step 1 (local consistency). First we solve the constrained penalized likelihood estimation problem on the subspace $\mathcal{T} = \{\beta \in \mathbb{R}^p : \beta_{S_0^c} = \mathbf{0}\}$. Then on $\mathcal{T}$, any estimator of β_0 must be of the form $\hat{\beta} = (\hat{\beta}_{S_0}^\top, \mathbf{0}^\top)^\top$. We first show the existence of such a $\hat{\beta}$ to the optimization problem,

$$\hat{\beta} = \arg \max_{\beta} Q_1(\beta) = \arg \max_{\beta} \left\{ \ell_n(\beta) + \sum_{g=1}^{s_0} \text{pen}(\beta_g), \beta_{S_0^c} = \mathbf{0} \right\}, \quad (\text{A.7})$$

where $\text{pen}(\beta_g)$ is defined as in (A.1) and $\hat{\beta}$ satisfies

$$\|\hat{\beta} - \beta_0\|_2 = O_p(\epsilon_n), \quad (\text{A.8})$$

with $\epsilon_n = (s_0 \log G/n)^{1/2}$ and $\beta_0 = (\beta_{0S_0}^\top, \mathbf{0}^\top)^\top$. Let $\Delta = (\Delta_{S_0}^\top, \mathbf{0}^\top)^\top$, where $\|\Delta_{S_0}\|_2 = C$. In order to establish (A.8), it suffices to show that, for some small $\varepsilon > 0$ and sufficiently large $C > 0$,

$$P \left\{ \sup_{\|\Delta_{S_0}\|_2=C} Q_1(\beta_0 + \epsilon_n \Delta) < Q_1(\beta_0) \right\} \geq 1 - \varepsilon, \quad (\text{A.9})$$

as $n \rightarrow \infty$. If (A.9) holds, this will imply that with probability at least $1 - \varepsilon$, there exists a local maximum inside the ball $\{\beta_0 + \epsilon_n \Delta : \|\Delta_{S_0}\|_2 \leq C\}$ as $n \rightarrow \infty$. In order to establish (A.9), we must show that with probability tending

to one and sufficiently large $C > 0$,

$$\begin{aligned}
& Q_1(\beta_0 + \epsilon_n \Delta) - Q_1(\beta_0) \\
&= [\ell_n(\beta_0 + \epsilon_n \Delta) - \ell_n(\beta_0)] + \sum_{g=1}^{s_0} [\text{pen}(\beta_{0g} + \epsilon_n \Delta_g) - \text{pen}(\beta_{0g})] \\
&\stackrel{\Delta}{=} \text{(I)} + \text{(II)} < 0,
\end{aligned} \tag{A.10}$$

as $n \rightarrow \infty$. We bound the terms in (A.10) separately. We first focus on bounding (I) $:= \ell_n(\beta_0 + \epsilon_n \Delta) - \ell_n(\beta_0)$. Let ∇_{S_0} and $\nabla_{S_0}^2$ denote the partial derivatives with respect to β_{0S_0} . By a Taylor expansion and the fact that $\beta_0 = (\beta_{0S_0}^\top, \mathbf{0}^\top)^\top$, we have

$$\begin{aligned}
\text{(I)} &= \epsilon_n \nabla_{S_0} \ell_n(\beta_0)^\top \Delta_{S_0} + \frac{\epsilon_n^2}{2} \Delta_{S_0}^\top \nabla_{S_0}^2 \ell_n(\tilde{\beta}) \Delta_{S_0} \\
&\stackrel{\Delta}{=} J_1 + J_2,
\end{aligned} \tag{A.11}$$

where $\tilde{\beta}$ lies on the line segment between β_0 and $\beta_0 + \epsilon_n \Delta$. Next, we bound the terms J_1 and J_2 in (A.11) separately. We have

$$\nabla_{S_0} \ell_n(\beta_0) = \mathbf{X}_{S_0}^\top \text{diag}\{\xi'(\mathbf{X}\beta_0)\}(\mathbf{Y} - b'(\xi(\mathbf{X}\beta_0))).$$

Let d_2 denote the largest diagonal entry in $\text{diag}(\xi'(\mathbf{X}\beta_0))$, which we know is finite and positive due to Assumption (A4)(ii). Then we have, for some $A > 0$,

$$\begin{aligned}
P(\|\nabla_{S_0} \ell_n(\beta_0)\|_2 > An\epsilon_n) &\leq P\left(\left\|\frac{1}{n} \mathbf{X}_{S_0}^\top (\mathbf{Y} - b'(\xi(\mathbf{X}\beta_0)))\right\|_2^2 > A^2 \epsilon_n^2 / d_2^2\right) \\
&\leq \frac{d_2^2 n}{A^2 s_0 \log G} \mathbb{E} \left\{ \sum_{g=1}^{s_0} \sum_{k=1}^{m_g} \left[\frac{1}{n} \sum_{i=1}^n x_{igk} (y_i - b'(\xi(\mathbf{x}_i^\top \beta_0))) \right]^2 \right\} \\
&= \frac{d_2^2}{A^2 s_0 \log G} \text{tr} \left(\frac{1}{n} \mathbf{X}_{S_0}^\top \boldsymbol{\Omega}(\beta_0) \mathbf{X}_{S_0} \right) \\
&\leq \frac{d_2^2}{A^2 s_0 \log G} \left(\sum_{g=1}^{s_0} m_g \right) \lambda_{\max} \left(\frac{1}{n} \mathbf{X}_{S_0}^\top \boldsymbol{\Omega}(\beta_0) \mathbf{X}_{S_0} \right) \\
&\leq \frac{d_2^2 s_0 m_{\max}}{A^2 s_0 \log G} \lambda_{\max} \left(\frac{1}{n} \mathbf{X}_{S_0}^\top \boldsymbol{\Omega}(\beta_0) \mathbf{X}_{S_0} \right) \\
&\lesssim \frac{d_2^2 \log n}{A^2 \log G} \lambda_{\max} \left(\frac{1}{n} \mathbf{X}_{S_0}^\top \boldsymbol{\Omega}(\beta_0) \mathbf{X}_{S_0} \right) \rightarrow 0 \quad \text{as } n \rightarrow \infty,
\end{aligned} \tag{A.12}$$

where the last inequality of the display uses Assumption (A2) that $m_{\max} = O(\log n \wedge (\log G / \log n))$ and Assumption (A3)(ii) that $\lambda_{\max}(n^{-1} \mathbf{X}_{S_0}^\top \boldsymbol{\Omega}(\beta_0) \mathbf{X}_{S_0}) = o(\log G / \log n)$. Therefore, from (A.12), we have that $\|\nabla_{S_0} \ell_n(\beta_0)\|_2 = O_p(n\epsilon_n)$, and so an upper bound for J_1 in (A.11) is

$$J_1 \leq |J_1| \leq \epsilon_n \|\nabla_{S_0} \ell_n(\beta_0)\|_2 \|\Delta_{S_0}\|_2 = An\epsilon_n^2 \|\Delta_{S_0}\|_2.$$

for some $A > 0$. We also have by Assumption (A3)(ii) that

$$\begin{aligned} J_2 &= -\frac{\epsilon_n^2}{2} \mathbf{\Delta}_{S_0}^\top \left\{ \mathbf{X}_{S_0}^\top \mathbf{\Sigma}(\tilde{\beta}) \mathbf{X}_{S_0} \right\} \mathbf{\Delta}_{S_0} \\ &\leq -\frac{\epsilon_n^2}{2} \lambda_{\min} \left(\mathbf{X}_{S_0}^\top \mathbf{\Sigma}(\tilde{\beta}) \mathbf{X}_{S_0} \right) \|\mathbf{\Delta}_{S_0}\|_2^2 \\ &\lesssim -\frac{n\epsilon_n^2}{2} \|\mathbf{\Delta}_{S_0}\|_2^2. \end{aligned}$$

Therefore, an upper bound for (I) in (A.10) is

$$(I) \lesssim An\epsilon_n^2 \|\mathbf{\Delta}_{S_0}\|_2 - \frac{n\epsilon_n^2}{2} \|\mathbf{\Delta}_{S_0}\|_2^2 \quad (\text{A.13})$$

Now we turn our attention to upper-bounding (II) in (A.10). We first focus on bounding each of the summands in (II) from above. We first rewrite $\text{pen}(\beta_g)$ in (A.1) as

$$\text{pen}(\beta_g) = -\lambda_1 \|\beta_g\|_2 - \log[p_\theta^*(\beta_g)],$$

where $p_\theta^*(\beta_g)$ is defined as in (A.2). Therefore, for any one group g ,

$$\begin{aligned} &\text{pen}(\beta_{0g} + \epsilon_n \mathbf{\Delta}_g) - \text{pen}(\beta_{0g}) \\ &= \lambda_1 \{ \|\beta_{0g}\|_2 - \|\beta_{0g} + \epsilon_n \mathbf{\Delta}_g\|_2 \} + [\log(p_\theta^*(\beta_{0g})) - \log(p_\theta^*(\beta_{0g} + \epsilon_n \mathbf{\Delta}_g))] \\ &\leq \lambda_1 \epsilon_n \|\mathbf{\Delta}_g\|_2 - [\log(p_\theta^*(\beta_{0g} + \epsilon_n \mathbf{\Delta}_g)) - \log(p_\theta^*(\beta_{0g}))] \\ &\leq \lambda_1 \epsilon_n \|\mathbf{\Delta}_g\|_2 - \epsilon_n \left(\frac{\partial p_\theta^*(\tilde{\beta}_g)}{\partial \tilde{\beta}_g} \right)^\top \mathbf{\Delta}_g, \\ &\quad \text{where } \tilde{\beta}_g \text{ lies on line segment between } \beta_{0g} \text{ and } \beta_{0g} + \epsilon_n \mathbf{\Delta}_g, \\ &\leq \epsilon_n \|\mathbf{\Delta}_g\|_2 \left\{ \lambda_1 + \left\| \frac{\partial p_\theta^*(\tilde{\beta}_g)}{\partial \tilde{\beta}_g} \right\|_2 \right\}, \end{aligned} \quad (\text{A.14})$$

We examine the second term in brackets in (A.14). Noting that

$$p_\theta^*(\beta_g) = \frac{1}{1 + \left(\frac{1-\theta}{\theta}\right) \left(\frac{\lambda_0}{\lambda_1}\right)^{m_g} \exp[-(\lambda_0 - \lambda_1)\|\beta_g\|_2]} := \frac{1}{1 + c(\beta_g)}, \quad (\text{A.15})$$

we can see that when $\beta_g \neq \mathbf{0}_{m_g}$,

$$\frac{\partial p_\theta^*(\beta_g)}{\partial \beta_g} = \frac{(\lambda_0 - \lambda_1)c(\beta_g)}{(1 + c(\beta_g))^2} \times \frac{\beta_g}{\|\beta_g\|_2}.$$

Note that

$$\frac{(\lambda_0 - \lambda_1)c(\beta_g)}{(1 + c(\beta_g))^2} < (\lambda_0 - \lambda_1)c(\beta_g) \prec 1,$$

since the exponent term $\exp[-(\lambda_0 - \lambda_1)\|\beta_g\|_2]$ in $c(\beta_g)$ decays faster than $(\lambda_0 - \lambda_1)$ multiplied by the other terms in $c(\beta_g)$ as $n \rightarrow \infty$. Since it must be that $\tilde{\beta}_g$ in (A.14) is a nonzero vector, we have that for large n ,

$$\left\| \frac{\partial p_\theta^*(\tilde{\beta}_g)}{\partial \tilde{\beta}_g} \right\|_2 < \left\| \frac{\tilde{\beta}_g}{\|\tilde{\beta}_g\|_2} \right\|_2 = 1,$$

and so we can further bound (A.14) from above by $\epsilon_n(1 + \lambda_1)\|\Delta_g\|_2$. Therefore, an upper bound for (II) in (A.10) is

$$\begin{aligned} \text{(II)} &\lesssim \epsilon_n \sum_{g=1}^{s_0} (1 + \lambda_1) \|\Delta_g\|_2 \leq s_0 \epsilon_n (1 + \lambda_1) \max_{1 \leq g \leq s_0} \|\Delta_g\|_2 \\ &\leq s_0 \epsilon_n (1 + \lambda_1) \|\Delta_{S_0}\|_2. \end{aligned} \quad (\text{A.16})$$

Combining the upper bounds in (A.13) and (A.16), it is clear that

$$\text{(I)} + \text{(II)} \lesssim An\epsilon_n^2 \|\Delta_{S_0}\|_2 - \frac{n\epsilon_n^2}{2} \|\Delta_{S_0}\|_2^2 + s_0 \epsilon_n (1 + \lambda_1) \|\Delta_{S_0}\|_2,$$

and the right-hand side can be made negative by choosing $C := \|\Delta_{S_0}\|_2$ large enough. For large $C > 0$, the negative term above is the dominating term. This proves (A.10), and thus, $\hat{\beta} = (\hat{\beta}_{S_0}^\top, \mathbf{0}^\top)^\top$ is a local maximum over the subspace $\mathcal{T} = \{\beta \in \mathbb{R}^p : \beta_{S_0^c} = \mathbf{0}\}$ as $n \rightarrow \infty$.

Step 2 (sparsity). Since $\log \pi(\beta)$ is a nonconvex function, the KKT conditions (A.6) only give necessary conditions for $\hat{\beta}$ from Step 1 to be a local maximum on $\mathcal{T}$. However, by suitably modifying Theorem 1 of Fan and Lv (2011) to the grouped GLM setting with the log SSGL prior as the penalty function, we can also obtain *sufficient* conditions for $\hat{\beta} \in \mathcal{T}$ from Step 1 to be a strict local maximum over *all* of $\mathbb{R}^p$. This will be the case if, for sufficiently large n , $\hat{\beta}$ from Step 1 satisfies

$$\lambda_{\max} \left\{ \nabla_g^2 \ell_n(\hat{\beta}) + \nabla^2 \text{pen}(\hat{\beta}_g) \right\} < 0 \quad \text{for all } \hat{\beta}_g \neq \mathbf{0}_{m_g}, g \in S_0, \quad (\text{A.17})$$

where ∇_g^2 denotes the second partial derivative with respect to $\hat{\beta}_g$, $\text{pen}(\beta_g)$ is the penalty function in (A.1), and

$$\|\mathbf{X}_g^\top \text{diag}(\xi'(\mathbf{X}\hat{\beta}))[\mathbf{Y} - b'(\xi(\mathbf{X}\hat{\beta}))]\|_2 < \lambda_\theta^*(\mathbf{0}_{m_g}) \quad \text{for all } g \in S_0^c. \quad (\text{A.18})$$

That is, we require the KKT second order sufficiency condition to hold as $n \rightarrow \infty$, and the second inequality in (A.6) has to hold with a *strict* inequality.

We first prove (A.17). For $\hat{\beta}_g \neq \mathbf{0}_{m_g}$, we have by (6) and the chain rule for

second derivatives (with $f(\mathbf{u}) = \text{pen}(\mathbf{u})$ and $\mathbf{u} = g(\mathbf{v}) = \|\mathbf{v}\|_2$) that

$$\begin{aligned} & \nabla_g^2 \ell_n(\hat{\beta}) + \nabla^2 \text{pen}(\hat{\beta}_g) \\ &= -\mathbf{X}_g^\top \Sigma(\hat{\beta}) \mathbf{X}_g + p_\theta^*(\hat{\beta}_g)[1 - p_\theta^*(\hat{\beta}_g)](\lambda_0 - \lambda_1)^2 \frac{\hat{\beta}_g \hat{\beta}_g^\top}{\|\hat{\beta}_g\|_2^2} \\ & \quad - \lambda_\theta^*(\hat{\beta}_g) \left[\frac{1}{\|\hat{\beta}_g\|_2} \mathbf{I}_{m_g} - \frac{\hat{\beta}_g \hat{\beta}_g^\top}{\|\hat{\beta}_g\|_2^3} \right]. \end{aligned} \quad (\text{A.19})$$

Note that $\lambda_{\min}(\|\hat{\beta}_g\|_2^{-1} \mathbf{I}_{m_g} - \|\hat{\beta}_g\|_2^{-3} \hat{\beta}_g \hat{\beta}_g^\top) = 0$, since $\mathbf{I}_{m_g} - \|\hat{\beta}_g\|_2^{-2} \hat{\beta}_g \hat{\beta}_g^\top$ is the annihilator matrix. Using the fact that $\lambda_{\max}(-\mathbf{A}) = -\lambda_{\min}(\mathbf{A})$ for a square matrix $\mathbf{A}$, it follows from (A.19) that

$$\begin{aligned} & \lambda_{\max} \left\{ \nabla_g^2 \ell_n(\hat{\beta}) + \nabla^2 \text{pen}(\hat{\beta}_g) \right\} \\ & \leq -\lambda_{\min} \left(\mathbf{X}_g^\top \Sigma(\hat{\beta}) \mathbf{X}_g \right) + p_\theta^*(\hat{\beta}_g)[1 - p_\theta^*(\hat{\beta}_g)](\lambda_0 - \lambda_1)^2 \\ & \lesssim -n + \frac{c(\hat{\beta}_g)(\lambda_0 - \lambda_1)^2}{(1 + c(\hat{\beta}_g))^2}, \quad \text{where } c(\beta_g) \text{ is as in (A.15)} \\ & \leq -n + c(\hat{\beta}_g)(\lambda_0 - \lambda_1)^2 \\ & \lesssim -n + 1 < 0, \end{aligned}$$

for sufficiently large n , where we used Assumption (A3)(ii) in the third line and the fact that given our choices of hyperparameters for $(\theta, \lambda_0, \lambda_1)$ in Theorem 1, $c(\beta_g)(\lambda_0 - \lambda_1)^2 \prec 1$ for any $\beta_g \neq \mathbf{0}_{m_g}$ in the last line. Thus, we have shown that (A.17) holds for large n .

Next, we prove (A.18). It suffices to show that as $n \rightarrow \infty$,

$$P \left(\exists g \in S_0^c, \left\| \frac{\partial \ell_n(\hat{\beta})}{\partial \beta_g} \right\|_2 \geq \lambda_\theta^*(\mathbf{0}_{m_g}) \right) \rightarrow 0.$$

To ease the notation, let $h_g(\hat{\beta}) = \partial \ell_n(\hat{\beta}) / \partial \beta_g$. By a Taylor expansion around β_0 and the fact that $\hat{\beta}_{0S_0^c} = \mathbf{0}$, we have

$$h_g(\hat{\beta}) = h_g(\beta_0) + \nabla_{S_0} h_g(\tilde{\beta})^\top (\hat{\beta}_{S_0} - \beta_{0S_0}), \quad (\text{A.20})$$

where $\tilde{\beta}$ lies on the line segment between β_0 and $\hat{\beta}$. By (A.20), we have that

$$\begin{aligned}
& P\left(\exists g \in S_0^c, \|h_g(\hat{\beta})\|_2 \geq \lambda_\theta^*(\mathbf{0}_{m_g})\right) \\
& \leq P\left(\exists g \in S_0^c, \|h_g(\beta_0)\|_2 \geq \frac{\lambda_\theta^*(\mathbf{0}_{m_g})}{2}\right) \\
& \quad + P\left(\exists g \in S_0^c, \|\nabla_{S_0} h_g(\tilde{\beta})^\top (\hat{\beta}_{S_0} - \beta_{0S_0})\|_2 \geq \frac{\lambda_\theta^*(\mathbf{0}_{m_g})}{2}\right) \\
& \leq \sum_{g=s_0+1}^G P\left(\|h_g(\beta_0)\|_2 \geq \frac{\lambda_\theta^*(\mathbf{0}_{m_g})}{2}\right) \\
& \quad + \sum_{g=s_0+1}^G P\left(\|\nabla_{S_0} h_g(\tilde{\beta})^\top (\hat{\beta}_{S_0} - \beta_{0S_0})\|_2 \geq \frac{\lambda_\theta^*(\mathbf{0}_{m_g})}{2}\right) \\
& \triangleq T_1 + T_2. \tag{A.21}
\end{aligned}$$

We will bound each of the terms in (A.21) from above separately. First note that since $\lambda_0 = (1 - \theta)/\theta \asymp G^c, c > 2$, and $\lambda_1 \asymp n^{-1}$, we have that for some constant $M > 0$,

$$p_\theta^*(\mathbf{0}_{m_g}) = \frac{1}{1 + \left(\frac{1-\theta}{\theta}\right) \left(\frac{\lambda_0}{\lambda_1}\right)^{m_g}} = \frac{1}{1 + Mn\lambda_0^{m_g+1}},$$

and thus,

$$\begin{aligned}
\lambda_\theta^*(\mathbf{0}_{m_g}) &= \lambda_1 p_\theta^*(\mathbf{0}_{m_g}) + \lambda_0 [1 - p_\theta^*(\mathbf{0}_{m_g})] \\
&= \frac{\lambda_1 + \lambda_0 (Mn\lambda_0^{m_g+1})}{1 + Mn\lambda_0^{m_g+1}} > \lambda_0 \left(\frac{Mn\lambda_0^{m_g+1}}{1 + Mn\lambda_0^{m_g+1}} \right) > \frac{\lambda_0}{2}, \tag{A.22}
\end{aligned}$$

for large n .

We first examine each of the summands in T_1 in (A.21). Since all of the entries of $\mathbf{X}$ satisfy $|x_{ij}| = O(\log G)$ (by Assumption (A3)(i)), we have that for some constant $D > 0$ and any $g \in \{1, \dots, G\}$, $\|\mathbf{X}_g\|_{2,\infty} \leq D \log G \sqrt{m_g} \leq D \log G \sqrt{m_{\max}}$. Thus, in light of (A.22), we have that for sufficiently large n ,

$$\begin{aligned}
& P\left(\|h_g(\beta_0)\|_2 \geq \frac{\lambda_\theta^*(\mathbf{0}_{m_g})}{2}\right) < P\left(\|\mathbf{X}_g\|_{2,\infty} \|\mathbf{Y} - b'(\xi(\mathbf{X}\beta_0))\|_2 \geq \frac{\lambda_0}{2}\right) \\
& \leq P\left(\|\mathbf{Y} - b'(\xi(\mathbf{X}\beta_0))\|_2 \geq \frac{\lambda_0}{2D \log G \sqrt{m_{\max}}}\right) \\
& \leq 2 \exp\left(-\frac{\lambda_0^2 / (4D^2 (\log G)^2 m_{\max})}{2[2nG + L\lambda_0 / (D \log G \sqrt{m_{\max}})]}\right) \\
& = 2 \exp\left(-\frac{\lambda_0^2}{16D^2 (\log G)^2 m_{\max} nG + 8L\lambda_0 D \log G \sqrt{m_{\max}}}\right) \tag{A.23}
\end{aligned}$$

where we used Assumption (A4)(i) and the Bernstein inequality (specifically, Lemma 2.2.11 of van der Vaart and Wellner (2006)). Thus, from (A.23), we have as an upper bound for T_1 in (A.21),

$$T_1 < 2 \exp \left(\log G - \frac{\lambda_0^2}{16D^2(\log G)^2 m_{\max} n G + 8L\lambda_0 D \log G \sqrt{m_{\max}}} \right) \rightarrow 0, \quad (\text{A.24})$$

where we used the fact that $\lambda_0 \asymp G^c, c > 2$, and Assumptions (A1) and (A2) that $G \gg n$ and $m_{\max} = O(\log n \wedge (\log G / \log n))$. Thus, the second term in the exponent of (A.24) dominates the first term as $n \rightarrow \infty$.

We next examine each of the summands in T_2 in (A.21). In light of (A.22) that $\lambda_\theta^*(\mathbf{0}_{m_g}) > \lambda_0/2$, we have that for each $g \in S_0^c$ and sufficiently large n ,

$$\begin{aligned} & P \left(\|\nabla_{S_0} h_g(\tilde{\beta})^\top (\hat{\beta}_{S_0} - \beta_{0S_0})\|_2 \geq \frac{\lambda_\theta^*(\mathbf{0}_{m_g})}{2} \right) \\ & < P \left(\|\nabla_{S_0} h_g(\tilde{\beta})^\top (\hat{\beta}_{S_0} - \beta_{0S_0})\|_2 \geq \frac{\lambda_0}{4} \right) \\ & \leq P \left(\|\nabla_{S_0} h_g(\tilde{\beta})\|_{2,\infty} \|\hat{\beta}_{S_0} - \beta_{0S_0}\|_2 \geq \frac{\lambda_0}{4} \right) \\ & \leq P \left(\|\mathbf{X}_{S_0}^\top \Sigma(\tilde{\beta}) \mathbf{X}_g\|_{2,\infty} \|\hat{\beta}_{S_0} - \beta_{0S_0}\|_2 \geq \frac{\lambda_0}{4} \right) \\ & \lesssim \frac{4n\epsilon_n}{\lambda_0}, \end{aligned} \quad (\text{A.25})$$

To arrive at (A.25), we used the fact that $\nabla_{S_0} h_g(\tilde{\beta}) = \mathbf{X}_{S_0}^\top \Sigma(\tilde{\beta}) \mathbf{X}_g$ in the second line. In the last line, we used Assumption (A3)(iii) that for all $g \in S_0^c$, $\|\mathbf{X}_{S_0}^\top \Sigma(\tilde{\beta}) \mathbf{X}_g\|_{2,\infty} = O(n)$, the fact that $\|\hat{\beta}_{S_0} - \beta_{0S_0}\|_2 = O_p(\epsilon_n)$, and Markov's inequality. From (A.25), the fact that $\lambda_0 \asymp G^c, c > 2$, and $G \gg n$, we have as an upper bound for T_2 in (A.21),

$$T_2 \prec \frac{4Gn\epsilon_n}{\lambda_0} \prec \frac{Gn}{G^c} = \frac{n}{G^{c-1}} \rightarrow 0 \text{ as } n \rightarrow \infty. \quad (\text{A.26})$$

Combining (A.21), (A.24), and (A.26) allows us to conclude that (A.18) holds for sufficiently large n , and the proof is finished. $\square$

B Proofs for Section 4

To prove Theorems 2 and 3, we follow the proof strategy of Jeong and Ghosal (2021). However, a crucial difference is that we study an *absolutely continuous* spike-and-slab prior in this paper, whereas Jeong and Ghosal (2021) prove their results for a *point-mass* spike-and-slab prior. The continuous nature of the SSGL requires it to be handled quite differently from the point mass prior of Jeong and Ghosal (2021). In particular, since SSGL (10) puts zero probability

on exactly sparse configurations of β , we need to instead rely on notions of “approximate” sparsity to appropriately bound the approximation error.

We first define some necessary notation. Let f_i be the density of y_i belonging to the exponential family (1), where the natural parameter θ_i is related to the grouped covariates $\mathbf{x}_i = (\mathbf{x}_{i1}^\top, \dots, \mathbf{x}_{iG}^\top)^\top$ through (3). Denote f_{0i} analogously with true natural parameter θ_{0i} . Then the joint densities are $f = \prod_{i=1}^n f_i$ and $f_0 = \prod_{i=1}^n f_{0i}$. We define the average squared Hellinger metric between f and f_0 as $H_n^2(f, f_0) = n^{-1} \sum_{i=1}^n H^2(f_i, f_{0i})$, where $H(f_i, f_{0i}) = (\int (\sqrt{f_i} - \sqrt{f_{0i}})^2)^{1/2}$ is the Hellinger distance between f_i and f_{0i} . The Kullback-Leibler (KL) divergence and KL variation between f_i and f_{0i} are given by $K(f_{0i}, f_i) = \int f_{0i} \log(f_{0i}/f_i)$ and $V(f_{0i}, f_i) = \int f_{0i} (\log(f_{0i}/f_i) - K(f_{0i}, f_i))^2$ respectively. For a set Ω with semimetric d , the ε -covering number $N(\varepsilon, \Omega, d)$ is the minimal number of d -balls of radius ε needed to cover Ω . Meanwhile, the ε -packing number for Ω , denoted by $D(\varepsilon, \Omega, d)$, is the maximal number of d balls of radius ε needed to cover Ω .

B.1 Proof of Theorem 2

We first prove a lemma that is needed to prove Theorem 2.

Lemma 1 (evidence lower bound). *Assume the same setup as Theorem 2. Then $\sup \mathbb{P}_0(\mathcal{E}_n^c) \rightarrow 0$, where the set $\mathcal{E}_n$ is defined as*

$$\mathcal{E}_n = \left\{ \int \prod_{i=1}^n \frac{f_i(y_i)}{f_{0i}(y_i)} d\Pi(\beta) \geq e^{-C_1 n \epsilon_n^2} \right\}, \quad (\text{B.1})$$

for some constant $C_1 > 0$ and $\epsilon_n^2 = s_0 \log G/n$.

Proof. By Lemma 10 of Ghosal and van der Vaart (2007), this statement will be proven if we can show that

$$\Pi(\mathcal{B}_n) \gtrsim e^{-C_1 n \epsilon_n^2}, \quad (\text{B.2})$$

where the set $\mathcal{B}_n$ is defined as the event,

$$\mathcal{B}_n = \left\{ \frac{1}{n} \sum_{i=1}^n K(f_{0i}, f_i) \leq \epsilon_n^2, \frac{1}{n} \sum_{i=1}^n V(f_{0i}, f_i) \leq \epsilon_n^2 \right\}. \quad (\text{B.3})$$

As established in Lemma 1 of Jeong and Ghosal (2021), the KL divergence and KL variation for the exponential family (1) is given by

$$\begin{aligned} K(f_{0i}, f_i) &= (\theta_{0i} - \theta_i) b'(\theta_{0i}) - b(\theta_{0i}) + b(\theta_i), \\ V(f_{0i}, f_i) &= b''(\theta_{0i}) (\theta_i - \theta_{0i})^2. \end{aligned}$$

Recalling that $\theta_i = \xi(\mathbf{x}_i^\top \beta)$ and $\theta_{0i} = \xi(\mathbf{x}_i^\top \beta_0)$ where $\xi = (h \circ b')^{-1}$, we have by Taylor expansion in $\mathbf{x}_i^\top \beta$ at $\mathbf{x}_i^\top \beta_0$ that

$$\begin{aligned} & \max \{K(f_{0i}, f_i), V(f_{0i}, f_i)\} \\ & \leq b''(\theta_{0i}) (\xi'(\mathbf{x}_i^\top \beta_0))^2 (\mathbf{x}_i^\top \beta - \mathbf{x}_i^\top \beta_0)^2 + o((\mathbf{x}_i^\top \beta - \mathbf{x}_i^\top \beta_0)^2) \\ & = (h^{-1})'(\mathbf{x}_i^\top \beta_0) \xi'(\mathbf{x}_i^\top \beta_0) (\mathbf{x}_i^\top \beta - \mathbf{x}_i^\top \beta_0)^2 + o((\mathbf{x}_i^\top \beta - \mathbf{x}_i^\top \beta_0)^2). \end{aligned}$$

Note that h^{-1} and ξ are strictly increasing. Thus, it must be the case that if $\max_{1 \leq i \leq n} \{(h^{-1})'(\mathbf{x}_i^\top \beta_0) \xi'(\mathbf{x}_i^\top \beta_0) (\mathbf{x}_i^\top \beta - \mathbf{x}_i^\top \beta_0)^2\} \leq \epsilon_n^2 \rightarrow 0$, then $|\mathbf{x}_i^\top \beta - \mathbf{x}_i^\top \beta_0| \rightarrow 0$ for all $1 \leq i \leq n$. Recall $\Omega(\beta_0) = \text{diag}\{(h^{-1})'(\mathbf{X}\beta_0)\} \text{diag}\{\xi'(\mathbf{X}\beta_0)\}$. Then both $n^{-1} \sum_{i=1}^n K(f_{0i}, f_i)$ and $n^{-1} \sum_{i=1}^n V(f_{0i}, f_i)$ can be bounded above by a constant multiple of $n^{-1} \|\Omega^{1/2}(\beta_0) \mathbf{X}(\beta - \beta_0)\|_2^2$ on $\mathcal{B}_n$.

Thus, for sufficiently large n , we have for some constants $\tilde{C}_2, C_2 > 0$,

$$\begin{aligned} \Pi(\mathcal{B}_n) &\geq \Pi\left(\frac{1}{n} \|\Omega^{1/2}(\beta_0) \mathbf{X}(\beta - \beta_0)\|_2^2 \leq \tilde{C}_2^2 \epsilon_n^2\right) \\ &\geq \Pi\left(\lambda_{\max}\left(\frac{1}{n} \mathbf{X}_g^\top \Omega(\beta_0) \mathbf{X}_g\right) \left(\sum_{g=1}^G \|\beta_g - \beta_{0g}\|_2\right)^2 \leq \tilde{C}_2^2 \epsilon_n^2\right) \\ &\geq \Pi\left(\sum_{g=1}^G \|\beta_g - \beta_{0g}\|_2 \leq \frac{C_2 \sqrt{s_0}}{\sqrt{n}}\right) \\ &\geq \Pi\left(\sum_{g \in S_0} \|\beta_g - \beta_{0g}\|_2 \leq \frac{C_2 \sqrt{s_0}}{2\sqrt{n}}\right) \Pi\left(\sum_{g \in S_0^c} \|\beta_g - \beta_{0g}\|_2 \leq \frac{C_2 \sqrt{s_0}}{2\sqrt{n}}\right) \\ &\geq \prod_{g \in S_0} \Pi\left(\|\beta_g - \beta_{0g}\|_2^2 \leq \frac{C_2^2}{4n}\right) \prod_{g \in S_0^c} \Pi\left(\|\beta_{S_0^c}\|_2^2 \leq \frac{C_2^2 s_0}{4n(G-s_0)}\right), \quad (\text{B.4}) \end{aligned}$$

where we used Assumption (B3)(ii) in the third inequality and the Cauchy-Schwarz inequality in the fifth inequality. Arguing similarly as in (D.17)-(D.20) in the proof of Theorem 2 of Bai et al. (2022), we have that

$$\prod_{g \in S_0} \Pi\left(\|\beta_g - \beta_{0g}\|_2^2 \leq \frac{C_2^2}{4n}\right) \geq \theta^{s_0} e^{-\lambda_1 \|\beta_{S_0}\|_2} e^{-\lambda_1 C_2/2\sqrt{n}} \prod_{g \in S_0} \frac{C_g}{m_g!} \left(\frac{\lambda_1 C_2}{2s_0 \sqrt{n}}\right)^{m_g},$$

where $C_g = 2^{-m_g} \pi^{-(m_g-1)/2} [\Gamma((m_g+1)/2)]^{-1}$ and

$$\prod_{g \in S_0^c} \Pi\left(\|\beta_g\|_2^2 \leq \frac{C_2^2 s_0}{4n(G-s_0)}\right) \geq (1-\theta)^{G-s_0} \left[1 - \frac{4nGm_{\max}(m_{\max}+1)}{\lambda_0^2 C_2^2 s_0}\right]^{G-s_0}.$$

Thus, a lower bound on (B.4) is

$$\begin{aligned} \Pi(\mathcal{B}_n) &\geq \theta^{s_0} (1-\theta)^{G-s_0} e^{-\lambda_1 \|\beta_{S_0}\|_2} e^{-\lambda_1 C_2/2\sqrt{n}} \left[1 - \frac{4nGm_{\max}(m_{\max}+1)}{\lambda_0^2 C_2^2 s_0}\right]^{G-s_0} \\ &\quad \times \prod_{g \in S_0} \frac{C_g}{m_g!} \left(\frac{\lambda_1 C_2}{2s_0 \sqrt{n}}\right)^{m_g}. \quad (\text{B.5}) \end{aligned}$$

Now, since $\lambda_0 = (1-\theta)/\theta \asymp G^c, c > 2$, and $s_0 = o(n^{1/2}/\log G)$ and $m_{\max} = O(\log n \wedge (\log G/\log n))$ by Assumption (A2), we have that

$$\left[1 - \frac{4nGm_{\max}(m_{\max}+1)}{\lambda_0^2 C_2^2 s_0}\right]^{G-s_0} \gtrsim \left[1 - \frac{1}{G-s_0}\right]^{G-s_0} \rightarrow e^{-1} \text{ as } n \rightarrow \infty.$$

Furthermore, $\theta \asymp \frac{1}{G^c+1}$, and thus,

$$(1 - \theta)^{G-s_0} \gtrsim \left(1 - \frac{1}{G-s_0}\right)^{G-s_0} \rightarrow e^{-1}.$$

Thus, using Assumption (B4) that $\|\beta_0\|_\infty = O(\log G)$ and $\|\beta_{S_0}\|_2 \leq \sqrt{s_0}\|\beta_0\|_\infty$, we can further lower bound (B.5) as

$$\Pi(\mathcal{B}_n) \gtrsim (G^c + 1)^{-s_0} e^{-\lambda_1(\sqrt{s_0} \log G + C_1/2\sqrt{n})} \prod_{g \in S_0} \frac{C_g}{m_g!} \left(\frac{\lambda_1 C_2}{2s_0 \sqrt{n}}\right)^{m_g},$$

and since $\lambda_1 \asymp 1/n$, we have

$$\begin{aligned} & -\log \Pi(\mathcal{B}_n) \\ & \lesssim 2cs_0 \log G + \sqrt{s_0} \log G + \sum_{g \in S_0} \left[\log(m_g!) - \log(C_g) + m_g \log \left(\frac{2s_0 \sqrt{n}}{\lambda_1 C_2} \right) \right]. \end{aligned} \quad (\text{B.6})$$

Similar to (D.24) in the proof of Theorem 2 of Bai et al. (2022), the first term of (B.6) can be shown to be the dominating term, and thus, $-\log \Pi(\mathcal{B}_n) \lesssim n\epsilon_n^2$. This establishes (B.2) and the proof is finished. $\square$

Proof of Theorem 2 We follow the proof strategy of Theorem 2 in Jeong and Ghosal (2021) but work with the generalized dimensionality $|\nu(\beta)|$ (20). Let $\mathcal{E}_n$ be the event defined in (B.1), where $\epsilon_n = (s_0 \log G/n)^{1/2}$. Define the set $\mathcal{A}_n = \{\beta : |\nu(\beta)| \leq K_1 s_0\}$, where $K_1 \geq C_1 \vee 1$, and C_1 is the constant in Lemma 1. Then we have

$$\mathbb{E}_0 \Pi(\mathcal{A}_n^c \mid \mathbf{Y}) \leq \mathbb{E}_0 \Pi(\mathcal{A}_n^c \mid \mathbf{Y}) \mathbb{1}_{\mathcal{E}_n} + \mathbb{P}_0(\mathcal{E}_n^c).$$

By Lemma 1, $\mathbb{P}_0(\mathcal{E}_n^c) \rightarrow 0$ as $n \rightarrow \infty$, so it suffices to show that $\mathbb{E}_0 \Pi(\mathcal{A}_n^c \mid \mathbf{Y}) \mathbb{1}_{\mathcal{E}_n} \rightarrow 0$. Note that

$$\Pi(\mathcal{A}_n^c \mid \mathbf{Y}) = \frac{\int_{\mathcal{A}_n^c} \prod_{i=1}^n \frac{f_i(y_i)}{f_{0i}(y_i)} d\Pi(\beta)}{\int \prod_{i=1}^n \frac{f_i(y_i)}{f_{0i}(y_i)} d\Pi(\beta)}. \quad (\text{B.7})$$

On the set $\mathcal{E}_n$, the denominator in (B.7) is bounded below by $e^{-C_1 n \epsilon_n^2}$ by Lemma 1. On the other hand, an upper bound for the expected value of the numerator is

$$\mathbb{E}_0 \left(\int_{\mathcal{A}_n^c} \prod_{i=1}^n \frac{f_i(y_i)}{f_{0i}(y_i)} d\Pi(\beta) \right) \leq \int_{\mathcal{A}_n^c} d\Pi(\beta) = \Pi(|\beta| > K_1 s_0). \quad (\text{B.8})$$

Now, since $\pi(\beta_g \mid \theta) < 2\theta C_g \lambda_1^{m_g} e^{-\lambda_1 \|\beta_g\|_2^2}$ for all $\|\beta_g\|_2 > \omega_g$, we have (arguing as in (D.30) of Bai et al. (2022)) that

$$\begin{aligned} \Pi(|\nu(\beta)| > K_1 s_0) &\leq \sum_{S: |S| > C_3 s_0} 2^{|S|} \theta^{|S|} \left\{ \prod_{g \in S} \int_{\|\beta_g\|_2 > \omega_g} C_g \lambda_1^{m_g} e^{-\lambda_1 \|\beta_g\|_2^2} d\beta_g \right\} \\ &\quad \times \left\{ \prod_{g \in S^c} \int_{\|\beta_g\|_2 \leq w_g} \pi(\beta_g) d\beta_g \right\} \\ &\lesssim \sum_{S: |S| > K_1 s_0} \theta^{|S|}, \end{aligned}$$

where we bounded the first bracketed term from above as $\prod_{g \in S} (1/n)^{m_g} \leq n^{-|S|}$ and the second bracketed term from above by one. Now, since $\theta < 1/G^2$, we have

$$\begin{aligned} \sum_{S: |S| > K_1 s_0} \theta^{|S|} &\leq \sum_{k=\lfloor C_1 s_0 \rfloor + 1}^G \binom{G}{k} \left(\frac{1}{G^2}\right)^k \leq \sum_{k=\lfloor K_1 s_0 \rfloor + 1}^G \left(\frac{e}{kG}\right)^k \\ &< \sum_{k=\lfloor K_1 s_0 \rfloor + 1}^G \left(\frac{e}{G \lfloor K_1 s_0 \rfloor + 1}\right)^k \\ &\leq G^{-(\lfloor K_1 s_0 \rfloor + 1)} \lesssim e^{-K_1 n \epsilon_n^2}. \quad (\text{B.9}) \end{aligned}$$

Combining (B.7)-(B.9) and Lemma 1 gives $\mathbb{E}_0 \Pi(\mathcal{A}_n^c \mid \mathbf{Y}) \mathbb{1}_{\mathcal{E}_n} \prec e^{-(K_1 - C_1) n \epsilon_n^2} \rightarrow 0$, since $K_1 > C_1$. This completes the proof. $\square$

B.2 Proof of Theorem 3

We first prove a lemma which verifies that the prior puts sufficient mass on the event that the generalized dimensionality $|\nu(\beta)|$ (20) equals s_0 , where s_0 is the true number of nonzero groups in β_0 .

Lemma 2. *Assume the same setup as Theorem 2. Then for some constant $C_3 > 0$,*

$$\Pi(|\nu(\beta)| = s_0) \geq e^{-C_3 n \epsilon_n^2}$$

Proof. Note that $\pi(\beta_g) \geq \Psi(\beta_g \mid \lambda_0)$ for all β_g . Further, for $\beta_g \sim \Psi(\beta_g \mid \lambda_0)$, we have that $\|\beta_g\|_2$ follows a gamma distribution with shape parameter m_g and

scale parameter λ_0 . Thus,

$$\begin{aligned}
\Pi(\|\beta_g\|_2 > \omega_g) &\geq \int_{\omega_g}^{\infty} \frac{1}{\Gamma(m_g)\lambda_0^{m_g}} x^{m_g-1} e^{-x/\lambda_0} dx \\
&\geq \frac{1}{\Gamma(m_g)\lambda_0^{m_g}} \int_{\lambda_0}^{\infty} x^{m_g-1} e^{-x/\lambda_0} dx \\
&\geq \frac{1}{\Gamma(m_g)\lambda_0} \int_{\lambda_0}^{\infty} e^{-x/\lambda_0} dx \\
&= \frac{e^{-1}}{\Gamma(m_g)} \geq \frac{e^{-1}}{\Gamma(m_{\max})} > \frac{e^{-1}}{m_{\max}!}.
\end{aligned} \tag{B.10}$$

Meanwhile,

$$\begin{aligned}
\Pi(\|\beta_g\|_2 \leq \omega_g) &\geq \int_0^{\omega_g} \frac{1}{\Gamma(m_g)\lambda_0^{m_g}} x^{m_g-1} e^{-x/\lambda_0} dx \\
&= 1 - \int_{\omega_g}^{\infty} \frac{1}{\Gamma(m_g)\lambda_0^{m_g}} x^{m_g-1} e^{-x/\lambda_0} dx \\
&\geq 1 - \exp(-\lambda_0\omega_g + m_{\max}) \\
&\geq 1 - \frac{1}{G - s_0},
\end{aligned} \tag{B.11}$$

where we used the tail bound for the gamma density on p. 29 of Boucheron et al. (2013) and the inequality $1 + x - \sqrt{1 + 2x} \geq (x - 1)/2$ for $x > 0$ (similarly as in Ning et al. (2020)) to bound the integral in the second line from above by $\exp(-\lambda_0\omega_g + m_{\max})$.

From (B.10)-(B.11), we can bound $\Pi(|\nu(\beta)| = s_0)$ from below as

$$\begin{aligned}
\Pi(|\nu(\beta)| = s_0) &\geq \left(\frac{e^{-1}}{m_{\max}!} \right)^{s_0} \left(1 - \frac{1}{G - s_0} \right)^{G - s_0} \\
&\gtrsim \left(\frac{e^{-1}}{m_{\max}^{m_{\max}}} \right)^{s_0},
\end{aligned}$$

where we used the facts that $(1 - 1/(G - s_0))^{G - s_0} \rightarrow e^{-1}$ as $n \rightarrow \infty$, and $x! \leq x^x$ in the second line of the display. Therefore, we see that for sufficiently large n ,

$$-\log \Pi(|\nu(\beta)| = s_0) \leq s_0 m_{\max} \log m_{\max} + s_0 \lesssim s_0 \log G, \tag{B.12}$$

by Assumptions (A1) and (A2) that $G \gg n$ and $m_{\max} = O(\log n \wedge (\log G / \log n))$. This proves the lemma. $\square$ $\square$

Proof of Theorem 3 The proof of Theorem 3 is broken up into two parts and follows the proof strategy of Theorems 2 and 3 of Jeong and Ghosal (2021). In the first part, we establish the posterior contraction rate in terms of the average squared Hellinger metric. In the second part, we convert the Hellinger contraction rate result to a contraction rate for the regression coefficients themselves.

Step 1: Hellinger contraction rate. We first show that there exists $C_4 > 0$ such that

$$\mathbb{E}_0 \Pi(\beta \in \mathbb{R}^p : H_n(\beta, \beta_0) > C_4 \epsilon_n \mid \mathbf{Y}) \rightarrow 0 \text{ as } n \rightarrow \infty. \quad (\text{B.13})$$

Define the set $\mathcal{A}_n = \{\beta \in \mathbb{R}^p : |\nu(\beta)| \leq K_1 s_0\}$, where K_1 is the constant from Theorem 2. Then for every $\epsilon > 0$,

$$\begin{aligned} \mathbb{E}_0 \Pi(\beta \in \mathbb{R}^p : H_n(\beta, \beta_0) > \epsilon \mid \mathbf{Y}) \\ \leq \mathbb{E}_0 \Pi(\beta \in \mathcal{A}_n : H_n(\beta, \beta_0) > \epsilon \mid \mathbf{Y}) \mathbb{1}_{\mathcal{E}_n} + \mathbb{E}_0 \Pi(\mathcal{A}_n^c \mid \mathbf{Y}) + \mathbb{P}_0 \mathcal{E}_n^c, \end{aligned}$$

where $\mathcal{E}_n$ is in the event in (B.1). By Lemma 1 and Theorem 2, the second and third terms on the right-hand side above tend to zero uniformly over β_0 . Hence, in order to prove (B.13), we only focus on the first term. That is, it suffices to show that for $\epsilon = C\epsilon_n$, where $C > 0$ and $\epsilon_n = (s_0 \log G/n)^{1/2}$,

$$\mathbb{E}_0 \Pi(\beta \in \mathcal{A}_n : H_n(\beta, \beta_0) > \epsilon \mid \mathbf{Y}) \text{ as } n \rightarrow \infty, \quad (\text{B.14})$$

uniformly over β_0 . Define the sieve,

$$\mathcal{A}_n^* = \{\beta \in \mathcal{A}_n : \|\beta - \beta_0\|_2 \leq G^M\}, \quad (\text{B.15})$$

where $M > 0$ is defined as in Assumption (B3)(ii). By Lemma 2 of Ghosal and van der Vaart (2007), there exists a test function φ_n such that for any β_1 with $H_n(\beta_0, \beta_1) > \epsilon$,

$$\mathbb{E}_0 \varphi_n \leq \exp(-n\epsilon^2/2), \quad \sup_{\beta: H_n(\beta, \beta_1) \leq \epsilon/18} \mathbb{E}_\beta(1 - \varphi_n) \leq \exp(-n\epsilon^2/2).$$

We want to use Lemma 9 of Ghosal and van der Vaart (2007). In order to do so, we need to show that $\log N(\epsilon_n/36, \mathcal{A}_n^*, H_n) \lesssim n\epsilon_n^2$ for $\epsilon_n = (s_0 \log G/n)^{1/2}$. Since the squared Hellinger distance is bounded by one half of the KL divergence, we have by a Taylor expansion that for any $\beta_1, \beta_2 \in \mathcal{A}_n^*$,

$$\begin{aligned} H^2(f_{\beta_1, i}, f_{\beta_2, i}) &\leq \frac{K(f_{\beta_1, i}, f_{\beta_2, i})}{2} \\ &= \frac{(h^{-1})'(\mathbf{x}_i^\top \beta_1) \xi'(\mathbf{x}_i^\top \beta_1)}{2} (\mathbf{x}_i^\top \beta_1 - \mathbf{x}_i^\top \beta_2)^2 + o((\mathbf{x}_i^\top \beta_1 - \mathbf{x}_i^\top \beta_2)^2). \end{aligned}$$

It then follows that

$$\begin{aligned} H_n^2(f_{\beta_1}, f_{\beta_2}) &= \frac{1}{2n} \|\boldsymbol{\Omega}^{1/2}(\beta_1) \mathbf{X}(\beta_1 - \beta_2)\|_2^2 + o(\|\mathbf{X}(\beta_1 - \beta_2)\|_2^2) \\ &\leq \frac{1}{2} \max_{1 \leq g \leq G} \left(\lambda_{\max} \left(\frac{1}{n} \mathbf{X}_g^\top \boldsymbol{\Omega}(\beta_1) \mathbf{X}_g \right) \right) \left(\sum_{g=1}^G \|\beta_{1g} - \beta_{2g}\|_2 \right)^2 \\ &\quad + o(\|\mathbf{X}(\beta_1 - \beta_2)\|_2^2) \\ &\lesssim G \log G \|\beta_1 - \beta_2\|_2^2 + \|\mathbf{X}\|_{2, \infty}^2 \|\beta_1 - \beta_2\|_2^2 \\ &\lesssim G^2 \|\beta_1 - \beta_2\|_2^2, \end{aligned}$$

where we used Assumption (B3)(ii) that $\lambda_{\max}(n^{-1}\mathbf{X}_g^\top \boldsymbol{\Omega}(\boldsymbol{\beta}_1)\mathbf{X}_g) \lesssim \log G$ and the fact that all of the entries of $\mathbf{X}$ are uniformly bounded, and thus, $\|\mathbf{X}\|_{2,\infty} \lesssim G^{1/2}$. Thus, for some $C_5 > 0$, we can bound $N(\epsilon_n/36, \mathcal{A}_n^*, H_n)$ from above by

$$\begin{aligned} N(\epsilon_n/36, \mathcal{A}_n^*, H_n) &\leq N\left(\frac{C_5\epsilon_n}{36G}, \mathcal{A}_n^*, \|\cdot\|_2\right) \\ &= N\left(\frac{C_5\epsilon_n}{36G}, \{\boldsymbol{\beta} \in \mathcal{A}_n, \|\boldsymbol{\beta} - \boldsymbol{\beta}_0\|_2 \leq G^M\}, \|\cdot\|_2\right). \end{aligned} \quad (\text{B.16})$$

Denote by $\tilde{\omega}$,

$$\tilde{\omega} = \frac{1}{\lambda_0 - \lambda_1} \log \left[\frac{1 - \theta}{\theta} \frac{\lambda_0^{m_{\max}}}{\lambda_1^{m_{\max}}} \right],$$

It is clear that given our choice of hyperparameters for $(\theta, \lambda_0, \lambda_1)$, the threshold ω_g in (21) is an increasing function of m_g . Therefore,

$$\begin{aligned} &\{\boldsymbol{\beta} \in \mathcal{A}_n, \|\boldsymbol{\beta} - \boldsymbol{\beta}_0\|_2 \leq G^M\} \\ &\subset \{\boldsymbol{\beta} \in \mathbb{R}^p : \|\boldsymbol{\beta}_{S_0} - \boldsymbol{\beta}_{0S_0}\|_2 \leq G^M, \text{ and } \|\boldsymbol{\beta}_g\|_2 \leq \tilde{\omega} \text{ for all } g \in S_0^c\} \\ &\subset \{\|\boldsymbol{\beta}_{S_0} - \boldsymbol{\beta}_{0S_0}\|_2 \leq G^M\} \times [-\tilde{\omega}, \tilde{\omega}]^{G-K_1s_0}. \end{aligned}$$

By Lemma S4 of Shen et al. (2024), for any set of the form $E = A \times [-\delta, \delta]^{Q-s} \subset \mathbb{R}^Q$ where $A \subset \mathbb{R}^s$ and $s < Q$,

$$D(\epsilon, E, \|\cdot\|_2) \leq D((1 - T^{-1})^{1/2}\epsilon, A, \|\cdot\|_2),$$

if $\delta < \epsilon/(2[T(Q-s)]^{1/2})$, for some $T > 1$, where D denotes the packing number. In order to use Lemma A4 of Shen et al. (2024), we can check that for sufficiently large n and $T = 2$,

$$\tilde{\omega} = \frac{1}{\lambda_0 - \lambda_1} \log \left[\frac{1 - \theta}{\theta} \frac{\lambda_0^{m_{\max}}}{\lambda_1^{m_{\max}}} \right] \lesssim \frac{m_{\max} \log G}{G^2} < \frac{C_5\epsilon_n/36G}{2[2T(G-s_0)]^{1/2}},$$

where we used Assumptions (A1)-(A2) and the fact that $\lambda_0 = (1 - \theta)/\theta \asymp p^c$, $c > 2$, and $\lambda_1 \asymp 1/n$. Hence, by Lemma S4 of Shen et al. (2024) and the fact that we can upper bound the covering number by the packing number, we can further upper bound (B.16) by

$$\begin{aligned} &\binom{G}{K_1s_0} D\left(\frac{C_5\epsilon_n}{36G\sqrt{2}}, \{\|\boldsymbol{\beta}_S - \boldsymbol{\beta}_{0S}\|_2 \leq G^M\}, \|\cdot\|_2\right) \\ &\leq \binom{G}{K_1s_0} \left(\frac{108\sqrt{2}G^{M+1}}{C_5\epsilon_n}\right)^{K_1s_0}, \end{aligned} \quad (\text{B.17})$$

noting that on the set $\mathcal{A}_n$, $|S| \leq K_1s_0$. From (B.16)-(B.17), we have that

$$\begin{aligned} \log N(\epsilon_n/36, \mathcal{A}_n^*, H_n) &\leq K_1s_0 \log G + K_1s_0 \left[(M+1) \log G + \log(108\sqrt{2n}) \right] \\ &\lesssim n\epsilon_n^2, \end{aligned} \quad (\text{B.18})$$

where we used the fact that $\binom{G}{K_1 s_0} \leq G^{K_1 s_0}$. Having established (B.18), we can now use Lemma 9 of Ghosal and van der Vaart (2007), which implies that for every $\epsilon > \epsilon_n$, there exists a test $\bar{\varphi}_n$ such that

$$\mathbb{E}_0 \bar{\varphi}_n \leq \frac{1}{2} \exp(C_6 n \epsilon_n^2 - n \epsilon / 2), \quad \sup_{\beta \in \mathcal{A}_n^* : H_n(\beta, \beta_0) > \epsilon} \mathbb{E}_\beta (1 - \bar{\varphi}_n) \leq \exp(-n \epsilon^2 / 2).$$

for some $C_6 > 0$. By Lemma 2, $\Pi(|\nu(\beta)| = s_0) \geq e^{-C_3 n \epsilon_n^2}$ for some $C_3 > 0$, and thus,

$$\begin{aligned} & \mathbb{E}_0 \Pi(\beta \in \mathcal{A}_n : H_n(\beta, \beta_0) > \epsilon \mid \mathbf{Y}) \\ & \leq \mathbb{E}_0 \Pi(\beta \in \mathcal{A}_n : H_n(\beta, \beta_0) > \epsilon \mid \mathbf{Y}) \mathbb{1}_{\mathcal{E}_n}(1 - \bar{\varphi}_n) + \mathbb{E}_0 \bar{\varphi}_n \\ & \leq \left\{ \sup_{\beta \in \mathcal{A}_n^* : H_n(\beta, \beta_0) > \epsilon} \mathbb{E}_\beta (1 - \bar{\varphi}_n) + \Pi(\mathcal{A}_n \setminus \mathcal{A}_n^*) \right\} e^{C_3 n \epsilon_n^2} + \mathbb{E}_0 \bar{\varphi}_n. \end{aligned}$$

All of the terms in the display except $\Pi(\mathcal{A}_n \setminus \mathcal{A}_n^*) e^{C_3 n \epsilon_n^2}$ go to zero by choosing $\epsilon = C_4 \epsilon_n$ for a sufficiently large $C_4 > C_3$. To complete the proof then, we must show that

$$\Pi(\mathcal{A}_n \setminus \mathcal{A}_n^*) e^{C_3 n \epsilon_n^2} \rightarrow 0 \text{ as } n \rightarrow \infty. \quad (\text{B.19})$$

To establish (B.19), note that

$$\begin{aligned} \Pi(\mathcal{A}_n \setminus \mathcal{A}_n^*) &= \Pi(\beta \in \mathbb{R}^p : |\nu(\beta)| \leq K_1 s_0, \|\beta - \beta_0\|_2 > G^M) \\ &\leq \Pi(\|\beta - \beta_0\|_2 > G^M) \\ &\leq \sum_{g \in S_0} \Pi(\|\beta_g\|_2 > G^M - \|\beta_{0g}\|_2) + \sum_{g \in S_0^c} \Pi(\|\beta_g\|_2 > G^M). \end{aligned} \quad (\text{B.20})$$

We examine each of the terms in (B.20) separately. Notice that $\pi(\beta_g) \leq \Psi(\beta_g \mid \lambda_1)$ for all β_g . Moreover, for $\Psi(\beta_g \mid \lambda_1)$, $\|\beta_g\|_2$ follows a gamma distribution with shape parameter m_g and scale parameter λ_1 . Letting X_g denote a gamma random variable with shape and scale parameters m_g and λ_1 , and recalling that $M \geq 1$, we have

$$\begin{aligned} \sum_{g \in S_0} \Pi(\|\beta_g\|_2 > G^M - \|\beta_{0g}\|_2) &\leq \sum_{g \in S_0} P(X_g > G^M - \|\beta_{0g}\|_2) \\ &\leq \sum_{g \in S_0} P(X_g > G^M - C_7 \sqrt{G} \log G) \\ &\leq \sum_{g \in S_0} \exp \left[-\lambda_1 (G^M - C_7 \sqrt{G} \log G) + m_{\max} \right] \\ &= \exp \left[-\lambda_1 G^M + \lambda_1 C_7 \sqrt{G} \log G + m_{\max} + \log s_0 \right], \end{aligned} \quad (\text{B.21})$$

where in the second line, we used the fact that by Assumption (B4), $\|\beta_{0g}\|_2 \leq \sqrt{G} \|\beta_0\|_\infty \leq C_7 \sqrt{G} \log G$, for some constant $C_7 > 0$. In the third line of the

display, we used same gamma density tail bound that we used to prove (B.11). Using a similar argument as (B.21), we also have

$$\begin{aligned} \sum_{g \in S_0^c} \Pi(\|\beta_g\|_2 > G^M) &\leq \sum_{g \in S_0^c} P(X_g > G^M) \\ &\leq \sum_{g \in S_0^c} \exp[-\lambda_1 G^M + m_{\max}] \\ &= \exp[-\lambda_1 G^M + m_{\max} + \log(G - s_0)]. \end{aligned} \quad (\text{B.22})$$

Combining (B.20)-(B.22), we have that for sufficiently large n ,

$$\begin{aligned} \Pi(\mathcal{A}_n \setminus \mathcal{A}_n^*) &\leq 2 \exp[-\lambda_1 G^M + \lambda_1 C_7 \sqrt{G} \log G + \log G + m_{\max}] \\ &\lesssim \exp(-C_4 s_0 \log G), \end{aligned} \quad (\text{B.23})$$

where we can choose $C_4 > C_3$ so that $-\lambda_1 G^M$ is the dominating term in the first line of the display, and further, $G^M \gg C_4 s_0 n \log G$ when n is large. Thus, from (B.23), we can see that (B.19) holds, and we have established the posterior contraction rate with respect to the average squared Hellinger metric (B.13).

Step 2: Contraction rate for the regression coefficients. We first define the uniform compatibility number ϕ_1 as

$$\phi_1(s; \mathbf{\Omega}) = \inf_{\beta: 1 \leq |\nu(\beta)| \leq s} \frac{\|\mathbf{\Omega}^{1/2} \mathbf{X} \beta\|_2 |\nu(\beta)|^{1/2}}{n^{1/2} \|\beta\|_1},$$

and the smallest scaled singular value ϕ_2 as

$$\phi_2(s; \mathbf{\Omega}) = \inf_{\beta: 1 \leq |\nu(\beta)| \leq s} \frac{\|\mathbf{\Omega}^{1/2} \mathbf{X} \beta\|_2}{n^{1/2} \|\beta\|_2}.$$

According to Theorem 3 of Jeong and Ghosal (2021), as long as

$$\frac{s_0^2 \log G \|\mathbf{X}\|_{\max}}{\phi_1^2((K_2 + 1)s_0; \mathbf{\Omega}(\beta_0))} = o(n), \quad (\text{B.24})$$

then for some $C_8 > 0$,

$$\{\beta : H_n(\beta, \beta_0) \leq C_4 \epsilon_n\} \subset \left\{ \beta : \|\beta - \beta_0\|_2 \leq \frac{C_8 \epsilon_n}{\phi_2((K_2 + 1)s_0; \mathbf{\Omega}(\beta_0))} \right\}, \quad (\text{B.25})$$

where $\epsilon_n = (s_0 \log G/n)^{1/2}$. Notice that for any set S such that $s := |S| \leq (K_2 + 1)s_0$, we have

$$\begin{aligned} \phi_2(s; \mathbf{\Omega}(\beta_0)) &= \inf_{\beta: 1 \leq |\nu(\beta)| \leq s} \frac{\|\mathbf{\Omega}^{1/2}(\beta_0) \mathbf{X} \beta\|_2}{\|\beta\|_2 n^{1/2}} \\ &\geq \frac{[\lambda_{\min}(\mathbf{X}_S^\top \mathbf{\Omega}(\beta_0) \mathbf{X}_S)]^{1/2} \|\beta\|_2}{\|\beta\|_2 n^{1/2}} \\ &= \left[\lambda_{\min} \left(\frac{1}{n} \mathbf{X}_S^\top \mathbf{\Omega}(\beta_0) \mathbf{X}_S \right) \right]^{1/2} \gtrsim 1, \end{aligned} \quad (\text{B.26})$$

by Assumption (B3)(ii). We also have $\|\mathbf{X}\|_{\max} \lesssim \log G$ by Assumption (B3)(i), and $\phi_1(s; \boldsymbol{\Omega}) \geq \phi_2(s; \boldsymbol{\Omega})$. It follows from (B.26) that

$$\frac{s_0^2 \log G \|\mathbf{X}\|_{\max}}{\phi_1^2((K_2 + 1)s_0; \boldsymbol{\Omega}(\boldsymbol{\beta}_0))} \lesssim \frac{s_0^2 (\log G)^2}{\phi_2^2((K_2 + 1)s_0; \boldsymbol{\Omega}(\boldsymbol{\beta}_0))} \lesssim s_0^2 (\log G)^2.$$

Thus, by Assumption (A2) that $s_0 = o(n^{1/2}/\log G)$, it must be that (B.24) holds. Thus, for some $K_3 > 0$, we have by (B.25) and Theorem 2 that

$$\mathbb{E}_0 \Pi(\boldsymbol{\beta} : \|\boldsymbol{\beta} - \boldsymbol{\beta}_0\|_2 \leq K_3 \epsilon_n \mid \mathbf{Y}) \geq \Pi(\boldsymbol{\beta} : H_n(\boldsymbol{\beta}, \boldsymbol{\beta}_0) \leq C_4 \epsilon_n) \rightarrow 1 \text{ as } n \rightarrow \infty,$$

where the smallest scaled singular value $\phi_2((K_2 + 1)s_0; \boldsymbol{\Omega}(\boldsymbol{\beta}_0))$ is a positive constant (by (B.26)) that is absorbed into the constant K_3 . We are done. $\square$

B.3 Proof of Proposition 1

Throughout, we use P_0 and $\mathbb{E}_0$ to respectively denote the probability and expectation operators under the law induced by the model $\mathbf{Y} \sim \mathcal{N}_n(\boldsymbol{\beta}_0, \mathbf{I}_n)$.

The posterior contraction proof in Proposition 1 follows a similar strategy as the proof of Theorem 7 in Castillo et al. (2015). However, it should be noted that Castillo et al. (2015) deal exclusively with the case where $\boldsymbol{\beta}_0 = (\beta_{01}, \dots, \beta_{0n})^\top$ is the zero vector, i.e. $\boldsymbol{\beta}_0 = \mathbf{0}_n$, and the *individual* regression coefficients $\beta_{0i}, i = 1, \dots, n$, are endowed with *univariate* Laplace(0, λ^{-1}) priors. In contrast, we consider the *grouped* scenario where $\boldsymbol{\beta}_0 = (\boldsymbol{\beta}_{01}^\top, \dots, \boldsymbol{\beta}_{0G}^\top)^\top$ with $\boldsymbol{\beta}_{01} \neq \mathbf{0}_{m_1}$, and the groups $\boldsymbol{\beta}_{0g}, g = 1, \dots, G$, are endowed with independent *multivariate* Laplace priors (9).

Proof of Proposition 1 We first prove (25). Let $\mathbf{Y} = (\mathbf{Y}_1^\top, \dots, \mathbf{Y}_G^\top)^\top$, where the indices of $\mathbf{Y}_g$ correspond to those of $\boldsymbol{\beta}_{0g}$ in $\boldsymbol{\beta}_0 = (\boldsymbol{\beta}_{01}^\top, \dots, \boldsymbol{\beta}_{0G}^\top)^\top$ for all $g = 1, \dots, G$. When $\mathbf{Y} \sim \mathcal{N}_n(\boldsymbol{\beta}_0, \mathbf{I}_n)$, the group lasso estimator has a closed form,

$$\hat{\boldsymbol{\beta}}_g = \left(1 - \frac{\lambda}{\|\mathbf{Y}_g\|_2}\right)_+ \mathbf{Y}_g, \text{ for all } g = 1, \dots, G,$$

where $x_+ = \min\{x, 0\}$. Therefore, we must have $\|\hat{\boldsymbol{\beta}}_g\|_2 < \|\mathbf{Y}_g\|_2$ for all $g \in \{1, \dots, G\}$. We can decompose the expected squared ℓ_2 risk as

$$\mathbb{E}_0 \|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}_0\|_2^2 = \mathbb{E}_0 \|\hat{\boldsymbol{\beta}}_1 - \boldsymbol{\beta}_{01}\|_2^2 + \sum_{g=2}^G \mathbb{E}_0 \|\hat{\boldsymbol{\beta}}_g\|_2^2, \quad (\text{B.27})$$

We first upper-bound the first term in (B.27). Since $\mathbf{Y}_1 \sim \mathcal{N}_m(\boldsymbol{\beta}_{01}, \mathbf{I}_m)$, $\|\hat{\boldsymbol{\beta}}_1\|_2 < \|\mathbf{Y}_1\|_2$, and $\hat{\boldsymbol{\beta}}_1 = \mathbf{0}_m$ when $\|\mathbf{Y}_1\|_2 \leq \lambda$, we have

$$\begin{aligned} \mathbb{E}_0 \|\hat{\boldsymbol{\beta}}_1 - \boldsymbol{\beta}_{01}\|_2^2 &\leq 2\mathbb{E}_0 \|\mathbf{Y}_1 - \boldsymbol{\beta}_{01}\|_2^2 + 2\mathbb{E}_0 \|\hat{\boldsymbol{\beta}}_1 - \mathbf{Y}_1\|_2^2 \\ &\leq 2m + 2\mathbb{E}_0 \{\|\mathbf{Y}_1\|_2^2 \mathbb{I}(\|\mathbf{Y}_1\|_2 \leq \lambda)\} + 2\mathbb{E}_0 \{\lambda^2 \mathbb{I}(\|\mathbf{Y}_1\|_2 > \lambda)\} \\ &\leq 2m + 2\lambda^2 + 2\lambda^2 P_0(\|\mathbf{Y}_1\|_2 > \lambda) \\ &\leq 2m + 4\lambda^2 \asymp 2 \log n, \end{aligned} \quad (\text{B.28})$$

since m is constant and $\lambda \asymp \sqrt{2 \log n}$. Next, we upper-bound the second term in (B.27). Again, using the fact that $\|\widehat{\beta}_g\|_2 < \|\mathbf{Y}_g\|_2$ and $\widehat{\beta}_g = \mathbf{0}_m$ when $\|\mathbf{Y}_g\|_2 \leq \lambda$, each of the summands can be upper-bounded as

$$\mathbb{E}_0 \|\widehat{\beta}_g\|_2^2 \leq \mathbb{E}_0 \{\|\mathbf{Y}_g\|_2^2 \mathbb{I}(\|\mathbf{Y}_g\|_2^2 > \lambda^2)\}.$$

Since $\mathbf{Y}_g \sim \mathcal{N}_m(\mathbf{0}_m, \mathbf{I}_m)$ for all $g \neq 1$, $\|\mathbf{Y}_g\|_2^2 \sim \chi_m^2$. Letting $\Gamma(a, x) = \int_x^\infty t^{a-1} e^{-t} dt$ denote the incomplete gamma function, we proceed to upper-bound the individual summands as

$$\begin{aligned} \mathbb{E}_0 \|\widehat{\beta}_g\|_2^2 &\leq \int_{\lambda^2}^\infty \left[\Gamma\left(\frac{m}{2}\right) \right]^{-1} \left(\frac{u}{2}\right)^{m/2} e^{-u/2} du \\ &= 2 \left[\Gamma\left(\frac{m}{2}\right) \right]^{-1} \Gamma\left(\frac{m}{2} + 1, \lambda^2\right) \\ &\leq 2 \left[\Gamma\left(\frac{m}{2}\right) \right]^{-1} \left(\frac{m}{2} + 1\right) e^{-\lambda^2} (\lambda^2)^{m/2} \quad \text{for large } n \\ &\asymp n^{-2} (\log n)^{m/2}, \end{aligned} \tag{B.29}$$

where we used the bound $\Gamma(a, x) \leq a e^{-x} x^{a-1}$ when $a \geq 1$ and $x > a$ (Pinelis, 2020) in the third line, and the fact that m is fixed and $\lambda \asymp \sqrt{2 \log n}$ in the last line. Combining (B.27)-(B.29) gives

$$\mathbb{E}_0 \|\widehat{\beta} - \beta_0\|_2^2 \lesssim 2 \log n + (G-1) n^{-2} (\log n)^{m/2} = 2 \log n + o(1) \quad \text{as } n \rightarrow \infty,$$

where we used the fact that $G = n/m$ and m is a constant. This completes the proof of (25).

Next, we prove (26). We follow the proof strategy of Theorem 7 in Castillo et al. (2015). According to Lemma 7.1 of Castillo and van der Vaart (2012), $\mathbb{E}_0 \Pi(\beta : \|\beta - \beta_0\|_2 < s_n \mid \mathbf{Y}) \rightarrow 0$ for any s_n for which there exists a sequence r_n such that

$$\frac{\Pi(\beta : \|\beta - \beta_0\|_2 < s_n)}{\Pi(\beta : \|\beta - \beta_0\|_2 < r_n)} e^{2r_n^2} = o(1). \tag{B.30}$$

Since the prior is $\beta_g \stackrel{iid}{\sim} \Psi(\beta_g \mid \lambda)$ for all $g = 1, \dots, G$, $\|\beta_1\|_2 \sim \text{Gamma}(m, \lambda)$ and $\sum_{g=1}^G \|\beta_g\|_2 \sim \text{Gamma}(n, \lambda)$ where $n = mG$.

We first upper-bound the numerator in (B.30). Since $\|\beta_0\|_2 = \|\beta_{01}\|_2 = L$ and $\|\beta - \beta_0\|_2 \geq \|\beta_1 - \beta_{01}\|_2$, we have that for sufficiently large n such that $s_n > L$,

$$\begin{aligned} \Pi(\beta : \|\beta - \beta_0\|_2 < s_n) &\leq \Pi(\beta_1 : \|\beta_1 - \beta_{01}\|_2 < s_n) \\ &\leq \Pi(\beta_1 : \left| \|\beta_1\|_2 - \|\beta_{01}\|_2 \right| < s_n) \\ &= \Pi(\beta_1 : \|\beta_1\|_2 < s_n + L) \\ &= \int_0^{s_n+L} \frac{\lambda^m}{\Gamma(m)} u^{m-1} e^{-\lambda u} du \\ &\leq \frac{[\lambda(s_n + L)]^m}{\Gamma(m+1)}. \end{aligned} \tag{B.31}$$

Next, we lower-bound the denominator in (B.30). For sufficiently large n so that $r_n > L$, we have

$$\begin{aligned}
\Pi(\boldsymbol{\beta} : \|\boldsymbol{\beta} - \boldsymbol{\beta}_0\|_2 < r_n) &\geq \Pi(\boldsymbol{\beta} : \|\boldsymbol{\beta}\|_2 + \|\boldsymbol{\beta}_0\|_2 < r_n) \\
&= \Pi(\boldsymbol{\beta} : \|\boldsymbol{\beta}\|_2 < r_n - L) \\
&\geq \Pi\left(\boldsymbol{\beta} : \sum_{g=1}^G \|\boldsymbol{\beta}_g\|_2 < r_n - L\right) \\
&= \int_0^{r_n-L} \frac{\lambda^n}{\Gamma(n)} u^{n-1} e^{-\lambda u} du \\
&\geq \frac{\lambda^n}{n\Gamma(n)} e^{-\lambda(r_n-L)} (r_n - L)^n. \tag{B.32}
\end{aligned}$$

Combining (B.30)-(B.32) gives

$$\begin{aligned}
\frac{\Pi(\boldsymbol{\beta} : \|\boldsymbol{\beta} - \boldsymbol{\beta}_0\|_2 < s_n)}{\Pi(\boldsymbol{\beta} : \|\boldsymbol{\beta} - \boldsymbol{\beta}_0\|_2 < r_n)} &\leq n\lambda^{m-n} \frac{\Gamma(n)}{\Gamma(m+1)} \frac{(s_n + L)^m}{(r_n - L)^n} e^{\lambda(r_n-L)} \\
&\lesssim e^{\lambda\sqrt{n}r_n} \left(\frac{s_n}{r_n}\right)^n \\
&= \exp\left[\lambda\sqrt{n}r_n - n\log\left(\frac{r_n}{s_n}\right)\right]. \tag{B.33}
\end{aligned}$$

Setting $r_n = \sqrt{n}(\lambda^{-1} \wedge 1)$ and $s_n = \delta r_n$ for δ small enough in (B.33) satisfies (B.30), and thus, by Lemma 7.1 of Castillo and van der Vaart (2012),

$$\mathbb{E}_0 \Pi\left(\boldsymbol{\beta} : \|\boldsymbol{\beta} - \boldsymbol{\beta}_0\|_2 < \delta\sqrt{n}\left(\frac{1}{\lambda} \wedge 1\right) \mid \mathbf{Y}\right) \rightarrow 0 \text{ as } n \rightarrow \infty.$$

By setting $\lambda \asymp \sqrt{2\log n}$ in the above display, we obtain the posterior contraction rate of $\sqrt{n/\log n}$ in (26). $\square$

C Additional details for the EM algorithm in Section 5.1

C.1 Derivation of the EM algorithm

Recall that the log-likelihood function $\ell_n(\boldsymbol{\beta})$ is given in (4). With the reparameterization of the hierarchical SSGL prior in (30) and the $\mathcal{B}(a, b)$ prior (27) on θ , we can write the log-posterior $\log \pi(\boldsymbol{\beta}, \theta \mid \mathbf{Y})$ in (28) as

$$\begin{aligned}
\log \pi(\boldsymbol{\beta}, \theta \mid \mathbf{Y}) &= \ell_n(\boldsymbol{\beta}) + \sum_{g=1}^G \log\left((1 - \gamma_g)\lambda_0^{m_g} e^{-\lambda_0\|\boldsymbol{\beta}_g\|_2} + \gamma_g\lambda_1^{m_g} e^{-\lambda_1\|\boldsymbol{\beta}_g\|_2}\right) \\
&\quad + \left(a - 1 + \sum_{g=1}^G \gamma_g\right) \log \theta + \left(b - 1 + p - \sum_{g=1}^G \gamma_g\right) \log(1 - \theta). \tag{C.1}
\end{aligned}$$

From (C.1), it is straightforward to verify that $\mathbb{E}[\gamma_g \mid \mathbf{Y}, \boldsymbol{\beta}, \theta] = p_g^*(\boldsymbol{\beta}_g, \theta)$, where $p_g^*(\boldsymbol{\beta}_g, \theta)$ is as in (31).

In the E-step of our EM algorithm, we treat γ as missing data and take expectation with respect to γ_g for all $g = 1, \dots, G$, holding the parameters $(\boldsymbol{\beta}, \theta)$ fixed at their previous values. That is, we compute $p_g^{*(t-1)} := p^*(\boldsymbol{\beta}_g^{(t-1)}, \theta^{(t-1)}) = \mathbb{E}[\gamma_g \mid \mathbf{Y}, \boldsymbol{\beta}^{(t-1)}, \theta^{(t-1)}]$ for all $g = 1, \dots, G$, given the previous estimates $(\boldsymbol{\beta}^{(t-1)}, \theta^{(t-1)})$.

For the M-step, we then maximize the following function:

$$\begin{aligned} \mathbb{E}[\log \pi(\boldsymbol{\beta}, \theta \mid \mathbf{Y}) \mid \boldsymbol{\beta}^{(t-1)}, \theta^{(t-1)}] &= \ell_n(\boldsymbol{\beta}) - \sum_{g=1}^G \lambda_g^{*(t-1)} \|\boldsymbol{\beta}_g\|_2 \\ &+ \left(a - 1 + \sum_{g=1}^G p_g^{*(t-1)} \right) \log \theta + \left(b - 1 + G - \sum_{g=1}^G p_g^{*(t-1)} \right) \log(1 - \theta), \end{aligned} \quad (\text{C.2})$$

where $\lambda_g^{*(t-1)} = \lambda_1 p_g^{*(t-1)} + \lambda_0(1 - p_g^{*(t-1)})$. From (C.2), it is easy to derive the update for θ in (32) by taking the derivative of (C.2) with respect to θ . The update for $\boldsymbol{\beta}$ in (33) is obtained by isolating the terms on the right-hand side of (C.2) that only depend on $\boldsymbol{\beta}$. We discuss the M-step for updating $\boldsymbol{\beta}$ in detail in the next section.

C.2 M-step for updating the regression coefficients

In the M-step of our EM algorithm, we need to solve the optimization problem (33), or equivalently,

$$\boldsymbol{\beta}^{(t)} = \arg \min_{\boldsymbol{\beta}} \left\{ -\ell_n(\boldsymbol{\beta}) + \sum_{g=1}^G \lambda_g^{*(t-1)} \|\boldsymbol{\beta}_g\|_2 \right\}, \quad (\text{C.3})$$

where $-\ell_n(\boldsymbol{\beta})$ is the negative of the log-likelihood in (4). While (C.3) may appear to be intractable, we can apply the standard iteratively reweighted least squares (IRLS) algorithm (McCullagh and Nelder, 1989) for GLMs to efficiently solve (C.3).

The IRLS algorithm in GLMs is based on a quadratic approximation of the negative log-likelihood. Denote the mean response as $\mu_i = b'(\theta_i)$, and let $V(\mu_i) = b''((b')^{-1}(\mu_i))$ be the variance function, where b is the cumulant function in (1). With the link function h in (3), we define the “working response” vector $\mathbf{Z}$ at the k th iteration of the optimization problem (C.3) as $\mathbf{Z} = \mathbf{X}\boldsymbol{\beta}^{(k)} + \boldsymbol{\zeta}^{(k)}$, where $\zeta_i^{(k)} = h'(\mu_i^{(k)})(y_i - \mu_i^{(k)})$, $i = 1, \dots, n$. Similarly, we define the weights matrix $\mathbf{W} = \text{diag}(w_1^{(k)}, \dots, w_n^{(k)})$, where the weights are $w_i^{(k)} = [V(\mu_i^{(k)})(h'(\mu_i^{(k)}))^2]^{-1}$. As shown in McCullagh and Nelder (1989), the negative log-likelihood for any GLM in the exponential dispersion family (1.1) can then be approximated as

$$-\ell_n(\boldsymbol{\beta}) \approx \frac{1}{2}(\mathbf{Z} - \mathbf{X}\boldsymbol{\beta})^\top \mathbf{W}(\mathbf{Z} - \mathbf{X}\boldsymbol{\beta}). \quad (\text{C.4})$$

Thus, substituting the approximation (C.4) for $-\ell_n(\boldsymbol{\beta})$ in (C.3), we see that the M-step (C.3) for $\boldsymbol{\beta}$ can alternatively be written as

$$\boldsymbol{\beta}^{(t)} = \arg \min_{\boldsymbol{\beta}} \left\{ \frac{1}{2}(\mathbf{Z} - \mathbf{X}\boldsymbol{\beta})^\top \mathbf{W}(\mathbf{Z} - \mathbf{X}\boldsymbol{\beta}) + \sum_{g=1}^G \lambda_g^{*(t-1)} \|\boldsymbol{\beta}_g\|_2 \right\}. \quad (\text{C.5})$$

With the quadratic approximation to the negative log-likelihood, (C.5) is now a standard linear regression model with a group lasso penalty (Yuan and Lin, 2006) and group-specific weights $\lambda_g^{*(t-1)}$ for each g th group. Solving (C.5) can be done with any standard block coordinate descent algorithm for the regular group lasso (Yuan and Lin, 2006) in linear regression. In particular, if the canonical link function $h = (b')^{-1}$ is used, then the Fisher information matrix and the negative Hessian matrix are equal, and we use the majorization-minimization (MM) algorithm in Breheny and Huang (2015) to solve (C.5). On the other hand, if a non-canonical link function is used, then we use the least squares approximation (LSA) approach of Wang and Leng (2007) with a group lasso penalty.

The SSGL MAP algorithm is on average slower than the numerical algorithms used for group lasso, group SCAD, or group MCP. This is because *each* iteration of the EM algorithm requires numerically solve the optimization problem (C.3) in the M-step. Therefore, if the EM algorithm takes B total iterations to converge, we expect the average runtime of the SSGL method to be on the order of B times slower than that for group lasso, group SCAD, and group MCP. However, we observed that our EM algorithm tended to converge within a small number of iterations (typically fewer than 10), at least in all of the simulation settings we considered. Thus, the computation time for SSGL for fixed λ_0 was still reasonably quick. For example, in Simulations 4 and 8 of the main manuscript where $p = 1600$, the SSGL runtime for any fixed λ_0 was under 30 seconds in all replications.

To further accelerate the computational efficiency for SSGL, we could employ alternative optimization strategies that obviate the need to introduce latent indicator variables. For example, we could potentially approximate the SSGL prior with a neuronized prior (Shin and Liu, 2022), which is a product of a Gaussian weight variable and a scale variable transformed from Gaussian via an activation function. This approximation would allow for a fast coordinate ascent algorithm which avoids the need to marginalize out latent indicator variables (Shin and Liu, 2022). We leave the investigation of more computationally efficient algorithms for SSGL for future work.

D Complete Gibbs sampling algorithms for Section 5.3

Recall that $\boldsymbol{\beta} = (\boldsymbol{\beta}_1^\top, \dots, \boldsymbol{\beta}_G^\top)^\top$, where $\boldsymbol{\beta}_g \in \mathbb{R}^{m_g}$, $g = 1, \dots, G$, and $\boldsymbol{\tau} = (\tau_1, \dots, \tau_G)^\top$ is defined in (5.14). Let $\mathcal{PG}(a, b)$ denote the Pólya-gamma distribution (Polson et al., 2013), and let $\mathcal{GIG}(a, b, c)$ denote the generalized inverse

Gaussian distribution with density function $f(x; a, b, c) \propto x^{a-1} e^{-(b/x+cx)/2}$. We define $\boldsymbol{\kappa} = (\kappa_1, \dots, \kappa_n)^\top$, where for $i = 1, \dots, n$,

$$\kappa_i = \begin{cases} y_i - 0.5 & \text{for logistic regression,} \\ 0.5(y_i - M) & \text{for Poisson regression.} \end{cases} \quad (\text{D.1})$$

We also define $\boldsymbol{\omega} = (\omega_1, \dots, \omega_n)^\top$, where for $i = 1, \dots, n$,

$$\begin{cases} \omega_i \sim \mathcal{PG}(1, \mathbf{x}_i^\top \boldsymbol{\beta}) & \text{for logistic regression,} \\ \omega_i \sim \mathcal{PG}(M, \mathbf{x}_i^\top \boldsymbol{\beta} - \log M) & \text{for Poisson regression.} \end{cases} \quad (\text{D.2})$$

Finally, define the vector $\mathbf{Z} = (z_1, \dots, z_n)^\top$, where for $i = 1, \dots, n$,

$$z_i = \begin{cases} \kappa_i / \omega_i & \text{for logistic regression,} \\ \kappa_i / \omega_i + \log M & \text{for Poisson regression,} \end{cases} \quad (\text{D.3})$$

and $\boldsymbol{\kappa}$ is defined as in (D.1).

The Gibbs sampler for grouped logistic or Poisson regression with the SSGL prior (5.4) is given in Algorithm 1.

When $p = \sum_{g=1}^G m_g$ is larger than n , we can use the fast sampling algorithm of Bhattacharya et al. (2016) to sample $\boldsymbol{\beta}^{(t)}$ in Step 4 of Algorithm 1 in $\mathcal{O}(n^2 p)$ time instead of $\mathcal{O}(p^3)$ time. This sampling algorithm is given in Algorithm 2.

E Simulation studies for negative binomial regression with a log link

Since the mean and variance are equal for the Poisson distribution, Poisson regression may be inappropriate for modeling count data that exhibit overdispersion. In this case, it may be more appropriate to use negative binomial regression. That is, we assume that $y_i \sim NB(\alpha, \mu_i)$, $i = 1, \dots, n$, where $\alpha > 0$ is a size parameter, and the probability density function (pdf) for each y_i is

$$f(y_i | \alpha, \mu_i) = \frac{\Gamma(y_i + \alpha)}{y_i! \Gamma(\alpha)} \left(\frac{\mu_i}{\mu_i + \alpha} \right)^{y_i} \left(\frac{\alpha}{\mu_i + \alpha} \right)^\alpha, \quad y_i = 0, 1, 2, \dots$$

Since $\text{Var}(y_i) = \mu_i + \mu_i^2/\alpha$, the negative binomial distribution is better equipped to handle overdispersed count data than the Poisson distribution.

For negative binomial distributed responses, the natural parameter is $\theta_i = \log(\mu_i/(\mu_i + \alpha))$. With the link function $h(u) = \log(u)$, we have $b(u) = -\alpha \log(1 - e^u)$ in (1). It is readily seen that $\mathbb{E}(y_i) = b'(\theta_i) = \mu_i$. For the function $\xi = (h \circ b')^{-1}$, we have $\xi(u) = -\log(\alpha e^{-u} + 1)$. Since $\xi(u) \neq u$, negative binomial regression with the log link is an example of a GLM with a *non*-canonical link function.

We now present simulation results in grouped negative binomial regression for the SSGL model (10). We compared our results to group lasso (GL), adaptive group lasso (AdGL), group SCAD (GSCAD), and group MCP (gMCP). We used

Algorithm 1 MCMC algorithm for grouped logistic and grouped Poisson regression with the SSGL prior

Input: Initial values $\beta^{(0)}$, $\theta^{(0)}$, $\tau^{(0)}$, T (number of MCMC samples to run), B (number of samples to discard as burn-in)

Output: MCMC samples for regression coefficients β

for $t = 1, \dots, T$ **do**

1. **for** $i = 1, \dots, n$ **do**

- i. Sample $\omega^{(t)}$ as in (D.2) and set $\Omega = \text{diag}\{\omega^{(t)}\}$
- ii. Update $\mathbf{Z}^{(t)}$ as in (D.3)

2. **for** $g = 1, \dots, G$ **do**

- i. Sample $\gamma_g^{(t)} \sim \text{Bernoulli}\left(\frac{\pi_1}{\pi_1 + \pi_0}\right)$, where
 $\pi_1 = \theta^{(t-1)} \lambda_1^{m_g+1} \exp\left(-\lambda_1^2 \tau_g^{(t-1)} / 2\right)$
and $\pi_0 = (1 - \theta^{(t-1)}) \lambda_0^{m_g+1} \exp\left(-\lambda_0^2 \tau_g^{(t-1)} / 2\right)$
- ii. Sample $\tau_g^{(t)} \sim \mathcal{GIG}\left(\frac{1}{2}, \|\beta_g^{(t-1)}\|_2^2, (\lambda_g^*)^2\right)$, where
 $\lambda_g^* = \gamma_g^{(t)} \lambda_1 + (1 - \gamma_g^{(t)}) \lambda_0$

3. Sample $\theta^{(t)} \sim \text{Beta}\left(a + \sum_{g=1}^G \gamma_g^{(t)}, b + G - \sum_{g=1}^G \tau_g^{(t)}\right)$

4. Sample $\beta^{(t)} \sim \mathcal{N}(\mu_\beta, \Sigma_\beta)$, where $\mu_\beta = \Sigma_\beta \mathbf{X}^\top \Omega \mathbf{Z}$,
 $\Sigma_\beta = (\mathbf{X}^\top \Omega \mathbf{X} + \mathbf{D}_\tau^{-1})^{-1}$, and $\mathbf{D}_\tau = \text{Bdiag}(\tau_1^{(t)} \mathbf{I}_{m_1}, \dots, \tau_G^{(t)} \mathbf{I}_{m_G})$

return $\beta^{(B+1)}, \dots, \beta^{(T)}$

the same methods for choosing the hyperparameters, regularization parameters, and weights as those described in Section 6.1 of the main article. For our simulation study, we fixed $\alpha = 1$ and generated the responses independently as $y_i \mid \mathbf{x}_i \sim NB(\alpha, \exp(\theta_i))$, $i = 1, \dots, n$. We conducted the following simulation studies:

Experiment 1. We set $n = 500$ and $G = 30$. The design matrix $\mathbf{X}$ and the regression coefficients β were generated the same way as in Experiment 1 in Section 6.3 of the main manuscript. Then we modeled

$$\log(\theta) = \mathbf{x}^\top \beta.$$

Experiment 2. We set $n = 500$ and $G = 30$. We generated the entries of the $n \times G$ design matrix $\mathbf{X}$ from independent $\text{Uniform}(-1, 1)$ random

Algorithm 2 Algorithm for sampling β in Step 4 of Algorithm 1 when $p > n$

Input: Most recent updates of Ω , D_τ , and $\mathbf{Z}$

Output: An exact sample of β in Step 4 of Algorithm 1

1. Set $\tilde{\mathbf{X}} = \Omega^{1/2}\mathbf{X}$
2. Sample $\mathbf{m} \sim \mathcal{N}(\mathbf{0}, D_\tau)$ and $\delta \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$ independently
3. Set $\mathbf{v} = \tilde{\mathbf{X}}\mathbf{m} + \delta$ and $\mathbf{v}^* = \Omega^{1/2}\mathbf{Z} - \mathbf{v}$
4. Set $\mathbf{A} = \tilde{\mathbf{X}}D_\tau\tilde{\mathbf{X}}^\top + \mathbf{I}_n$
5. Set $\mathbf{w} = \mathbf{A}^{-1}\mathbf{v}^*$
6. Set $\beta = \mathbf{m} + D_\tau\tilde{\mathbf{X}}^\top\mathbf{w}$

return β

variables. Then we modeled

$$\log(\theta) = 1.5 \sin(3x_1) - x_5 e^{0.5x_5^2}.$$

We represented each covariate as a six-term B-spline basis expansion.

We repeated each experiment 200 times. Table 1 reports our results averaged across the 200 replications. In Experiment 1, SSGL had the lowest MSE and the second lowest MSPE. Meanwhile, all five methods had a TPR of 1, indicating that all methods had equally high power to detect the significant groups in this particular setting. However, SSGL had the highest average TNR and precision, indicating superior overall group selection performance.

In Experiment 2, SSGL also had the lowest MSPE, indicating superior average prediction on out-of-sample data. In this experiment, all five methods once again had a TPR of 1. However, SSGL had the highest average TNR and precision, indicating that SSGL was the best suited for function selection in this particular generalized additive model.

References

- Bai, R., Moran, G. E., Antonelli, J. L., Chen, Y., and Boland, M. R. (2022). Spike-and-slab group lassos for grouped regression and sparse generalized additive models. *Journal of the American Statistical Association*, 117(537):184–197.
- Bhattacharya, A., Chakraborty, A., and Mallick, B. K. (2016). Fast sampling with Gaussian scale mixture priors in high-dimensional regression. *Biometrika*, 103(4):985–991.

Table 1: Simulation results for grouped negative binomial regression under the SSGL, GL, AdGL, GMCP, and GSCAD models, averaged across 200 replicates. The empirical standard error is reported in parentheses below the average.

Experiment 1					
	MSE	MSPE	TPR	TNR	Prec
SSGL	0.008 (0.006)	30.04 (10.56)	1 (0)	0.999 (0.007)	0.999 (0.029)
GL	0.029 (0.008)	29.43 (9.52)	1 (0)	0.387 (0.146)	0.255 (0.051)
AdGL	0.018 (0.008)	31.09 (12.32)	1 (0)	0.809 (0.119)	0.562 (0.180)
GMCP	0.012 (0.005)	31.15 (15.98)	1 (0)	0.892 (0.094)	0.705 (0.192)
GSCAD	0.013 (0.005)	31.31 (14.59)	1 (0)	0.702 (0.133)	0.434 (0.130)

Experiment 2				
	MSPE	TPR	TNR	Prec
SSGL	17.43 (5.09)	1 (0)	1 (0)	1 (0)
GL	18.35 (5.85)	1 (0)	0.383 (0.184)	0.114 (0.044)
AdGL	18.16 (5.47)	1 (0)	0.759 (0.140)	0.296 (0.186)
GMCP	17.52 (5.10)	1 (0)	0.921 (0.123)	0.695 (0.317)
GSCAD	17.49 (5.33)	1 (0)	0.790 (0.160)	0.363 (0.236)

Boucheron, S., Lugosi, G., and Massart, P. (2013). *Concentration Inequalities: A Nonasymptotic Theory of Independence*. Oxford University Press.

Breheny, P. and Huang, J. (2015). Group descent algorithms for nonconvex penalized linear and logistic regression models with grouped predictors. *Statistics and Computing*, 25(2):173–187.

Castillo, I., Schmidt-Hieber, J., and van der Vaart, A. (2015). Bayesian linear regression with sparse priors. *The Annals of Statistics*, 43(5):1986 – 2018.

Castillo, I. and van der Vaart, A. (2012). Needles and straw in a haystack: Posterior concentration for possibly sparse sequences. *The Annals of Statistics*, 40(4):2069 – 2101.

- Fan, J. and Lv, J. (2011). Nonconcave penalized likelihood with NP-dimensionality. *IEEE Transactions on Information Theory*, 57(8):5467–5484.
- Ghosal, S. and van der Vaart, A. (2007). Convergence rates of posterior distributions for noniid observations. *The Annals of Statistics*, 35(1):192 – 223.
- Jeong, S. and Ghosal, S. (2021). Posterior contraction in sparse generalized linear models. *Biometrika*, 108(2):367–379.
- McCullagh, P. and Nelder, J. (1989). *Generalized Linear Models, Second Edition*. Chapman & Hall, New York.
- Ning, B., Jeong, S., and Ghosal, S. (2020). Bayesian linear regression for multivariate responses under group sparsity. *Bernoulli*, 26(3):2353 – 2382.
- Pinelis, I. (2020). Exact lower and upper bounds on the incomplete gamma function. *Mathematical Inequalities & Applications*, 23(4):1261–1278.
- Polson, N. G., Scott, J. G., and Windle, J. (2013). Bayesian inference for logistic models using Pólya–gamma latent variables. *Journal of the American Statistical Association*, 108(504):1339–1349.
- Shen, Y., Solís-Lemus, C., and Deshpande, S. K. (2024). Estimating sparse direct effects in multivariate regression with the spike-and-slab lasso. *Bayesian Analysis (to appear)*.
- Shin, M. and Liu, J. S. (2022). Neuronized priors for Bayesian sparse linear regression. *Journal of the American Statistical Association*, 117(540):1695–1710.
- van der Vaart, A. W. and Wellner, J. A. (2006). *Weak Convergence and Empirical Process: With Applications to Statistics*. Springer.
- Wang, H. and Leng, C. (2007). Unified lasso estimation by least squares approximation. *Journal of the American Statistical Association*, 102:1039–1048.
- Yuan, M. and Lin, Y. (2006). Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 68(1):49–67.