

The graphical horseshoe estimator for inverse covariance matrices

Anindya Bhadra

www.stat.purdue.edu/~bhadra

Purdue University

Overview

- Goal: To develop a new Bayes estimator for sparse precision matrices for multivariate Gaussian data.
- A full Gibbs sampler for efficient sampling.
- Theoretical contrast with popular existing methods such as the graphical lasso and graphical SCAD.
- Numerical examples.
- *Joint work with Yunfan Li and Bruce Craig at Purdue. Supported by NSF Grant DMS-1613063.*

Some existing estimators

- Estimators that assume an unstructured precision matrix usually make an assumption of sparsity.
 - Frequentist: graphical lasso (Friedman et al. 2008, Biostatistics); graphical SCAD (Lam and Fan, 2009, AoS).
 - Bayesian: the Bayesian graphical lasso (Wang, 2012, BA); **the proposed graphical horseshoe estimator**.
- Estimators that assume an underlying structure (banding, latent factors, low rank)
 - Frequentist: Bickel and Levina (2008, AoS) and many others
 - Bayesian: Banerjee and Ghosal (2014, EJS (banded)); Pati et al. (2014, AoS (sparse latent factors))

Some existing estimators

- Let Ω be $p \times p$ positive definite. Observe i.i.d.

$$\mathbf{y}_k \sim \text{Normal}(\mathbf{0}, \Omega^{-1}),$$

for $k = 1, \dots, n$ and let $p > n$.

- Let $S = \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i'$. Penalized likelihood approaches maximize

$$\log(\det \Omega) - \text{tr}(S\Omega/n) - \sum_{i,j} \phi_\lambda(|\omega_{ij}|),$$

- Graphical lasso: $\phi_\lambda(|x|) = \lambda|x|$; $\lambda > 0$.
- Graphical SCAD:

$$\phi'_\lambda(|x|) = \lambda \left\{ 1_{\{|x| \leq \lambda\}} + \frac{(a\lambda - |x|)_+}{(a-1)\lambda} 1_{\{|x| > \lambda\}} \right\}; \quad \lambda > 0, a > 2.$$

Some existing estimators

- Wang (2012, BA) showed the frequentist glasso estimate is posterior mode under the prior

$$p(\Omega \mid \lambda) \propto \prod_{i < j} \{\text{DE}(\omega_{ij} \mid \lambda)\} \prod_{i=1}^p \{\text{EXP}(\omega_{ii} \mid \lambda/2)\} 1_{\Omega \in \mathcal{S}_p},$$

where $\mathcal{S}_p$ is the set of positive definite matrices.

- He also developed a block Gibbs sampler for a full Bayes solution - a strategy we will closely follow.
- The posterior mean estimate under this prior is known as the Bayesian graphical lasso (BGL).

Issues with the existing estimators

- Graphical lasso introduces bias in estimating large signals due to soft thresholding.
- Graphical SCAD penalty is non-convex and unbiased for large signals, but the estimate is not guaranteed to be positive definite in finite samples.
- We will also argue neither graphical lasso nor graphical SCAD provides strong enough shrinkage towards zero, resulting in poorer information-theoretic properties.

Our proposal

- With the constraint $\Omega \in S_p$, define element-wise ($i, j = 1, \dots, p$)

$$\begin{aligned}\omega_{ij} &\propto 1, \\ \omega_{ij:i < j} &\sim \text{Normal}(0, \lambda_{ij}^2 \tau^2), \\ \lambda_{ij:i < j} &\sim C^+(0, 1),\end{aligned}$$

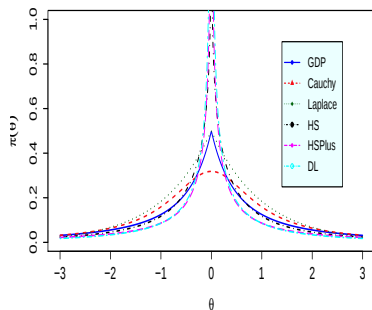
- That is, the prior on Ω is

$$p(\Omega \mid \tau) \propto \prod_{i < j} \text{Normal}(\omega_{ij} \mid 0, \lambda_{ij}^2 \tau^2) \prod_{i < j} C^+(\lambda_{ij} \mid 0, 1) 1_{\Omega \in S_p},$$

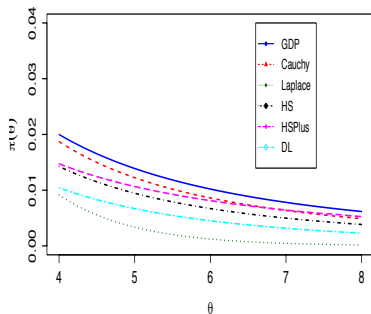
- A non-informative prior on the diagonal terms and independent global-local **horseshoe priors** on off-diagonal terms to ensure (i) efficient shrinkage of noise terms and (ii) not shrinking the signals.

Some examples of global-local priors

- The order of peakedness near zero: $HS+ \approx DL > HS > GDP = Laplace > Cauchy$



- The order of tail heaviness: $GDP > Cauchy > HS+ > HS > DL > Laplace$



Posterior simulation

- The posterior is

$$p(\Omega \mid Y, \Lambda, \tau) \propto |\Omega|^{\frac{n}{2}} \exp\left\{-\operatorname{tr}\left(\frac{1}{2}S\Omega\right)\right\} \prod_{i < j} \exp\left(-\frac{\omega_{ij}^2}{2\lambda_{ij}^2\tau^2}\right) 1_{\Omega \in \mathcal{S}_p}.$$

- Write

$$\Omega = \begin{pmatrix} \Omega_{(-p)(-p)} & \omega_{(-p)p} \\ \omega'_{(-p)p} & \omega_{pp} \end{pmatrix}, \quad S = \begin{pmatrix} S_{(-p)(-p)} & \mathbf{s}_{(-p)p} \\ \mathbf{s}'_{(-p)p} & s_{pp} \end{pmatrix}.$$

- Define $\beta = \omega_{(-p)p}$ and $\gamma = \omega_{pp} - \omega'_{(-p)p} \Omega_{(-p)(-p)}^{-1} \omega_{(-p)p}$.

Posterior simulation

- Wang (2012, BA) showed

$$p(\gamma, \beta \mid \Omega_{(-p)(-p)}, Y, \Lambda, \tau) \sim \text{Gamma}(n/2 + 1, s_{pp}/2) \times \text{Normal}(-C\mathbf{s}_{(-p)p}, C),$$

$$\text{where } C = \{s_{pp}\Omega_{(-p)(-p)}^{-1} + (\Lambda^*\tau^2)^{-1}\}^{-1}.$$

- Conditional posteriors of off-diagonal terms are normal, those of diagonal terms are inverse gamma.
- If the initial Ω is positive definite, ensures all subsequent iterations are also positive definite.

Posterior simulation

- The difference is, for BGL, $\lambda_{ij} \sim \text{Exp}(1)$ but for GHS $\lambda_{ij} \sim C^+(0, 1)$, a priori.
- Here, we follow the key data augmentation technique proposed by Makalic and Schmidt (2016, IEEE Sig. Proc. Letters):

$$\begin{array}{ll} \text{if } x^2 \mid a \sim \text{InvGamma}(1/2, 1/a) & \text{and } a \sim \text{InvGamma}(1/2, 1), \\ \text{then marginally, } x \sim & C^+(0, 1) \end{array}$$

- That is, a half-Cauchy is a mixture of two inverse gammas.
- BUT! Inverse gamma is conjugate to itself and to the variance parameter in a normal linear regression model.

Posterior simulation

- All conditional posteriors are available in closed form, leading to a full Gibbs sampler.
- They are either multivariate normal, gamma or inverse gamma.
- Computational complexity is $O(p^3)$.
- For full details, see Algorithm 1 of [Li, Craig and Bhadra \(2017, arXiv: 1707.06661\)](#).
- MATLAB code on github at <http://github.com/liyf1988/GHS>.

A useful lemma (Li, Craig and Bhadra, 2017)

- Modification of Barron and Clarke (1990, IEEE Trans. Inf. Theory).
- Let the true parameter be Ω_0 and $A_\epsilon = \{\Omega : D(p_{\Omega_0} || p_\Omega) \leq \epsilon\}$.
- Let $\nu(d\Omega)$ be the prior measure of Ω and $\nu_n(d\Omega) \propto \prod_{i=1}^n p_\Omega(y_i) \nu(d\Omega)$ be the posterior measure.
- $\hat{p}_n = \int p_\Omega \nu_n(d\Omega)$ be the posterior mean estimate of the density.
- The Cesàro-average risk R_n of the estimator $\hat{p}_n$ admits the following lower and upper bounds for all $\epsilon > 0$:

$$-\epsilon - \frac{1}{n} \log \nu(A_\epsilon) \leq R_n = \frac{1}{n} \sum_{j=1}^n \mathbb{E} D(p_{\Omega_0} || \hat{p}_j) \leq \epsilon - \frac{1}{n} \log \nu(A_\epsilon),$$

where $\mathbb{E}$ denotes an expectation with respect to the data distribution.

K-L risk bounds

- If we take $\epsilon = 1/n$, then R_n is a function of two things: (i) the sample size n and (ii) $\nu(A_{1/n})$, the prior measure of K-L information neighborhood of true Ω_0 .
- BUT! Recall the horseshoe density is unbounded at zero. When the true parameter is zero, it places **more mass** in the K-L neighborhood than **any prior that is continuous and bounded above**.
- This immediately includes the double exponential prior in BGL.
- The graphical SCAD estimate does not admit an interpretation as a posterior mode similar to BGL, but if we view the corresponding penalty as negative of the log of prior, it's easy to see that prior will be bounded above.

K-L risk bounds (Thm. 3.2; Li, Craig and Bhadra, 2017)

- For $\hat{p}_n$ under the graphical horseshoe prior,
$$p_0 \log \left\{ \frac{C_1}{Mn^{1/2}p} \log(2Mn^{1/2}p) \right\} + p_1 \log \frac{C_2}{n^{1/2}p} < \log \nu(A_{1/n}) < \\ p_0 \log \left\{ \frac{C_3}{Mn^{1/2}p} \log(2^{1/2}Mn^{1/2}p) \right\} + p_1 \log \frac{C_4}{n^{1/2}p},$$
 where p_0 is the number of zero elements in Ω_0 , p_1 is the number of nonzero elements in Ω_0 , and C_1, C_2, C_3, C_4 are constants.
- Suppose $p(\omega_{ij})$ is any other prior density that is continuous, bounded above, and strictly positive on a neighborhood of the true value ω_{ij0} . Then $p^2 \log \frac{K_1}{n^{1/2}p} < \log \nu(A_{1/n}) < p^2 \log \frac{K_2}{n^{1/2}p}$, where K_1 and K_2 are constants.

K-L risk bounds: summary

- The bounds for all methods diverge when $p^2/n \rightarrow \infty$.
- But the upper and lower bounds for GHS diverge slower when true Ω is sparse.
- For finite n and finite $p > n$, we see a nontrivial improvement over BGL and graphical SCAD.

- We are also able to show the graphical horseshoe estimate is nearly unbiased in estimating large signals.
- This is at a contrast with the BGL, which will leave a constant bias regardless of the signal strength due to soft thresholding.
- These results mirror the findings of lasso vs. horseshoe in the normal means model.

Numerical examples

We consider the following structures for true Ω :

- *Random*. A sparse matrix is generated. Then each off-diagonal element is randomly assigned a sign.
- *Hubs*. The rows/columns are partitioned into disjoint groups $\{G_k\}_1^K$. Within each group a sparse structure is generated, but no connections between groups.
- *Cliques*. A sparse decomposable graph that admits clique-separator decomposition.

Numerical examples ($p = 100, n = 50$)

nonzero pairs nonzero elem.	Random 35/4950 $\sim -\text{Unif}(0.2, 1)$					Hubs 90/4950 0.25				
	GL1	GL2	GSCAD	BGL	GHS	GL1	GL2	GSCAD	BGL	GHS
Stein's loss	10.20 (0.53)	13.42 (1.06)	10.05 (0.55)	80.92 (1.63)	6.44 (0.85)	10.12 (0.53)	12.78 (0.96)	10.01 (0.50)	77.85 (1.66)	12.56 (1.04)
F norm	4.33 (0.18)	5.30 (0.20)	4.31 (0.16)	5.58 (0.26)	3.31 (0.29)	3.95 (0.13)	4.63 (0.18)	3.94 (0.14)	5.97 (0.30)	3.96 (0.27)
TPR	.8246 (.0520)	.7097 (.0620)	.9977 (.0078)	.8709 (.0470)	.5903 (.0537)	.8649 (.0443)	.7333 (.0751)	.9987 (.0053)	.8513 (.0378)	.2687 (.0764)
FPR	.0947 (.0141)	.0374 (.0070)	.9955 (.0102)	.1055 (.0059)	.0004 (.0003)	.0919 (.0130)	.0281 (.0086)	.9976 (.0069)	.1189 (.0058)	.0013 (.0005)
Avg CPU time	0.30	0.35	6.24	40.94	38.32	0.14	0.16	4.01	35.44	41.58
nonzero pairs nonzero elem.	Cliques positive 30/4950 -0.45					Cliques negative 30/4950 0.75				
	GL1	GL2	GSCAD	BGL	GHS	GL1	GL2	GSCAD	BGL	GHS
Stein's loss	9.16 (0.55)	14.16 (1.06)	8.99 (0.52)	81.58 (2.51)	5.87 (0.93)	11.00 (0.43)	14.37 (1.02)	10.90 (0.43)	81.27 (1.98)	6.28 (1.09)
F norm	3.75 (0.16)	5.01 (0.16)	3.71 (0.17)	5.44 (0.33)	3.81 (0.41)	6.00 (0.14)	6.86 (0.16)	5.99 (0.14)	6.51 (0.20)	3.64 (0.36)
TPR	1 (0)	1 (0)	1 (0)	1 (0)	.7487 (.0427)	.9993 (.0047)	.9880 (.0221)	1 (0)	.9993 (.0047)	.9733 (.0421)
FPR	.0900 (.0098)	.0255 (.0056)	.9901 (.0177)	.1014 (.0052)	.0003 (.0003)	.0922 (.0135)	.0279 (.0084)	.9752 (.0219)	.1161 (.0051)	.0010 (.0005)
Avg CPU time	0.24	0.28	4.52	34.45	41.65	0.18	0.20	6.91	33.88	41.05

Summary and conclusions

- The proposed estimator performs better than competing methods in terms of Stein's loss (a finite sample estimate of the K-L risk). Also performs well in terms of estimation and variable selection.
- The full Gibbs sampler helps avoiding issues such as tuning an M-H sampler, notoriously difficult in high dimensions.
- We have provided some preliminary theory, but the rich theory developed for horseshoe priors in the normal means model is waiting to be applied here!
- Preprint: [arXiv: 1707.06661](https://arxiv.org/abs/1707.06661); Code: <http://github.com/liyf1988/GHS>