

Fast and Separable Estimation in High-dimensional Tensor Gaussian Graphical Models

Keqian Min, Qing Mai and Xin Zhang*
Florida State University

Abstract

In the tensor data analysis, the Kronecker covariance structure plays a vital role in unsupervised learning and regression. Under the Kronecker covariance model assumption, the covariance of an M -way tensor is parameterized as the Kronecker product of M individual covariance matrices. With normally distributed tensors, the key to high-dimensional tensor graphical models becomes the sparse estimation of the M inverse covariance matrices. Unable to maximize the tensor normal likelihood analytically, existing approaches often require cyclic updates of the M sparse matrices. For the high-dimensional tensor graphical models, each update step solves a regularized inverse covariance estimation problem that is computationally nontrivial. This computational challenge motivates our study of whether a non-cyclic approach can be as good as the cyclic algorithms in theory and practice. To handle the potentially very high-dimensional and high-order tensors, we propose a separable and parallel estimation scheme. We show that the new estimator achieves the same minimax optimal convergence rate as the cyclic estimation approaches. Numerically, the new estimator is much faster and often more accurate than the cyclic approach. Moreover, another advantage of the separable estimation scheme is its flexibility in modeling, where we can easily incorporate user-specified or specially structured covariances on any modes of the tensor. We demonstrate the efficiency of the proposed method through both simulations and a neuroimaging application. Supplementary materials are available online.

Keywords: Graphical Models; Kronecker Covariance; Sparse Precision Matrix; Tensor.

*Correspondence: Xin Zhang (xzhang8@fsu.edu). The authors would like to thank the Editor, the Associate Editor and two reviewers for helpful comments. Research for this paper was supported in part by grants CCF-1908969 from the National Science Foundation.

1 Introduction

Modeling high-dimensional tensor data, i.e., multi-way array-valued variable $\mathbf{X}_i \in \mathbb{R}^{p_1 \times \cdots \times p_M}$ for each sample $i = 1, \dots, n$, are frequently involved in many areas of research such as neuroimaging, computational biology, and signal processing. Statistical methods for tensor data analysis are becoming increasingly popular in recent years (see Chi & Kolda 2012, Zhou et al. 2013, Hoff 2015, Lock 2018, Pan et al. 2019, Pfeiffer et al. 2020, Bi et al. 2020, and among others). In this paper, we study the tensor Gaussian graphical models, which are valuable tools to reveal the dependence structure among the large number of random variables in a tensor. The tensor Gaussian graphical models are generalizations of the Gaussian graphical model for multivariate data. For background on graphical models and the conditional independence interpretation of the Gaussian graphical models, see, for example, Lauritzen (1996).

A straightforward way to study the dependence of variables in tensor is to vectorize the tensor into a vector and then estimate the inverse of the covariance matrix (i.e., the precision matrix). Under the sparsity assumption, many penalized methods have been proposed to obtain a sparse precision matrix (e.g., Yuan & Lin 2007, Friedman et al. 2008, Zhang & Zou 2014, Danaher et al. 2014, Witten et al. 2011, Cai et al. 2016, Molstad & Rothman 2018). However, the high dimensionality of the vectorized data is a great challenge to the existing methods. For example, a typical 3D MRI scan of the brain has dimension $p_1 \times p_2 \times p_3 = 256 \times 256 \times 256$. The vectorized data has over 16 million variables, making it unrealistic to estimate the precision matrix. Therefore, to exploit the tensor structural information, the parsimonious and interpretable Kronecker covariance structure is widely used in modeling tensor data. Under this assumption, the problem is then simplified to the estimation of three much smaller precision matrices of size $p_m \times p_m$, $m = 1, 2, 3$, in the MRI example.

Sparse precision matrix estimation under the Kronecker covariance structure has been gaining increasing attention in recent years: From matrix data (Zhu & Li 2018, Leng & Tang 2012, Yin & Li 2012, Zhou 2014) to higher-order tensors (Tsiligkaridis et al. 2013, He et al. 2014, Xu et al. 2017, Lyu et al. 2019). Assuming the precision matrices are sparse, the standard way is to use a penalized negative log-likelihood function. Then the estimators are obtained

through iterative algorithms to cyclic update one precision matrix while fixing the other $M - 1$ precision matrices at the estimated value from the previous iteration. Such cyclic update algorithms can be computationally costly, especially when the tensor dimensions or the tensor order are not small. Even for matrix Gaussian graphical models, as the dimensions increase, the computational cost increases drastically, especially when cross-validation is used for tuning parameter selection. In addition, when using the cyclicly iterative methods, all precision matrices are assumed to be sparse and estimated in a nested sequence. If we are only interested in the dependence of one particular mode, then the estimation accuracy would be a concern when other irrelevant precision matrices violate the sparsity assumption. Because all precision matrices are estimated through regularized optimization, the inaccurate estimates of non-sparse precision matrices can affect the estimation of the sparse precision matrices.

To tackle this challenge, we propose a non-cyclic and parallel approach for estimating the sparse tensor Gaussian graphical model with a relaxed assumption. Specifically, only the precision matrices of interest are assumed to be sparse. Each precision matrix is estimated by solving an independent ℓ_1 constrained minimization problem. For example, in some cases, such as modeling spatio-temporal data, some particular non-sparse correlation patterns can be adapted into our framework. The proposed method can also be incorporated into jointly estimating regression coefficients and covariances in tensor response regression. The proposed method improves the computation efficiency and the estimation accuracy in numerical studies. In theory, the proposed estimator enjoys the same optimal convergence rate in the graphical model literature.

The rest of the paper is organized as follows. In Section 2, we present some preliminaries on tensor algebra and introduce the tensor graphical model. In Section 3, we develop our separable non-iterative method and its theoretical properties. Simulation studies and a real data illustration are presented in Sections 4 and 5, respectively. Online Supplementary Materials contains proofs, R code, and applications to partially sparse models and tensor regression.

2 Background and Problem Set-up

2.1 Notation

We review some basic tensor notation and operations that commonly used (e.g., Kolda & Bader 2009). Multidimensional array $\mathbf{A} \in \mathbb{R}^{p_1 \times \dots \times p_M}$ is called a tensor of order M . The vectorization of a tensor $\mathbf{A}$ is denoted by $\text{vec}(\mathbf{A}) \in \mathbb{R}^{(\prod_{m=1}^M p_m) \times 1}$, with $A_{i_1, \dots, i_M}$ being its j -th element, where $j = 1 + \sum_{m=1}^M (i_m - 1) \prod_{m'=1}^{m-1} p_{m'}$. The Frobenius norm of $\mathbf{A}$ is defined as $\|\mathbf{A}\|_F = (\sum_{i_1, \dots, i_M} A_{i_1, \dots, i_M}^2)^{1/2}$. We define $p = \prod_{m=1}^M p_m$ and $p_{-m} = \prod_{m'=1, m' \neq m}^M p_{m'}$. The mode- m matricization of $\mathbf{A}$ is denoted by $\mathbf{A}_{(m)} \in \mathbb{R}^{p_m \times p_{-m}}$, which is obtained by combining the $(M-1)$ modes of the tensor similar to vectorization. The mode- m product of tensor $\mathbf{A}$ with a matrix $\boldsymbol{\alpha} \in \mathbb{R}^{d \times p_m}$ is defined as $\mathbf{A} \times_m \boldsymbol{\alpha}$ and it yields a tensor of size $p_1 \times \dots \times p_{m-1} \times d \times p_{m+1} \times \dots \times p_M$. Elementwise, we have $(\mathbf{A} \times_m \boldsymbol{\alpha})_{i_1, \dots, i_{m-1}, j, i_{m+1}, \dots, i_M} = \sum_{i_m=1}^{p_m} A_{i_1, \dots, i_m} \boldsymbol{\alpha}_{j, i_m}$. For a list of matrices $\{\boldsymbol{\alpha}\} = \{\boldsymbol{\alpha}_1, \dots, \boldsymbol{\alpha}_M\}$ with $\boldsymbol{\alpha}_m \in \mathbb{R}^{d_m \times p_m}$, we define $\mathbf{A} \times \{\boldsymbol{\alpha}\} = \mathbf{A} \times_1 \boldsymbol{\alpha}_1 \cdots \times_M \boldsymbol{\alpha}_M$. Let $\mathbf{Y} = \mathbf{A} \times \{\boldsymbol{\alpha}\}$, then $\mathbf{Y}_{(m)} = \boldsymbol{\alpha}_m \mathbf{A}_{(m)} (\boldsymbol{\alpha}_M \otimes \dots \otimes \boldsymbol{\alpha}_{m+1} \otimes \boldsymbol{\alpha}_{m-1} \otimes \dots \otimes \boldsymbol{\alpha}_1)^\top$ where $\otimes$ is the Kronecker product. Let $\{\boldsymbol{\alpha}\}_{-m}$ be the subset of $\{\boldsymbol{\alpha}\}$ without the m -th matrix.

2.2 Tensor Gaussian Graphical Model

The tensor random variable $\mathbf{Z} \in \mathbb{R}^{p_1 \times \dots \times p_M}$ follows a standard tensor normal distribution if all elements of $\mathbf{Z}$ are independent standard normal random variables. Let $\mathbf{X} = \boldsymbol{\mu} + \mathbf{Z} \times \{\boldsymbol{\Sigma}^{1/2}\}$, where $\{\boldsymbol{\Sigma}^{1/2}\} = \{\boldsymbol{\Sigma}_1^{1/2}, \dots, \boldsymbol{\Sigma}_M^{1/2}\}$ is a list of symmetric positive definite matrices, then $\mathbf{X}$ follows a tensor normal (TN) distribution with mean $\boldsymbol{\mu}$ and Kronecker separable covariance structure, denoted as $\mathbf{X} \sim \text{TN}(\boldsymbol{\mu}, \boldsymbol{\Sigma}_1, \dots, \boldsymbol{\Sigma}_M)$. If $M = 2$, the tensor normal distribution reduces to the matrix normal (MN) distribution (Gupta & Nagar (2018)). Suppose that a mode- M tensor $\mathbf{X} \in \mathbb{R}^{p_1 \times \dots \times p_M}$ follows a tensor normal distribution $\text{TN}(\mathbf{0}; \boldsymbol{\Sigma}_1, \dots, \boldsymbol{\Sigma}_M)$, then its probability density function is defined as

$$p(\mathbf{X} \mid \boldsymbol{\Sigma}_1, \dots, \boldsymbol{\Sigma}_M) = (2\pi)^{-p/2} \left\{ \prod_{m=1}^M |\boldsymbol{\Sigma}_m|^{-p_{-m}/2} \right\} \exp \left(-\frac{1}{2} \|\mathbf{X} \times \{\boldsymbol{\Sigma}^{-1/2}\}\|_F^2 \right), \quad (1)$$

where $\{\boldsymbol{\Sigma}^{-1/2}\} = \{\boldsymbol{\Sigma}_1^{-1/2}, \dots, \boldsymbol{\Sigma}_M^{-1/2}\}$. For $m = 1, \dots, M$, we call $\boldsymbol{\Omega}_m \equiv \boldsymbol{\Sigma}_m^{-1}$ the mode- m precision matrix, which is our target parameter.

To distinguish the population true parameter and the argument in optimization, we consider the model that $\mathbf{X}_1, \dots, \mathbf{X}_n$ are i.i.d. samples from $\text{TN}(\mathbf{0}; \Sigma_1^*, \dots, \Sigma_M^*)$. We study the problem of estimating the precision matrices $\{\Omega_1^*, \dots, \Omega_M^*\}$, where $\Omega_m^* = (\Sigma_m^*)^{-1}$ for $m = 1, \dots, M$. In particular, we focus on high-dimensional settings that Ω_m^* is sparse for $m = 1, \dots, M$. This sparse precision matrix assumption has the conditional independence interpretation that is analogous to the classical multivariate (vector) settings (Lauritzen 1996). Let $\mathbf{X}_{(m)}$ denote the mode- m matricization of $\mathbf{X}$ and $\mathbf{X}_{(m),i}$ denote the i -th row of $\mathbf{X}_{(m)}$. Then $[\Omega_m^*]_{i,j} = 0$ if and only if $\mathbf{X}_{(m),i}$ is independent of $\mathbf{X}_{(m),j}$ given $\mathbf{X}_{(m),k}, k \neq i, j$.

The precision matrices $\{\Omega_1^*, \dots, \Omega_M^*\}$ are identifiable up to $(M - 1)$ scaling constants. For example, it is easy to verify that $\text{TN}(\mathbf{0}; \Sigma_1^*, \Sigma_2^*, \dots, \Sigma_M^*)$ and $\text{TN}(\mathbf{0}; \Sigma_1^*/c, c\Sigma_2^*, \dots, \Sigma_M^*)$ have the same probability density function for any $c > 0$. To address the identifiability problem in the tensor normal model, we assume that $\|\Omega_m^*\|_F = 1$ for all $m = 1, \dots, M$, which can be easily satisfied by standardization. Alternatively, one can impose $(M - 1)$ constraints that $\|\Omega_m^*\|_F = 1$ for all $m = 2, \dots, M$, and absorb the scaling constant into Ω_1^* . Clearly, the scaling does not affect the sparsity pattern of the precision matrices and is a nuisance parameter in graphical model.

A standard approach to the tensor graphical model estimation is based on minimizing the penalized negative log-likelihood function,

$$L_n(\Omega_1, \dots, \Omega_M) = \frac{1}{p} \text{tr} \{ \mathbf{S} (\Omega_M \otimes \dots \otimes \Omega_1) \} - \sum_{m=1}^M \frac{1}{p_m} \log |\Omega_m| + \sum_{m=1}^M P_{\lambda_m}(\Omega_m), \quad (2)$$

where $\mathbf{S} = n^{-1} \sum_{i=1}^n \text{vec}(\mathbf{X}_i) \text{vec}(\mathbf{X}_i)^\top$ and $P_{\lambda_m}(\cdot)$ is a penalized function indexed by the tuning parameter $\lambda_m > 0, m = 1, \dots, M$. In this paper, we focus on the ℓ_1 penalty on the off-diagonal elements: $P_{\lambda_m}(\Omega_m) = \lambda_m \|\Omega_m\|_{1,\text{off}}$, where $\|\Omega_m\|_{1,\text{off}} = \sum_{i \neq j} |[\Omega_m]_{i,j}|$. The penalized model in (2) is referred to as the sparse tensor graphical model.

3 Proposed Method

3.1 The Separable Estimation Approach

In this section, we propose a scalable and parallelizable approach for solving the tensor graphical lasso problem that is computationally much more efficient than the widely used cyclic updating scheme. The idea is to estimate each Ω_m^* separately so that the method is robust and flexible in dealing with different modes of the tensors.

The proposed estimator is obtained through a penalized pseudo-likelihood approach. When estimating Ω_m^* , suppose that the other $M - 1$ precision matrices are fixed at a list of symmetric and positive definite matrices $\{\tilde{\Omega}\}_{-m}$. Then we can obtain the estimator $\hat{\Omega}_m$ by minimizing the following pseudo-log-likelihood function with respect to Ω_m ,

$$L_{nm}(\Omega_m; \{\tilde{\Omega}\}_{-m}) = \frac{1}{p_m} \text{tr} \left(\tilde{\mathbf{S}}_m \Omega_m \right) - \frac{1}{p_m} \log |\Omega_m| + \lambda_m \|\Omega_m\|_{1,\text{off}}, \quad (3)$$

where $\tilde{\mathbf{S}}_m = \frac{1}{np_{-m}} \sum_{i=1}^n \tilde{\mathbf{V}}_i^m (\tilde{\mathbf{V}}_i^m)^T$ and $\tilde{\mathbf{V}}_i^m = \mathbf{X}_{i(m)} (\tilde{\Omega}_M^{1/2} \otimes \cdots \otimes \tilde{\Omega}_{m+1}^{1/2} \otimes \tilde{\Omega}_{m-1}^{1/2} \cdots \otimes \tilde{\Omega}_1^{1/2})$. Recall that we use normalization to achieve $\|\tilde{\Omega}_m\|_F = \|\hat{\Omega}_m\|_F = 1$ for all m . For our separable estimation, we use

$$\tilde{\Omega}_m = \begin{cases} \left\{ \frac{1}{np_{-m}} \sum_{i=1}^n \mathbf{X}_{i(m)} (\mathbf{X}_{i(m)})^T \right\}^{-1}, & np_{-m} > p_m(p_m - 1)/2; \\ \mathbf{I}_{p_m}, & np_{-m} \leq p_m(p_m - 1)/2. \end{cases} \quad (4)$$

Clearly, the key advantage of our separable estimation is that the optimization problem (3) for Ω_m does not depend on the sparse estimators for other precision matrices $\{\hat{\Omega}\}_{-m}$. Moreover, the precision matrices $\{\tilde{\Omega}\}_{-m}$ used in optimization (4) are adaptive to the tensor dimensions and sample size. The number $n_m = np_{-m}$ is the effective sample size for estimating Ω_m , which has $p_m(p_m - 1)/2$ unique parameters.

We present the following results for the Kronecker separable covariance model and matrix/tensor normal distribution to gain more insight into our estimator. Similar results have been obtained in recent studies (Lyu et al. 2019, Pan et al. 2019, Drton et al. 2020).

Lemma 1. *Suppose that $\mathbf{X}_1, \dots, \mathbf{X}_n$ are i.i.d. from $\text{TN}(\mathbf{0}; \Sigma_1^*, \dots, \Sigma_M^*)$. When $np_{-j} > p_j$, $\sum_{i=1}^n \mathbf{X}_{i(j)} (\mathbf{X}_{i(j)})^T$ is positive definite with probability 1. When $\{\tilde{\Omega}\}_{-m} = \{\alpha\}_{-m}$ is non-*

Algorithm 1 Parallel algorithm for tensor Gaussian graphical models

1. **Input:** Tensor samples $\mathbf{X}_1, \dots, \mathbf{X}_n$, tuning parameters $\lambda_1, \dots, \lambda_M$.
 2. **Initialization:** For each $m \in \{1, \dots, M\}$, compute $\tilde{\Omega}_m$ using (4).
 3. **Sparse Estimation:** For each $m \in \{1, \dots, M\}$, solve the optimization problem (3) for $\hat{\Omega}_m$ using the glasso algorithm (Friedman et al. 2008).
-

stochastic, the sample covariance matrix $\tilde{\mathbf{S}}_m$ in (3) has expected value $E(\tilde{\mathbf{S}}_m \mid \{\tilde{\Omega}\}_{-m} = \{\alpha\}_{-m}) = \frac{1}{p_{-m}} \left\{ \prod_{j \neq m} \text{tr}(\Sigma_j^* \alpha_j) \right\} \Sigma_m^*$.

By Lemma 1, when $n_j > p_j(p_j - 1)/2$, the sample covariance matrix $\tilde{\mathbf{S}}_m$ in (3) is positive definite with probability 1. Moreover, it is an unbiased estimator for Σ_m^* after scaling when we simply use any fixed $\{\tilde{\Omega}\}_{-m}$, e.g. the identity matrix in (4). Of course, when the sample size increases, we expect the well-conditioned sample estimator $\tilde{\Omega}_m = \left\{ \frac{1}{np_{-m}} \sum_{i=1}^n \mathbf{X}_{i(m)} (\mathbf{X}_{i(m)})^T \right\}^{-1}$ improves estimation accuracy over the identity matrix.

After obtaining $\{\tilde{\Omega}\}_{-m}$ from (4), solving the optimization problem in (3) is the same as solving a classical graphical lasso problem (Friedman et al. 2008). When estimating a list of precision matrices, minimizing L_{nm} is independent of minimizing $L_{nm'}$. Thus, the set of optimization problems are separable and can be solved simultaneously by implementing parallel computing. We summarize our scalable and parallelizable algorithm in Algorithm 1. Even when comparing it with the one-iteration estimator from cyclic algorithms, we still see a decrease in computation cost in the numerical studies. The M sparse estimation problems in Algorithm 1 are parallel, making the proposed method differs from the one-iteration estimator. It has better performance, especially when some Ω_m^* is not very sparse and can not be estimated well by sparse solutions.

3.2 Theoretical Properties

We show that the proposed estimator converges to the true precision matrix at an optimal rate. Since each optimization problem is independent of the others and can be solved separately or simultaneously by parallel computing, we present the theoretical results for an arbitrary mode,

i.e. analysis only on $\widehat{\Omega}_m$.

We explicitly express the sample-based minimization to (3) and the proposed estimator $\widehat{\Omega}_m$, which requires normalization, as follows,

$$\widehat{M}_m(\{\widetilde{\Omega}\}_{-m}) = \arg \min_{\Omega_m} L_{nm}(\Omega_m; \{\widetilde{\Omega}\}_{-m}), \quad \widehat{\Omega}_m = \frac{\widehat{M}_m(\{\widetilde{\Omega}\}_{-m})}{\|\widehat{M}_m(\{\widetilde{\Omega}\}_{-m})\|_F}. \quad (5)$$

To characterize the sparsity of the true precision matrix Ω_m^* , we define a sparsity parameter $s_m = |\mathbb{S}_m| - p_m$, where $\mathbb{S}_m = \{(i, j) : [\Omega_m^*]_{i,j} \neq 0\}$. Then s_m is the number of non-zero off-diagonal elements in Ω_m^* . We construct the convergence theory for the proposed estimator in the following theorem, under two mild technical conditions.

Condition (C1). *Bounded Eigenvalues.* For any $m = 1, \dots, M$, there is a constant $C_1 > 0$ such that $0 < C_1 \leq \lambda_{\min}(\Sigma_m^*) \leq \lambda_{\max}(\Sigma_m^*) \leq 1/C_1 < \infty$, where $\lambda_{\min}(\cdot)$ and $\lambda_{\max}(\cdot)$ denote the minimal and maximal eigenvalues.

Condition (C2). *Tuning.* There is a constant $C_2 > 0$ such that the tuning parameter λ_m satisfies $1/C_2 \sqrt{\log p_m / (npp_m)} \leq \lambda_m \leq C_2 \sqrt{\log p_m / (npp_m)}$.

Theorem 1. Suppose that Conditions (C1) and (C2) hold. The proposed estimator $\widehat{\Omega}_m$ in (5) satisfies $\|\widehat{\Omega}_m - \Omega_m^*\|_F = O_P\left(\sqrt{\frac{(p_m + s_m) \log p_m}{np - m}}\right)$.

Condition (C1) requires that the eigenvalues of the true covariance matrices are bounded uniformly. It is a commonly used assumption to study the estimation consistency in graphical models (He et al. 2014, Xu et al. 2017, Lyu et al. 2019). Condition (C2) gives the growth rate of the tuning parameters in terms of n and p . This type of assumption is also widely used in penalized methods to control the estimation bias and sparsity (Zhou 2014, Lyu et al. 2019).

Theorem 1 shows that the proposed estimator $\widehat{\Omega}_m$ converges to the true precision matrix Ω_m^* at a rate of $\sqrt{(p_m + s_m) \log p_m / (np - m)}$ in Frobenius norm. The same convergence rate is achieved by the Tlasso estimator after one iteration and it is minimax-optimal (Cai et al. 2016, Lyu et al. 2019). When the true precision matrices are known and let $\{\widetilde{\Omega}\}_{-m} = \{\Omega^*\}_{-m}$, the effective sample size for estimating $\widetilde{\Sigma}_m$ is $np - m$ since $\widetilde{\mathbf{V}}_i^m$ has independent columns. From the

analysis in Cai et al. (2016), the best rate one can obtain is $\sqrt{(p_m + s_m) \log p_m / (np - m)}$ which matches ours. Therefore, the convergence rate in Theorem 1 is also minimax-optimal.

We compare the theoretical results with other existing works. Leng & Tang (2012) showed the existence of a local minimizer that enjoyed the same convergence rate as ours at $M = 2$. However, they did not show that their algorithm could guarantee the local minimizer. Yin & Li (2012) established the convergence rates at $\sqrt{p_2(p_1 + s_1)(\log p_1 + \log p_2)/n}$ which was slightly slower. To achieve the optimal estimation error, Xu et al. (2017) required the number of iteration to be no less than $C \log(np/(p_m s_m))$ where C is a positive constant. In contrast, our proposed method does not have convergence concerns and enjoys the same optimal convergence property.

3.3 Connection with existing approaches

Clearly when $M = 1$, our method reduces to the sparse Gaussian graphical model which has been extensively studied (e.g., Yuan & Lin 2007, Banerjee et al. 2008, Friedman et al. 2008, Rothman et al. 2008, Ravikumar et al. 2011). When $M = 2$, our method is then applicable to the sparse matrix graphical model. Various iterative algorithm can be viewed as the matrix version of our problem (see Leng & Tang 2012, Yin & Li 2012, Tsiligkaridis et al. 2013). The only exception is Zhou (2014), which proposed a non-cyclic algorithm for inverse correlation matrix estimation. Our method extends Zhou (2014) from matrix to tensor, and from inverse correlation to inverse covariance (which is easier to interpret in graphical models), and the plug-in estimator $\tilde{\Omega}_j$ is not fixed as $\mathbf{I}_{p_j}$ for $j \neq m$ for more efficient estimation in larger sample scenarios. Since the effective sample size of tensor data is generally larger than the actual sample size, we can get an improved estimation.

For the sparse tensor graphical model with a general $M \geq 2$, He et al. (2014), Lyu et al. (2019) used a ℓ_1 penalty and provided a coordinate descent-based algorithm; Xu et al. (2017) used a l_0 penalty and proposed an alternating gradient descent algorithm. To the best of our knowledge, all existing methods propose iterative cyclic-updating algorithms.

The proposed separable and parallelizable algorithm circumvents many limitations in the

	(p_1, p_2, p_3)	Separate	Oracle	Sequential	Cyclic
Model 1	(30,36,30)	1.17 (0.01)	0.84 (0.01)	3.10 (0.01)	13.24 (0.92)
Model 2	(100,100,100)	44.13 (0.07)	34.44 (0.11)	93.22 (0.05)	379.13 (1.88)
Model 3	(5,5,500)	14.21 (0.04)	13.90 (0.03)	13.94 (0.02)	120.91 (0.07)

Table 1: Averaged computation cost (in seconds) and standard errors (in parentheses) for Models 1–3 from 20 replicates.

cyclic algorithms. Firstly, the efficiency of iterative algorithms is a concern due to the high dimensionality of tensor data. The new algorithm significantly reduces the computational cost through parallelism. Even comparing with the first iteration in the cyclic algorithms, we still see a decrease in computation cost in numerical studies. Secondly, the choice of initial value and the convergence speed of the algorithm should be considered for iterative approaches. The initialization step in Algorithm 1 offers a simple initial value with theoretical guarantees. Thirdly, with a well-chosen initial value for $\tilde{\Omega}_m$, we can finely tune λ_m at low computational cost for each mode. In practice, the more affordable and efficient tuning process can further improve performance. Finally, the cyclically iterative methods assume that all the precision matrices are sparse, which can be restrictive if we are only interested in estimating a subset of the precision matrices. In contrast, thanks to the independence of the optimization problems in our algorithm, we can be more flexible at scenarios with different sparsity levels. In Section B of the Supplementary Materials, we will discuss the applications to partially sparse graphical models, pre-specified covariance structured (e.g. longitudinal, spatial-temporal models), and joint model of mean and covariance in regression.

4 Simulations

We compare the proposed estimator (“*Separate*”) with three alternative solutions: (“*Oracle*”) The oracle estimator using the true parameter Ω_{-m}^* when estimates Ω_m^* in (3); (“*Cyclic*”) The cyclic-updates algorithm to the penalized MLE problem (2), where each sparse precision matrix is updated using current estimates of other sparse precision matrices and convergence is

(%)	Model 1				Model 2				Model 3				S.E.≤
	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.	
Frobenius norm loss													
AVG.	0.62	0.62	2.34	0.80	0.21	0.21	1.28	0.27	1.97	1.83	3.33	13.74	(0.01)
$m = 1$	0.59	0.59	3.74	0.73	0.21	0.21	2.29	0.27	0.50	0.30	0.79	1.87	(0.02)
$m = 2$	0.66	0.65	2.58	0.92	0.21	0.21	1.29	0.27	0.52	0.31	2.43	2.05	(0.02)
$m = 3$	0.62	0.62	0.71	0.75	0.22	0.21	0.26	0.27	4.88	4.88	6.77	37.29	(0.01)
Max norm loss													
AVG.	0.15	0.15	0.29	0.15	0.03	0.03	0.06	0.03	0.27	0.21	0.61	0.78	(0.01)
$m = 1$	0.15	0.14	0.44	0.15	0.03	0.03	0.10	0.03	0.22	0.14	0.37	0.81	(0.01)
$m = 2$	0.16	0.15	0.27	0.16	0.03	0.03	0.06	0.03	0.23	0.14	1.08	0.89	(0.01)
$m = 3$	0.16	0.16	0.17	0.16	0.04	0.03	0.04	0.04	0.37	0.36	0.38	0.64	(0.01)
True negative rate													
AVG.	66.21	66.19	19.28	41.08	81.88	81.42	26.09	62.18	45.96	46.02	54.34	27.32	(1.21)
$m = 1$	65.77	65.75	0.07	43.55	82.04	81.50	0.04	62.19	20.25	21.00	40.83	37.83	(2.2)
$m = 2$	67.95	68.06	2.29	36.46	81.58	81.10	2.76	62.04	23.83	24.33	42.17	42.83	(2.48)
$m = 3$	64.9	64.75	55.50	43.23	82.01	81.67	75.48	62.31	93.79	92.72	80.02	1.30	(0.52)
(%)	Model 4				Model 5				Model 6				S.E.≤
	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.	
Frobenius norm loss													
$m = 2$	0.67	0.66	2.78	1.64	-	-	-	-	-	-	-	-	(0.01)
$m = 3$	0.60	0.59	0.71	1.37	0.61	0.61	1.40	1.39	4.91	4.88	13.41	38.64	(0.02)
Max norm loss													
$m = 2$	0.16	0.16	0.29	0.17	-	-	-	-	-	-	-	-	(0.01)
$m = 3$	0.16	0.15	0.16	0.17	0.16	0.16	0.17	0.17	0.36	0.36	0.36	0.66	(0.01)
True negative rate													
$m = 2$	67.59	67.28	1.47	0.68	-	-	-	-	-	-	-	-	(0.35)
$m = 3$	64.48	64.56	48.63	0.81	65.77	65.63	0	0	91.34	92.17	47.34	0.18	(0.83)

Table 2: Comparison of means and the maximum standard errors (in parentheses) of different performance measures for Models 1–3 from 100 replicates. All methods have achieved 100% true positive rate (and hence not shown in the table). Comparison of means and the maximum standard errors (in parentheses) of different performance measures for Models 4–6 from 100 replicates. Results for $m = 2$ in Models 5 & 6 are not reported because Ω_2^* is not sparse.

often reached after 5–100 iterations; and (“*Sequential*”) The cyclic algorithm with only one iteration, which means $\hat{\Omega}_m$ is obtained using the sparse estimators $\hat{\Omega}_j$, $j < m$.

For fair comparison, we use the same R package *glasso* to compute the ℓ_1 penalized optimizations for all four estimators. We expect similar results if different penalties or packages

are used consistently for the four estimators. Additional implementation details are given in the Supplementary Materials.

We consider the following three types of covariance/precision matrices $\Sigma_m^*, \Omega_m^* \in \mathbb{R}^{p_m \times p_m}$.

- Triangle (TR) covariance. We set $[\Sigma_m^*]_{i,j} = \exp(-|h_i - h_j|/2)$ with $h_1 < h_2 < \dots < h_{p_m}$. The difference $h_i - h_{i-1}, i = 2, \dots, p_m$, is generated i.i.d. from $\text{Unif}(0.5, 1)$. This generated covariance matrix mimics autoregressive process of order one, i.e., AR(1).
- Autoregressive (AR) precision. We set $[\Omega_m^*]_{i,j} = 0.8^{|i-j|}$.
- Compound symmetry (CS) precision. We set $[\Omega_m^*]_{i,j} = 0.6$ for $i \neq j$ and $[\Omega_m^*]_{i,i} = 1$.

Note that the first model (TR) has a sparse precision matrix while the other two (AR and CS) have a non-sparse precision matrix.

We consider six models as follows. Models 1–3 are fully sparse models, and Models 4–6 are partially sparse models. In Models 3 & 6, we consider the unbalanced setting in dimension. In all models, we normalize the precision matrices such that $\|\Omega_m^*\|_F = 1$ for $m = 1, \dots, M$. For all the following models, we set $M = 3$ and $n = 100$.

- Model 1: $\Omega_1^*, \Omega_2^*, \Omega_3^*$ are all from the TR covariance model, $(p_1, p_2, p_3) = (30, 36, 30)$.
- Model 2: $\Omega_1^*, \Omega_2^*, \Omega_3^*$ are all from the TR covariance model, $(p_1, p_2, p_3) = (100, 100, 100)$.
- Model 3: $\Omega_1^*, \Omega_2^*, \Omega_3^*$ are all from the TR covariance model, $(p_1, p_2, p_3) = (5, 5, 500)$.
- Model 4: Same as Model 1, except for $\Omega_1^* = \text{AR}(0.8)$.
- Model 5: Same as Model 1, except for $\Omega_1^* = \Omega_2^* = \text{CS}(0.6)$.
- Model 6: Same as Model 3, except for $\Omega_1^* = \Omega_2^* = \text{CS}(0.6)$.

We also consider the six models but with sample size $n = 20$ and another two fully sparse models with $M = 4$ and $M = 5$. The results show similar findings and are included in the Supplementary Materials. To measure the estimation performance, we report the errors in Frobenius norm $\|\hat{\Omega}_m - \Omega_m^*\|_F$ and the errors in max norm $\|\hat{\Omega}_m - \Omega_m^*\|_{\max}$. We also report the

true positive rate and true negative rate to measure the performance of support recovery. For fully sparse models, we further report the averages of these criteria across all modes.

The computation time for Model 1–3 is summarized in Table 1. The proposed separable estimation approach is much faster than the cyclic-updating approach, especially as the dimension increases. The computational costs of Separable, Oracle and Sequential methods are roughly the same, as we expected. For Model 4–6, the computation time is not included because the sequential and cyclic methods failed in these partially sparse models.

The results of estimation accuracy are summarized in Table 2. For both fully sparse models and partially sparse models, we can observe that the separable estimator outperforms the sequential one-iteration estimator and cyclic estimator on every mode and is comparable to the oracle estimator. In particular, when the dimension is unbalanced in Models 3 & 6, the cyclic estimator performs much worse than the proposed method. Moreover, comparing Models 1 & 2 with Models 4 & 5, we notice that the cyclically iterative algorithm cannot improve the estimation accuracy in the iteration for partially sparse models. This confirms our concern that violation of sparsity assumption hurts the performance of iterative methods. It also indicates that the performance of the proposed method is more stable and consistent. The proposed method also shows superior performance in terms of support recovery.

5 Real Data Analysis

As an illustration, we apply the proposed method on the electroencephalography (EEG) data from a study to examine EEG correlates of genetic predisposition to alcoholism. The dataset is available at <http://kdd.ics.uci.edu/databases/eeg/>. In the study, 64 channels of electrodes were placed on the scalp at standard sites. There were 122 subjects among which 77 were alcohol individuals and 45 were nonalcoholic individuals. Each subject had 120 trials under exposure to different picture stimuli and measurements were collected from the electrodes which were sampled at 256 Hz (3.9-msec epoch) for 1 second. More collection details could be found in Zhang et al. (1995). We used the same part of the data that was analyzed in Li et al. (2010). The data focused on single stimulus condition and were averages from

all the corresponding trials. Therefore, each EEG sample was a 256×64 matrix. We further downsized it to 64×64 . We divided the dataset into an alcoholic group and a nonalcoholic group. For each group, we standardized the data and applied the proposed method to estimate the precision matrix for channels. Tuning parameter was chosen by five-fold cross-validation.

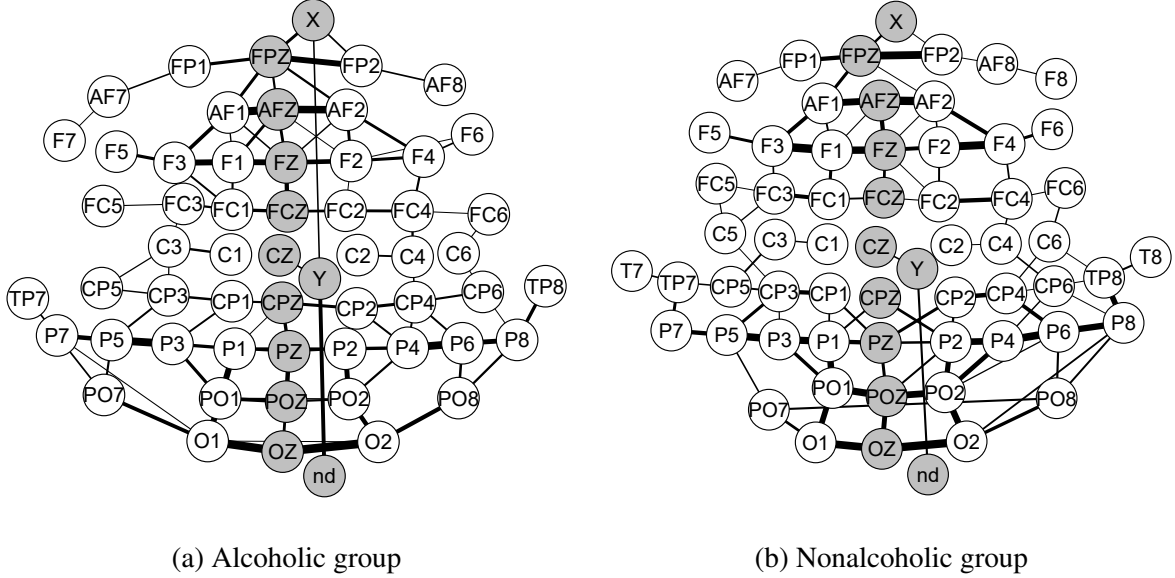


Figure 1: Correlation networks constructed by the proposed method using the estimated precision matrices among channels for alcoholic group and nonalcoholic group. Only the first 100 strongest correlations are displayed. Nodes are labeled with EEG electrode identifiers. The gray nodes are placed on the middle of the scalp whose left and right nodes are placed on the left and right of the scalp respectively. Nodes that are not correlated with the others are not included. The thickness of edges represents the magnitude of correlations.

The correlation networks identified by the proposed method are displayed in Figure 1. To have a clear insight of the graph pattern, we only present the top 100 correlations. We can observe that channels located close to each other tend to be correlated. For example, channels F4, F2, FZ, F1, F3 and F5 that are physically placed in line on the scalp are strongly correlated. The centered channels such as AFZ and OZ also have strong connections with the adjacent nodes AF1, AF2 and O1, O2. We also notice that the correlations among POZ and PO1, PO2 are stronger in nonalcoholic group than in alcoholic group. As POZ is located near the parietal-occipital junction, there is a large probability for it to be most sensitive to the visual stimuli

(Lei & Liao 2017). A possible explanation for the weaker correlations in alcoholic group is that alcohol can slow down the brain’s information processing (Tzambazis & Stough 2000).

The results from the cyclic algorithm and the colored version of the network graphs are given in the Supplementary Materials. While the cyclic algorithm can identify most of the strong correlation pairs, the total number of correlation pairs is about two or three times of the number identified by the proposed method. For nonalcoholic group, the proposed method identifies 387 and 425 correlation pairs respectively for the two modes, while the cyclic algorithm identifies 888 and 1598 pairs. For alcoholic group, the proposed method identifies 416 and 441 pairs and the cyclic algorithm identifies 947 and 1536 pairs. The separable method offers a more sparse solution that is easier to interpret.

Supplementary Materials

Supp.pdf This file contains proofs, additional numerical results, and application to partially sparse model. (pdf)

Rcode.zip The file includes R code files for simulation and real data analysis and a readme file for demonstration. (zip)

References

- Banerjee, O., Ghaoui, L. E. & d’Aspremont, A. (2008), ‘Model selection through sparse maximum likelihood estimation for multivariate gaussian or binary data’, *Journal of Machine learning research* **9**(Mar), 485–516.
- Bi, X., Tang, X., Yuan, Y., Zhang, Y. & Qu, A. (2020), ‘Tensors in statistics’, *Annual Review of Statistics and Its Application* **8**.
- Cai, T. T., Liu, W. & Zhou, H. H. (2016), ‘Estimating sparse precision matrix: Optimal rates of convergence and adaptive estimation’, *The Annals of Statistics* **44**(2), 455–488.
- Chi, E. C. & Kolda, T. G. (2012), ‘On tensors, sparsity, and nonnegative factorizations’, *SIAM Journal on Matrix Analysis and Applications* **33**(4), 1272–1299.
- Danaher, P., Wang, P. & Witten, D. M. (2014), ‘The joint graphical lasso for inverse covariance estimation across multiple classes’, *Journal of the Royal Statistical Society. Series B, Statistical methodology* **76**(2), 373.
- Drton, M., Kuriki, S. & Hoff, P. (2020), ‘Existence and uniqueness of the kronecker covariance mle’, *arXiv preprint arXiv:2003.06024*.

- Friedman, J., Hastie, T. & Tibshirani, R. (2008), ‘Sparse inverse covariance estimation with the graphical lasso’, *Biostatistics* **9**(3), 432–441.
- Gupta, A. K. & Nagar, D. K. (2018), *Matrix variate distributions*, Chapman and Hall/CRC.
- He, S., Yin, J., Li, H. & Wang, X. (2014), ‘Graphical model selection and estimation for high dimensional tensor data’, *Journal of Multivariate Analysis* **128**, 165–185.
- Hoff, P. D. (2015), ‘Multilinear tensor regression for longitudinal relational data’, *The annals of applied statistics* **9**(3), 1169.
- Kolda, T. G. & Bader, B. W. (2009), ‘Tensor decompositions and applications’, *SIAM review* **51**(3), 455–500.
- Lauritzen, S. L. (1996), *Graphical models*, Vol. 17, Clarendon Press.
- Lei, X. & Liao, K. (2017), ‘Understanding the influences of eeg reference: a large-scale brain network perspective’, *Frontiers in neuroscience* **11**, 205.
- Leng, C. & Tang, C. Y. (2012), ‘Sparse matrix graphical models’, *Journal of the American Statistical Association* **107**(499), 1187–1200.
- Li, B., Kim, M. K. & Altman, N. (2010), ‘On dimension folding of matrix-or array-valued statistical objects’, *The Annals of Statistics* **38**(2), 1094–1121.
- Lock, E. F. (2018), ‘Tensor-on-tensor regression’, *Journal of Computational and Graphical Statistics* **27**(3), 638–647.
- Lyu, X., Sun, W. W., Wang, Z., Liu, H., Yang, J. & Cheng, G. (2019), ‘Tensor graphical model: Non-convex optimization and statistical inference’, *IEEE transactions on pattern analysis and machine intelligence*.
- Molstad, A. J. & Rothman, A. J. (2018), ‘Shrinking characteristics of precision matrix estimators’, *Biometrika* **105**(3), 563–574.
- Pan, Y., Mai, Q. & Zhang, X. (2019), ‘Covariate-adjusted tensor classification in high dimensions’, *Journal of the American statistical association* **114**(527), 1305–1319.
- Pfeiffer, R. M., Kapla, D. B. & Bura, E. (2020), ‘Least squares and maximum likelihood estimation of sufficient reductions in regressions with matrix-valued predictors’, *International Journal of Data Science and Analytics* pp. 1–16.
- Ravikumar, P., Wainwright, M. J., Raskutti, G. & Yu, B. (2011), ‘High-dimensional covariance estimation by minimizing l1-penalized log-determinant divergence’, *Electronic Journal of Statistics* **5**, 935–980.
- Rothman, A. J., Bickel, P. J., Levina, E. & Zhu, J. (2008), ‘Sparse permutation invariant covariance estimation’, *Electronic Journal of Statistics* **2**, 494–515.
- Tsiligkaridis, T., Hero III, A. O. & Zhou, S. (2013), ‘On convergence of kronecker graphical lasso algorithms’, *IEEE transactions on signal processing* **61**(7), 1743–1755.
- Tzambazis, K. & Stough, C. (2000), ‘Alcohol impairs speed of information processing and simple and choice reaction time and differentially impairs higher-order cognitive abilities’, *Alcohol and Alcoholism* **35**(2), 197–201.

- Witten, D. M., Friedman, J. H. & Simon, N. (2011), ‘New insights and faster computations for the graphical lasso’, *Journal of Computational and Graphical Statistics* **20**(4), 892–900.
- Xu, P., Zhang, T. & Gu, Q. (2017), Efficient algorithm for sparse tensor-variate gaussian graphical models via gradient descent, in ‘Artificial Intelligence and Statistics’, pp. 923–932.
- Yin, J. & Li, H. (2012), ‘Model selection and estimation in the matrix normal graphical model’, *Journal of multivariate analysis* **107**, 119–140.
- Yuan, M. & Lin, Y. (2007), ‘Model selection and estimation in the gaussian graphical model’, *Biometrika* **94**(1), 19–35.
- Zhang, T. & Zou, H. (2014), ‘Sparse precision matrix estimation via lasso penalized d-trace loss’, *Biometrika* **101**(1), 103–120.
- Zhang, X. L., Begleiter, H., Porjesz, B., Wang, W. & Litke, A. (1995), ‘Event related potentials during object recognition tasks’, *Brain research bulletin* **38**(6), 531–538.
- Zhou, H., Li, L. & Zhu, H. (2013), ‘Tensor regression with applications in neuroimaging data analysis’, *Journal of the American Statistical Association* **108**(502), 540–552.
- Zhou, S. (2014), ‘Gemini: Graph estimation with matrix variate normal instances’, *The Annals of Statistics* **42**(2), 532–562.
- Zhu, Y. & Li, L. (2018), ‘Multiple matrix gaussian graphs estimation’, *Journal of the Royal Statistical Society. Series B, Statistical methodology* **80**(5), 927.

Supplementary Materials for “Fast and Separable Estimation in High-dimensional Tensor Gaussian Graphical Models”

Keqian Min, Qing Mai and Xin Zhang
Florida State University

A Additional Data Analysis Results

We present the colored version of the correlation networks constructed by the proposed method in Figure 1 and the correlation networks constructed by Tlasso for the alcoholic group and nonalcoholic group in Figure 2. Adjacent electrodes tend to be correlated.

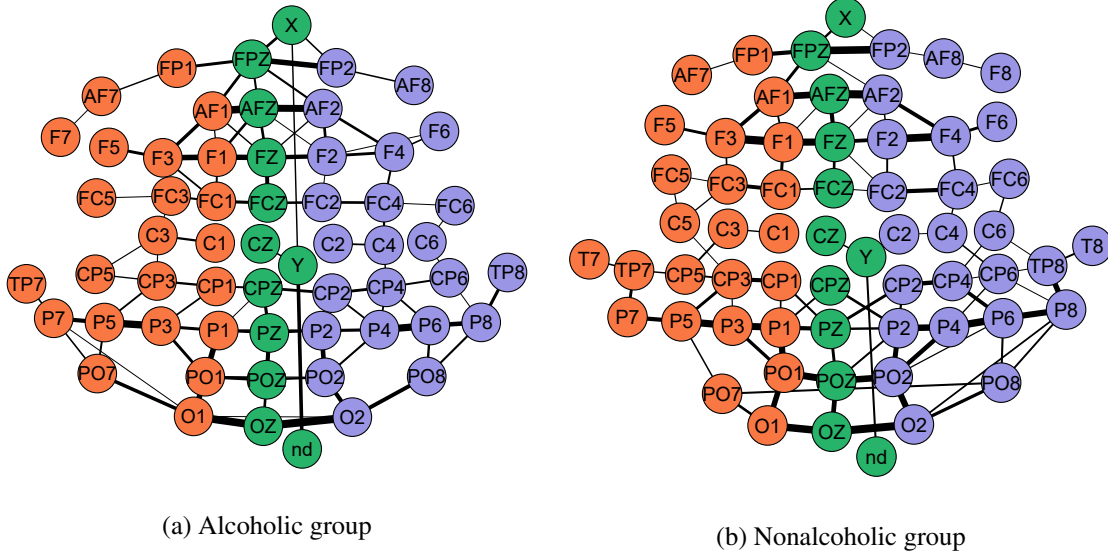


Figure 1: Correlation networks constructed by the proposed method using the estimated precision matrices among channels for alcoholic group and nonalcoholic group. Only the first 100 strongest correlations are displayed. Nodes are labeled with EEG electrode identifiers. Orange, green and purple nodes are placed on the left, middle and right of the scalp respectively. Nodes that are not correlated with the others are not included. The thickness of edges represents the magnitude of correlations.

B Application to Partially Sparse Model and Regression

A great motivation for considering partially sparse model comes from the special correlation structures used in modeling spatio-temporal data in environmental and longitudinal imaging studies (Genton 2007, George & Aban 2015, George et al. 2016, e.g). Typically, the joint covariance is described by the Kronecker product of a spatial covariance Σ_S and a temporal covariance Σ_T . Let $\mathbf{Y}_i \in \mathbb{R}^{T \times L}$ be the

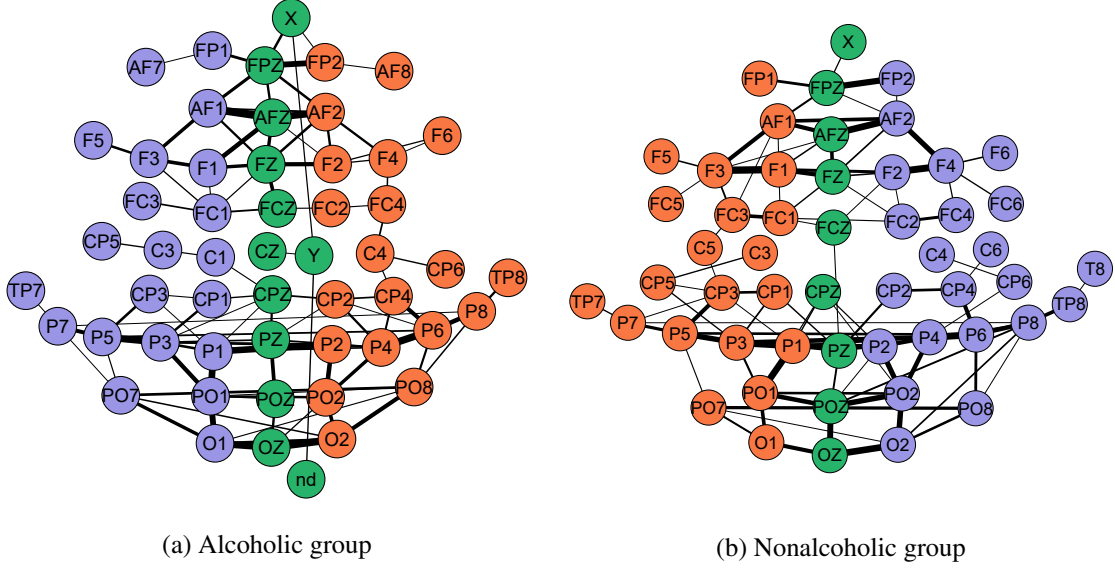


Figure 2: Correlation networks constructed by Tlasso using the estimated precision matrices among channels for alcoholic group and nonalcoholic group. Only the first 100 strongest correlations are displayed. Nodes are labeled with EEG electrode identifiers. Orange, green and purple nodes are placed on the left, middle and right of the scalp respectively. Nodes that are not correlated with the others are not included. The thickness of edges represents the magnitude of correlations.

observed response of subject i and Y_{itl} be the data observed at time t_j and location s_l with $j = 1, \dots, J$ and $l = 1, \dots, L$. George & Aban (2015), George et al. (2016) considered the following linear model

$$\mathbf{Y}_i = \mathbf{U}_i \boldsymbol{\beta} + \boldsymbol{\epsilon}_i, \quad \boldsymbol{\epsilon}_i \sim \text{MN}(\mathbf{0}; \boldsymbol{\Sigma}_T, \boldsymbol{\Sigma}_S) \quad (\text{B.1})$$

where $\mathbf{U}_i$ is the predictor and $\boldsymbol{\beta}$ is the associated coefficient.

Many special covariance structures have been widely used for modeling $\boldsymbol{\Sigma}_S$ and $\boldsymbol{\Sigma}_T$. For spatial covariance $\boldsymbol{\Sigma}_S$, a popular covariance structure is the Matérn covariance (Matérn 2013, Guttormp & Gneiting 2006). For temporal covariance $\boldsymbol{\Sigma}_T$, some possibilities are as follows (Diggle et al. 2002, Fitzmaurice et al. 2012). Let $\sigma_{T,jk}$ be the element in j th row and k th column in $\boldsymbol{\Sigma}_T$ and denote ϕ, ρ, v, τ as some non-negative parameters. The exponential correlation model assumes $\sigma_{T,jk} = \tau^2 \cdot \exp(-\phi|t_j - t_k|)$. In addition, if the time points t_j are equally spaced, then this model becomes the first order autoregressive model $\sigma_{T,jk} = \tau^2 \rho^{|j-k|}$. The uniform correlation model assumes that $\boldsymbol{\Sigma}_T = (\nu^2 + \tau^2) \{\rho \mathbf{J}_T + (1-\rho) \mathbf{I}_T\}$, where $\rho = \frac{\nu^2}{\nu^2 + \tau^2}$, $\mathbf{I}$ is the identity matrix and $\mathbf{J}$ is the matrix with all elements equal to one. This corresponds to the compound symmetry structure. The hybrid model combines the previous two models.

While some of the covariance structures and their inverses are not sparse, it motivates us to consider the estimation of a partially sparse model where we assume that

$$\text{there exists an index set } M_0 \text{ such that } \boldsymbol{\Omega}_m^* \text{ is sparse for } m \in M_0. \quad (\text{B.2})$$

When $M_0 = \{1, \dots, M\}$, the partially sparse model reduces to a fully sparse model. The proposed method can be naturally applied to partially sparse models as long as the precision matrices to be estimated belong to the subset M_0 . In addition, the structural information is helpful in finding a better

surrogate $\tilde{\Omega}_j$. For example, if Σ_j^* follows the autoregressive model, we only need to estimate two parameters, τ and ρ . However, the iterative methods are not applicable here.

If we generalize $\mathbf{Y}_i, \mathbf{U}_i$ in the linear model B.1 to mode- M tensors and assume ϵ_i follows the tensor normal distribution, $\text{TN}(\mathbf{0}; \Sigma_1^*, \dots, \Sigma_M^*)$, the proposed method can be applied to jointly estimate regression coefficient and conditional covariances in the tensor response regression. As an extension from the multivariate regression with covariance estimation (Rothman et al. 2010), a standard approach is to use negative log-likelihood function with penalties and iterate minimizing it with respect to β and $\{\Omega^*\}$. When β is fixed at previous estimation, the problem reduces to estimating the tensor graphical model. Considering the iterative scheme for updating β and $\{\Omega^*\}$, it is important that we can have a fast algorithm to update $\{\Omega^*\}$ especially in high-dimensional scenarios. Moreover, the partially sparse model can also be assumed when modeling ϵ_i . The proposed method is a better choice than the other iterative methods considering its flexibility, low computational cost and high estimation accuracy.

C Additional Implementation Details and Simulation Results

For the Sequential and Cyclic algorithm, we choose a recent R package Tlasso (which is indeed iteratively optimization with glasso). To choose the tuning parameters, we generate a validation set with the same sample size as the training set. For the separable and oracle estimators, λ_m is chosen to produce the largest log-likelihood calculated using the validation sample and $\{\tilde{\Omega}\}_{-m}$. For the sequential and cyclic estimators, as suggested by Lyu et al. (2019), λ_m is set to be of the form $\eta\sqrt{\log p_m/(npp_m)}$ where η is the tuning parameter. We choose η such that the log-likelihood calculated using the validation sample and $\{\hat{\Omega}\}$ is maximized.

We consider Models 7–12 that have same setting with Models 1–6 except that $n = 20$. We also consider another two models with $n = 20$:

- Model 13: $\Omega_m^*, m = 1, \dots, 4$ are all from the TR covariance model, $(p_1, p_2, p_3, p_4) = (30, 30, 40, 40)$.
- Model 14: $\Omega_m^*, m = 1, \dots, 5$ are all from the TR covariance model, $(p_1, p_2, p_3, p_4, p_5) = (20, 20, 20, 20, 20)$.

	(p_1, p_2, p_3)	Separate	Oracle	Sequential	Cyclic
Model 7	(30,36,30)	0.41 (0.01)	0.32 (0.01)	1.61 (0.01)	7.07 (0.17)
Model 8	(100,100,100)	10.37 (0.01)	8.09 (0.01)	22.40 (0.19)	89.86 (0.06)
Model 9	(5,5,500)	13.62 (0.07)	13.48 (0.05)	12.57 (0.03)	1955.60 (7.77)
Model 13	(30,30,40,40)	9.02 (0.02)	6.97 (0.01)	25.69 (0.02)	101.65 (0.07)
Model 14	(20,20,20,20,20)	17.19 (0.03)	12.72 (0.05)	55.54 (0.28)	224.38 (0.013)

Table 1: Averaged computation cost (in seconds) and standard errors (in parentheses) for Models 7–8 and Models 13–14 from 20 replicates.

Due to the high computational cost of Tlasso algorithm, the maximum number of iteration is set at 5 for Models 3, 6, 9, 12, 13 and 14. The computation time for the fully sparse models is summarized in Table 1. The results of estimation accuracy are summarized in Table 2 and 3. The separable estimator

outperforms the sequential one-iteration estimator and cyclic estimator on every mode and is comparable to the oracle estimator.

(%)	Model 7				Model 8				Model 9				S.E.≤	
	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.		
Frobenius norm loss														
AVG.	1.41	1.37	5.22	1.78	0.49	0.48	2.86	0.61	4.28	3.95	7.38	49.14	(0.02)	
$m = 1$	1.34	1.30	8.29	1.62	0.49	0.48	5.11	0.61	1.18	0.72	1.74	18.55	(0.05)	
$m = 2$	1.49	1.46	5.77	2.06	0.48	0.47	2.88	0.60	1.22	0.68	5.39	18.74	(0.04)	
$m = 3$	1.40	1.36	1.59	1.67	0.49	0.49	0.60	0.62	10.44	10.43	15.02	110.12	(0.03)	
Max norm loss														
AVG.	0.34	0.33	0.65	0.34	0.08	0.08	0.14	0.08	0.61	0.46	1.33	6.65	(0.02)	
$m = 1$	0.33	0.32	1.00	0.33	0.07	0.07	0.22	0.07	0.53	0.33	0.84	8.13	(0.03)	
$m = 2$	0.34	0.34	0.59	0.34	0.07	0.07	0.13	0.07	0.55	0.31	2.37	8.27	(0.03)	
$m = 3$	0.36	0.34	0.37	0.35	0.08	0.08	0.09	0.08	0.76	0.76	0.78	3.56	(0.01)	
True negative rate														
AVG.	66.24	66.05	18.69	40.97	80.42	80.78	25.96	62.57	48.92	46.92	53.52	25.01	(1.31)	
$m = 1$	65.37	65.49	0.08	43.57	80.36	80.97	0.04	62.62	27.58	22.08	39.33	25.00	(2.50)	
$m = 2$	68.05	67.74	2.13	36.20	80.17	80.46	2.76	62.31	25.58	25.08	40.67	50.00	(2.70)	
$m = 3$	65.29	64.92	53.87	43.14	80.74	80.90	75.07	62.79	93.60	93.59	80.55	0.04	(0.43)	
(%)	Model 10				Model 11				Model 12				S.E.≤	
	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.		
Frobenius norm loss														
$m = 2$	1.50	1.47	6.20	3.67	-	-	-	-	-	-	-	-	(0.02)	
$m = 3$	1.35	1.32	1.63	3.05	1.39	1.36	3.13	3.12	10.48	10.47	30.86	115.01	(0.05)	
Max norm loss														
$m = 2$	0.35	0.34	0.66	0.37	-	-	-	-	-	-	-	-	(0.01)	
$m = 3$	0.34	0.33	0.35	0.37	0.35	0.34	0.39	0.39	0.76	0.76	0.83	3.86	(0.01)	
True negative rate														
$m = 2$	68.10	67.97	1.40	0.73	-	-	-	-	-	-	-	-	(0.03)	
$m = 3$	64.58	64.10	46.79	0.99	64.27	65.07	0	0	93.16	93.18	45.82	0	(1.01)	

Table 2: Comparison of means and the maximum standard errors (in parentheses) of different performance measures for Models 7–9 from 100 replicates. All methods have achieved 100% true positive rate (and hence not shown in the table). Comparison of means and the maximum standard errors (in parentheses) of different performance measures for Models 10–12 from 100 replicates. Results for $m = 2$ in Models 11 & 12 are not reported because Ω_2^* is not sparse.

(%)	Model 13				Model 14				S.E. $\leq$
	Sep.	Ora.	Seq.	Cyc.	Sep.	Ora.	Seq.	Cyc.	
Frobenius norm loss									
AVG.	0.22	0.22	1.16	0.29	0.11	0.11	0.64	0.15	(0.00)
$m = 1$	0.20	0.21	2.12	0.24	0.11	0.11	1.36	0.15	(0.01)
$m = 2$	0.20	0.20	1.26	0.24	0.11	0.11	0.83	0.15	(0.01)
$m = 3$	0.24	0.24	0.98	0.35	0.11	0.11	0.54	0.15	(0.00)
$m = 4$	0.24	0.24	0.27	0.35	0.11	0.11	0.31	0.15	(0.00)
$m = 5$	-	-	-	-	0.11	0.11	0.13	0.15	(0.01)
Max norm loss									
AVG.	0.05	0.05	0.14	0.05	0.03	0.03	0.11	0.03	(0.00)
$m = 1$	0.05	0.05	0.24	0.05	0.03	0.03	0.22	0.03	(0.00)
$m = 2$	0.05	0.05	0.15	0.05	0.03	0.03	0.14	0.03	(0.00)
$m = 3$	0.05	0.05	0.10	0.05	0.03	0.03	0.09	0.03	(0.00)
$m = 4$	0.05	0.05	0.06	0.05	0.03	0.03	0.05	0.03	(0.00)
$m = 5$	-	-	-	-	0.03	0.03	0.03	0.03	(0.01)
True negative rate									
AVG.	66.9	67.54	15.89	39.64	58.44	56.63	9.37	23.42	(0.44)
$m = 1$	65.02	65.39	0.00	45.34	58.04	56.53	0.00	23.21	(0.87)
$m = 2$	64.02	66.34	0.07	44.82	58.83	56.63	0.00	23.49	(0.82)
$m = 3$	69.54	69.37	2.60	34.35	58.25	56.36	0.08	23.89	(0.74)
$m = 4$	69.02	69.05	60.88	34.04	58.36	56.7	2.53	22.9	(0.85)
$m = 5$	-	-	-	-	58.71	56.94	44.25	23.61	(1.65)

Table 3: Comparison of means and the maximum standard errors (in parentheses) of different performance measures for Models 13–14 from 100 replicates. All methods have achieved 100% true positive rate (and hence not shown in the table).

D Theoretical Analysis

D.1 Proof of Lemma 1

The first part of the results is analogous the Lemma 1 in Pan et al. (2019). For completeness, we provide a similar proof. By stacking the n observations together, we denote $\mathbf{X}_{(j)} = (\mathbf{X}_{1(j)}, \dots, \mathbf{X}_{n(j)})$. Then we have $\sum_{i=1}^n \mathbf{X}_{i(j)}(\mathbf{X}_{i(j)})^T = \mathbf{X}_{(j)}\mathbf{X}_{(j)}^T$. Since $\mathbf{X}_{(j)}\mathbf{X}_{(j)}^T$ is positive semi-definite and $\text{rank}(\mathbf{X}_{(j)}\mathbf{X}_{(j)}^T) = \text{rank}(\mathbf{X}_{(j)})$, it remains to show that $\text{rank}(\mathbf{X}_{(j)}) = p_j$. Because $\mathbf{X}_1, \dots, \mathbf{X}_n$ are independent samples, when $np_{-j} > p_j$, $\mathbf{X}_{(j)}$ is row full-rank with probability 1. Therefore, $\text{rank}(\mathbf{X}_{(j)}) = p_j$. Thus, $\sum_{i=1}^n \mathbf{X}_{i(j)}(\mathbf{X}_{i(j)})^T$ is positive definite with probability 1. The rest of Lemma 1 can be easily verified by Lemma S.2 in Lyu et al. (2019).

D.2 Proof of Theorem 1

For technical reasons, we split the data so that estimating $\tilde{\Omega}_j (j \neq m)$ using (4) in the first step and formulating the graphical lasso problem in the second step use separate and independent datasets. Without loss of generality, we assume that the sample size for each subset is n . The splitting procedure is not included in our numerical studies. Our simulation results also show that sample splitting is not necessary for practical use.

For a list of fixed symmetric and positive definite matrices $\{\alpha\}_{-m}$ where $\alpha_j \in \mathbb{R}^{p_j \times p_j} (j \neq m)$, we define the population objective function for Ω_m corresponding to (3) as

$$L_m(\Omega_m; \{\alpha\}_{-m}) = \frac{1}{p} \left\{ \text{tr} \left(E \left(\text{vec}(\mathbf{X}) \text{vec}(\mathbf{X})^T \right) (\alpha_M \otimes \cdots \otimes \alpha_{m+1} \otimes \Omega_m \otimes \alpha_{m-1} \otimes \cdots \otimes \alpha_1) \right) \right\} - \frac{1}{p_m} \log |\Omega_m|.$$

Define the population minimization function to $L_m(\Omega_m; \{\alpha\}_{-m})$ with parameter $\{\alpha\}_{-m}$ as

$$M_m(\{\alpha\}_{-m}) = \arg \min_{\Omega_m} L_m(\Omega_m; \{\alpha\}_{-m}). \quad (\text{D.1})$$

Define $\mathbb{B}(\Omega_m^*)$ as the set containing Ω_m^* and its neighborhood for some sufficiently large radius $r > 0$,

$$\mathbb{B}(\Omega_m^*) = \{\Omega \in \mathbb{R}^{p_m \times p_m} : \Omega = \Omega^T; \Omega \succ 0; \|\Omega - \Omega_m^*\|_F \leq r\}.$$

We first present several useful technical results.

Proposition D.1. *For any fixed $\alpha_j (j \neq m)$ satisfying $\text{tr}(\Sigma_j^* \alpha_j) \neq 0$, the population minimization function defined in (D.1) satisfies $M_m(\{\alpha\}_{-m}) = p_{-m} \left\{ \prod_{j \neq m} \text{tr}(\Sigma_j^* \alpha_j) \right\}^{-1} \Omega_m^*$.*

Proposition D.1 can be easily verified by Theorem 3.1 in Lyu et al. (2019). By Proposition D.1, $M_m(\{\alpha\}_{-m})$ is proportional to Ω_m^* . Since we standardize both $M_m(\{\alpha\}_{-m})$ and $\widehat{M}_m(\{\alpha\}_{-m})$ to have unit Frobenius norm, the idea of proving Theorem 1 is to derive the upper bound for the difference between $M_m(\{\alpha\}_{-m})$ and $\widehat{M}_m(\{\alpha\}_{-m})$.

The proof of Theorem 1 mainly relies on the following proposition which is obtained from the proof of Theorem 3.4 and Theorem 3.5 in Lyu et al. (2019). Denote $\mathbf{S}'_m = \frac{1}{np_{-m}} \sum_{i=1}^n \mathbf{V}_i^m (\mathbf{V}_i^m)^T$ and $\mathbf{V}_i^m = \mathbf{X}_{i(m)} (\alpha_m^{1/2} \otimes \cdots \otimes \alpha_{m+1}^{1/2} \otimes \alpha_{m-1}^{1/2} \cdots \otimes \alpha_1^{1/2})$.

Proposition D.2 (Lyu et al. (2019), Theorems 3.4 and 3.5). *Suppose that Conditions (C1) and (C2) hold. For any fixed symmetric and positive definite matrix $\{\alpha\}_{-m}$ satisfying that either $\alpha_j \in \mathbb{B}(\Omega_j^*)$ or $C_3 \leq \text{tr}(\Sigma_j^* \alpha_j) / p_j \leq 1/C_3$ and $\|\alpha_j\|_2 \leq C_4$ for some positive constants C_3 and C_4 , we have*

$$P \left(\left\| \widehat{M}_m(\{\alpha\}_{-m}) - M_m(\{\alpha\}_{-m}) \right\|_F \geq \gamma_1 \sqrt{\frac{(p_m + s_m) \log p_m}{np_{-m}}} \right) \rightarrow 0$$

where $\widehat{M}_m(\cdot)$ and $M_m(\cdot)$ are defined in (D.1) and (5). Moreover, the estimator $\widehat{\Omega}_m$ satisfies

$$P \left(\left\| \widehat{\Omega}_m - \Omega_m^* \right\|_F \geq \gamma_2 \sqrt{\frac{(p_m + s_m) \log p_m}{np_{-m}}} \right) \rightarrow 0$$

where $\widehat{\Omega}_m = \widehat{M}_m(\{\alpha\}_{-m}) / \|\widehat{M}_m(\{\alpha\}_{-m})\|_F$ and γ_1, γ_2 are positive constants depending on r, C_1, C_2, C_3, C_4 .

Proof of Theorem 1. We first check the conditions on $\widetilde{\Omega}_j$ for $j \neq m$.

When $\widetilde{\Omega}_j = \mathbf{I}_{p_j}$, under Condition (C1), we have

$$0 < C_1 \leq \text{tr}(\Sigma_j^* \widetilde{\Omega}_j) / p_j = \text{tr}(\Sigma_j^*) / p_j \leq 1/C_1 < \infty.$$

In addition, $\|\tilde{\Omega}_j\|_2 = 1 < \infty$.

When $\tilde{\Omega}_j$ is obtained from (4), since $\|\tilde{\Omega}_j\|_F = 1$, it is easy to check that

$$\|\tilde{\Omega}_j - \Omega_j^*\|_F \leq \|\tilde{\Omega}_j\|_F + \|\Omega_j^*\|_F = 2.$$

Therefore, we have $\tilde{\Omega}_j \in \mathbb{B}(\Omega_j^*)$.

We can let $C_3 = C_1$, $C_4 = 1$ and denote the set $\mathbb{B}_1(\Omega_j^*)$ as

$$\mathbb{B}_1(\Omega_j^*) = \{\Omega \in \mathbb{R}^{p_j \times p_j} : \Omega = \Omega^T; \Omega \succ 0; C_1 \leq \text{tr}(\Sigma_j^* \Omega) / p_j \leq 1/C_1; \|\Omega\|_2 \leq 1\} \cup \mathbb{B}(\Omega_j^*).$$

Then we have

$$P(\tilde{\Omega}_j \in \mathbb{B}_1(\Omega_j^*)) = 1.$$

By taking $r = 1$, the constants γ_1, γ_2 in Proposition D.2 only depend on C_1, C_2 . Therefore, by applying Proposition D.2 we have

$$\begin{aligned} P\left(\left\|\widehat{M}_m(\{\tilde{\Omega}\}_{-m}) - M_m(\{\tilde{\Omega}\}_{-m})\right\|_F \geq C_5 \sqrt{\frac{(p_m + s_m) \log p_m}{np_{-m}}}\right) \\ = P\left(\left\|\widehat{M}_m(\{\tilde{\Omega}\}_{-m}) - M_m(\{\tilde{\Omega}\}_{-m})\right\|_F \geq C_5 \sqrt{\frac{(p_m + s_m) \log p_m}{np_{-m}}} \mid \tilde{\Omega}_j \in \mathbb{B}_1(\Omega_j^*)\right) \rightarrow 0, \end{aligned}$$

where C_5 is a positive constant depending on C_1 and C_2 . This implies

$$\left\|\widehat{M}_m(\{\tilde{\Omega}\}_{-m}) - M_m(\{\tilde{\Omega}\}_{-m})\right\|_F = O_P\left(\sqrt{\frac{(p_m + s_m) \log p_m}{np_{-m}}}\right).$$

Similarly, we can obtain

$$\left\|\widehat{\Omega}_m - \Omega_m^*\right\|_F = O_P\left(\sqrt{\frac{(p_m + s_m) \log p_m}{np_{-m}}}\right).$$

□

References

- Diggle, P., Diggle, P. J., Heagerty, P., Liang, K. Y., Heagerty, P. J. & Zeger, S. (2002), *Analysis of longitudinal data*, Oxford University Press.
- Fitzmaurice, G. M., Laird, N. M. & Ware, J. H. (2012), *Applied longitudinal analysis*, Vol. 998, John Wiley & Sons.
- Genton, M. G. (2007), ‘Separable approximations of space-time covariance matrices’, *Environmetrics: The official journal of the International Environmetrics Society* **18**(7), 681–695.
- George, B. & Aban, I. (2015), ‘Selecting a separable parametric spatiotemporal covariance structure for longitudinal imaging data’, *Statistics in medicine* **34**(1), 145–161.
- George, B., Denney Jr, T., Gupta, H., Dell’Italia, L. & Aban, I. (2016), ‘Applying a spatiotemporal model for longitudinal cardiac imaging data’, *The annals of applied statistics* **10**(1), 527.

- Guttorp, P. & Gneiting, T. (2006), ‘Studies in the history of probability and statistics xlix on the matern correlation family’, *Biometrika* **93**(4), 989–995.
- Lyu, X., Sun, W. W., Wang, Z., Liu, H., Yang, J. & Cheng, G. (2019), ‘Tensor graphical model: Non-convex optimization and statistical inference’, *IEEE transactions on pattern analysis and machine intelligence* .
- Matérn, B. (2013), *Spatial variation*, Vol. 36, Springer Science & Business Media.
- Pan, Y., Mai, Q. & Zhang, X. (2019), ‘Covariate-adjusted tensor classification in high dimensions’, *Journal of the American statistical association* **114**(527), 1305–1319.
- Rothman, A. J., Levina, E. & Zhu, J. (2010), ‘Sparse multivariate regression with covariance estimation’, *Journal of Computational and Graphical Statistics* **19**(4), 947–962.