

Statistica Sinica Preprint No: SS-2023-0131	
Title	Adaptive Block Banding Precision Matrix Estimation For Multivariate Longitudinal Data
Manuscript ID	SS-2023-0131
URL	http://www.stat.sinica.edu.tw/statistica/
DOI	10.5705/ss.202023.0131
Complete List of Authors	Chunhui Liang, Wenqing Ma and Yanyuan Ma
Corresponding Authors	Wenqing Ma
E-mails	wenqingma@cnu.edu.cn

ADAPTIVE BLOCK BANDING PRECISION MATRIX ESTIMATION FOR MULTIVARIATE LONGITUDINAL DATA

Chunhui Liang¹, Wenqing Ma² and Yanyuan Ma³

¹ *Northeast Normal University*

² *Capital Normal University*

³ *The Pennsylvania State University*

Abstract: We propose an estimator for precision matrices with the structure of Banded Kronecker Sparse forms (BKS). BKS takes advantage of the special feature of a precision matrix, which has the form of the kronecker product of an adaptively banded matrix and a sparse matrix, both are positive definite. Such precision matrix arises frequently in practice in finance data, medical data and time series data. We achieve the adaptive bandedness via a specially designed penalty, and enforce the sparsity via lasso. We apply a computationally efficient procedure named Alternative Convex Search (ACS) algorithm to implement BKS. We establish the computational convergence and show the statistical guarantee through establishing the asymptotic rate. Our extensive simulation studies indicate the superior finite sample performance of BKS in comparison to existing methods. Additionally, we apply BKS to EEG and ADHD datasets, wherein it outperforms other methods in capturing the banding sparsity characteristics of the precision matrix.

Key words and phrases: Adaptive banding structure; Biconvex function; Convex optimization; Kronecker product; Sparse precision matrix.

1. Introduction

Matrix-valued data are becoming increasingly common in modern data collection procedures. This type of data can be found, for instance, in neuro image data, finance data, and time series data. In this paper, we focus on multivariate longitudinal data, which is a special type of matrix-valued data. This type of data contains multiple outcomes of interest for each subject, with repeated measurements taken over time. It is natural to represent the resulting data in a matrix format with two dimensions corresponding to variables and time. We are aware that the estimation of covariance or precision matrices is a fundamental problem in multivariate data analysis, including techniques such as principle component analysis, discriminant analysis, and regression analysis. Our objective is to estimate the precision matrix of the vectorized version of multivariate longitudinal data. However, considering the characteristics of multivariate longitudinal data, the covariance matrix in the two dimensions of the observed matrix contains different structural information. It not only incorporates the structural information between multiple outcomes at a fixed time point but also incorporates the time-series correlation structure among different time points for each response variable. The precision matrix exhibits the same characteristics. Due to the presence of complex structural information and the typically large dimension of the observed matrices, obtaining an efficient estimator for the precision matrix becomes increasingly challenging as the number of unknown parameters increases quadratically with the vector length.

To address the challenges posed by high-dimensionality when estimating a large di-

mensional precision matrix based on high-dimensional vector observations, the existing literature has proposed two general approaches by incorporating sparse structural assumptions. The first approach involves directly imposing sparsity on the precision matrix through methods, like the graphical lasso (e.g., Yuan and Lin (2007); Banerjee et al. (2008); Friedman et al. (2008)). These methods are suitable for estimating an unstructured precision matrix. In the case of variables with a natural ordering, such as time-series data, the second approach involves directly imposing a banded structure on the precision matrix. This approach has been explored in works like Yu and Bien (2017) and Furrer and Bengtsson (2007). The asymptotic validity of the inversion procedure was later established by Bickel and Levina (2008b) under the assumption of equal bandwidth for all rows and by Cai et al. (2010) under a general bandedness assumption. However, these methods treat the vector data as a whole and assume a sparse structure based on the characteristics of the vector data. Consequently, they are not suitable for modeling multivariate longitudinal data. This is because the aforementioned sparse structures cannot simultaneously capture the correlation structures among response variables and observed time points.

To estimate the precision matrix for multivariate longitudinal data, various methods have been proposed. These methods have shown success when the dimension of the observed matrix is not very large ((Kim and Zimmerman, 2012; Lee et al., 2020)). One approach involves utilizing modified Cholesky block decomposition to reparameterize the covariance. The purpose of this decomposition is to ensure positive definiteness, rather

than reducing the number of parameters. However, when dealing with high-dimensional data, structural assumptions are necessary to reduce the complexity of the model. Taking into consideration the characteristic that the correlation between measurements for any two time points decays with increasing time distance, Qian et al. (2020) and Qian et al. (2021) propose a regularized estimator for the precision matrix. This is achieved through a modified Cholesky block decomposition and by penalizing the log-likelihood with a penalty function that encourages a block banded structure in the lower triangle block matrix. However, the adaptive block banded precision matrix estimator (ABR) introduced by Qian et al. (2021) has a significant computational drawback. The running time of ABR increases dramatically as the number of rows or columns in the observation matrix increases. Additionally, the theoretical properties of ABR are based on the assumption that the repeated measurements are finite.

To address the challenges involved in estimating the precision matrix for high-dimensional multivariate longitudinal data, it is common to introduce a separability assumption on the covariance matrix. This assumption represents the precision matrix as a Kronecker product structure with two smaller matrices. To capture various types of structural information, different approaches have been proposed for these smaller matrices (Tsiligkaridis and Hero (2013), Greenewald and Hero (2015), Leng and Tang (2012), Leng and Pan (2018), Zhang et al. (2023), Dai et al. (2023)). Tsiligkaridis and Hero (2013) and Greenewald and Hero (2015) assume both matrices to be low-rank, Leng and Tang (2012) assume both matrices to be sparse, and further assume normality and propose the Sparse

Matrix Graphical Model (SMGM) estimator by penalizing the log-likelihood. However, these methods only consider sparsity in the precision matrix or correlation matrix, overlooking the potential bandedness property. Zhang et al. (2023) and Dai et al. (2023) assume a banded structure for both of these matrices. These methods are particularly suitable for space-time data where both the row and column variables in the observed matrix have a natural ordering. In this paper, we address the situation where the observed variables of interest may not have a natural ordering. To capture this complexity, we assume that one of the two matrices is sometimes sparse, while the other is banded. For instance, each column of the original matrix data may have a sparse precision matrix, such as in cases where a matrix column represents measurements at different brain locations. Conversely, different columns may correspond to measurements taken at different times, resulting in a banded precision matrix due to the decreasing time relation. Estimating a large-dimensional precision matrix is a challenging task due to the quadratic increase in the number of unknown parameters along the vector length. However, the kronecker product form mentioned earlier effectively reduces the number of parameters and enables contemporary methods to simultaneously consider bandedness and sparsity.

The remainder of this paper is structured as follows. Section 2 introduces the specific model setting, as well as the BKS estimator and the ACS algorithm. The algorithmic convergence of ACS is also established in this section. The theoretical properties of BKS are provided in Section 3. Section 4 presents simulation studies, while Section 5 presents real data analysis. Lastly, Section 6 offers concluding remarks. For the proofs

and technical derivations, please refer to the Supplement Materials.

2. Model and Estimation

In this section, we will provide a description of the model and the motivation behind our statistical model. Additionally, we delve into the computational aspects of the estimation procedure.

2.1 Model setup

Let $\mathbf{Y}_i \equiv (\mathbf{Y}_{i1.}, \dots, \mathbf{Y}_{iJ.})$ be the $K \times J$ random matrix associated with individual i across all time. We assume $\mathbf{Y}_1, \dots, \mathbf{Y}_n$ are independent and identically distributed (iid). Each matrix $\mathbf{Y}_i$ is vectorized to form $\text{vec}(\mathbf{Y}_i) \equiv (\mathbf{Y}_{i1.}^T, \dots, \mathbf{Y}_{iJ.}^T)^T \in \mathcal{R}^{KJ}$, and we assume that the mean and covariance of $\text{vec}(\mathbf{Y}_i)$ are $E\{\text{vec}(\mathbf{Y}_i)\} = \mathbf{0}$ and $\text{cov}\{\text{vec}(\mathbf{Y}_i)\} = \Sigma$, respectively. The covariance matrix Σ captures the correlations between any two responses at different times, including the temporal correlation for a fixed response variable $\text{cor}(Y_{ijk}, Y_{ij'k})$, the variable correlation for a fixed time point $\text{cor}(Y_{ijk}, Y_{ijk'})$, and the correlation between any two response variables at different time points $\text{cor}(Y_{ijk}, Y_{ij'k'})$, where $j \neq j' \in \{1, \dots, J\}$, $k \in \{1, \dots, K\}$. Let us divide Σ into J^2 size $K \times K$ block matrices. We denote the (j, l) block as $\Sigma^{(jl)}$, where j and l range from 1 to J . The diagonal block matrix $\Sigma^{(jj)} = \text{cov}(\mathbf{Y}_{ij.})$ represents the variance-covariance structure between the K responses at the j th time point. Similarly, $\Sigma^{(jl)} = \text{cov}(\mathbf{Y}_{ij.}, \mathbf{Y}_{il.})$ denotes the covariance between the K responses at the j th and the l th time points. We assume that

the correlation structure information can be separated into two dimensions: variable and time. Specifically, the assumed temporal correlation is the same for all responses, and the assumed variable correlation is the same for all time points. Thus, the covariance matrix for multivariate longitudinal data can be represented as a Kronecker product, $\Sigma = \mathbf{R} \otimes \mathbf{W}$, where $\mathbf{R} \in \mathcal{R}^{J \times J}$ and $\mathbf{W} \in \mathcal{R}^{K \times K}$. Under this assumption, we can write $\Sigma^{(jl)}$ as $r_{jl} \mathbf{W}$ for all j, l range from 1 to J . Here, r_{jl} represents a constant that denotes the signal amplification at different time points. In this case, the Kronecker structure $\mathbf{R} \otimes \mathbf{W}$ is non-identifiability. For convenience, we set $\mathbf{W} = \Sigma^{(11)} / \Sigma_{1,1}$, where $\Sigma_{1,1}$ represents the entry at position $(1, 1)$ in the matrix Σ . Our interest is in estimating the precision matrix $\mathbf{R}^{-1} \otimes \mathbf{W}^{-1}$. However, in high-dimensional situations, the number of variables and time points both may exceed the sample size. To obtain a stable and efficient estimator for the precision matrix, it is necessary to introduce additional structure assumptions for $\mathbf{R}^{-1}$ and $\mathbf{W}^{-1}$. To obtain a positive definite estimator for $\mathbf{R}^{-1}$, we perform a cholesky decomposition such that $\mathbf{R}^{-1} \equiv \mathbf{L}^T \mathbf{L}$, where $\mathbf{L}$ is a lower triangle matrix and $L_{j,j} > 0$, $j = 1, \dots, J$. The elements in $\mathbf{R}^{-1}$ represent the conditional correlation between any two time points, given the other observed time points, for the response variable. In practice, as the distance between two time points increases, the conditional correlation will decrease. Therefore, we assume that $\mathbf{R}^{-1}$ is a banded matrix with a bandwidth of d , where d is significantly smaller than J . The elements in $\mathbf{W}^{-1}$ represent the conditional correlation between any two response variables, given the other variables for a specific time point. For $\mathbf{W}^{-1}$, we assume that it is a sparse matrix. Furthermore, we provide

a detailed introduction to the Kronecker structure through two examples in Supplement Material S5.

Remark 1. The assumption that the true covariance or precision matrix is separable plays a crucial role in our model framework and should be evaluated during the data preprocessing stage. In this study, we follow Zhang et al. (2023) and employ the projection-based bootstrap test method introduced by Aston et al. (2017). This method is theoretically guaranteed and computationally fast in high-dimensional settings. Moreover, as a distribution-free approach, it is suitable for our framework.

2.2 Estimation

To estimate $\mathbf{R}^{-1} \otimes \mathbf{W}^{-1}$, we only need to estimate $\mathbf{R}^{-1}$ and $\mathbf{W}^{-1}$ based on the iid observations $\mathbf{Y}_1, \dots, \mathbf{Y}_n$. We consider minimizing the target function

$$l(\mathbf{W}^{-1}, \mathbf{R}^{-1}) = -J \log |\mathbf{W}^{-1}| - K \log |\mathbf{R}^{-1}| + \frac{1}{n} \sum_{i=1}^n \text{tr}(\mathbf{Y}_i^T \mathbf{W}^{-1} \mathbf{Y}_i \mathbf{R}^{-1}).$$

Obviously, up to a constant, $l(\mathbf{W}^{-1}, \mathbf{R}^{-1})$ is the negative loglikelihood of $(\mathbf{W}^{-1}, \mathbf{R}^{-1})$ under the assumption that $\text{vec}(\mathbf{Y}_i)$'s are normally distributed with mean zero. However, we do not assume normality here, so we view $l(\mathbf{W}^{-1}, \mathbf{R}^{-1})$ as a general loss function.

We incorporate a lasso penalty to take into account of the sparsity of $\mathbf{W}^{-1}$. Specifically, we add the penalty term $\lambda_1 \|\text{vec}(\mathbf{W}^{-1})\|_1$ to the loss function. To account for the positive definiteness and banded structure of $\mathbf{R}^{-1}$, we adopt the same methodology as Yu and Bien (2017). Begin with, we utilize the Cholesky decomposition of $\mathbf{R}^{-1}$ to express

it as $\mathbf{R}^{-1} = \mathbf{L}^T \mathbf{L}$, where $\mathbf{L}$ is a lower triangle matrix with $L_{j,j} > 0$ and $L_{j,l} = 0$ for all $j - l > d$. To incorporate the banded structure for $\mathbf{L}$, we consider each row of $\mathbf{L}$ separately. In the lower triangle form of $\mathbf{L}$, the j th row exclusively includes potentially non-zero elements $L_{j,1}, \dots, L_{j,j-1}$ and a positively definite element $L_{j,j}$. Roughly speaking, the banded structure implies that a smaller column index, denoted by l , suggests a higher probability for the entry $L_{j,l}$ to be zero. Thus to encourage more likely zeros corresponding to smaller column index l , we consider the penalty

$$p(\mathbf{L}_{j,\cdot}) \equiv \sum_{l=1}^{j-1} \sqrt{\sum_{q=1}^l L_{j,q}^2}. \quad (2.1)$$

We can see that this is actually a group lasso penalty, where the l th group is the subvector formed by the first l elements of the j th row. Because of the relation $(\sum_{q=1}^l L_{j,q}^2)^{1/2} \leq (\sum_{q=1}^{l+1} L_{j,q}^2)^{1/2}$, a zero value corresponding to an index l automatically implies a zero value for all indices $< l$. In other words, the sparsity penalty in (2.1) automatically leads to a banded structure on the j th row. Taking into account all rows, we thus incorporate a penalty $\lambda_2 \sum_{j=2}^J p(\mathbf{L}_{j,\cdot})$ to factor in the banded structure on $\mathbf{L}$, or equivalently, the banded structure on $\mathbf{R}^{-1}$.

Combining the above analysis, we propose to estimate $\mathbf{W}^{-1}$ and $\mathbf{L}$ through minimizing

$$\begin{aligned} Q(\mathbf{W}^{-1}, \mathbf{L}, \boldsymbol{\lambda}) = & -J \log |\mathbf{W}^{-1}| - 2K \log |\mathbf{L}| + \frac{1}{n} \sum_{i=1}^n \text{tr}(\mathbf{Y}_i^T \mathbf{W}^{-1} \mathbf{Y}_i \mathbf{L}^T \mathbf{L}) \\ & + \lambda_1 J \|\text{vec}(\mathbf{W}^{-1})\|_1 + \lambda_2 K \sum_{j=2}^J p(\mathbf{L}_{j,\cdot}) \end{aligned} \quad (2.2)$$

subject to the positive-definite constraint on $\mathbf{W}^{-1}$, where $p(\mathbf{L}_{j,\cdot})$ is defined in (2.1), and $\boldsymbol{\lambda} \equiv (\lambda_1, \lambda_2)^T$ contains the turning parameters.

At any $\boldsymbol{\lambda}$, $Q(\mathbf{W}^{-1}, \mathbf{L}, \boldsymbol{\lambda})$ is biconvex in $(\mathbf{W}^{-1}, \mathbf{L})$, i.e., it is a convex function of $\mathbf{W}^{-1}$ when $\mathbf{L}$ is fixed, and is a convex function of $\mathbf{L}$ when $\mathbf{W}^{-1}$ is fixed. We thus solve the optimization problem in (2.2) by alternate convex search (ACS), i.e., we alternately minimize $Q(\mathbf{W}^{-1}, \mathbf{L}, \boldsymbol{\lambda})$ with respect to $\mathbf{W}^{-1}$ or $\mathbf{L}$ while keeping the other matrix fixed. Specifically, at the s th step, we first fix $\mathbf{L}$ at $\hat{\mathbf{L}}^{(s)}$, and update the estimator of $\mathbf{W}^{-1}$ by solving

$$(\hat{\mathbf{W}}^{-1})^{(s+1)} = \underset{\mathbf{W}^{-1} > 0}{\operatorname{argmin}} \left\{ -\log |\mathbf{W}^{-1}| + \frac{1}{nJ} \sum_{i=1}^n \operatorname{tr} \{ \mathbf{Y}_i^T \mathbf{W}^{-1} \mathbf{Y}_i (\hat{\mathbf{L}}^{(s)})^T \hat{\mathbf{L}}^{(s)} \} + \lambda_1 \|\operatorname{vec}(\mathbf{W}^{-1})\|_1 \right\}, \quad (2.3)$$

where $\mathbf{W}^{-1} > 0$ means $\mathbf{W}^{-1}$ is positive definite. We then fix $\mathbf{W}^{-1}$ at $(\hat{\mathbf{W}}^{-1})^{(s+1)}$ and update the estimator of $\mathbf{L}$ by minimizing

$$-2 \log |\mathbf{L}| + \frac{1}{nK} \sum_{i=1}^n \operatorname{tr} \{ \mathbf{Y}_i^T (\hat{\mathbf{W}}^{-1})^{(s+1)} \mathbf{Y}_i \mathbf{L}^T \mathbf{L} \} + \lambda_2 \sum_{j=1}^J p(\mathbf{L}_{j,\cdot}), \quad (2.4)$$

subject to $L_{j,j} > 0$, and $L_{j,j'} = 0, 1 \leq j < j' \leq J$. We repeat the above optimization steps (2.3)-(2.4) until $\|(\hat{\mathbf{W}}^{-1})^{(s+1)} - (\hat{\mathbf{W}}^{-1})^{(s)}\|_F + \|\hat{\mathbf{L}}^{(s+1)} - \hat{\mathbf{L}}^{(s)}\|_F < \varepsilon$, where ε is a pre-determined sufficiently small constant. We then set $\hat{\mathbf{W}}^{-1} = (\hat{\mathbf{W}}^{-1})^{(s+1)}$ and $\hat{\mathbf{L}} = \hat{\mathbf{L}}^{(s+1)}$ as the final estimators.

Next, we discuss the details in solving (2.3) and (2.4) respectively. (2.3) is a well studied problem (Yuan and Lin, 2006; Rothman et al., 2008) and here, we adopt the graphical

lasso (glasso) algorithm (Friedman et al., 2008), which guarantees the positive definiteness of the matrix $(\widehat{\mathbf{W}}^{-1})^{(s+1)}$. To investigate the minimization problem of (2.4), we write $\mathbf{Y}_i^* \equiv ((\widehat{\mathbf{W}}^{-1})^{(s+1)})^{1/2} \mathbf{Y}_i$, $\mathbf{Y}^* \equiv \{(\mathbf{Y}_1^*)^T, \dots, (\mathbf{Y}_n^*)^T\}^T$, and $\mathbf{L}_j \equiv \mathbf{L}_{j,1:j}^T$ as a j -dimensional column vector formed by the first j elements on the j th row of the matrix $\mathbf{L}$ for a generic matrix $\mathbf{L}$. We then obtain $\sum_{i=1}^n \text{tr}(\mathbf{Y}_i^T (\widehat{\mathbf{W}}^{-1})^{(s+1)} \mathbf{Y}_i \mathbf{L}^T \mathbf{L}) = \sum_{i=1}^n \text{tr}((\mathbf{Y}_i^*)^T (\mathbf{Y}_i^*) \mathbf{L}^T \mathbf{L}) = \sum_{j=1}^J \|\mathbf{Y}_{:,1:j}^* \mathbf{L}_j\|_2^2$. Thus (2.4) can be equivalently written as

$$-2 \sum_{j=1}^J \log L_{j,j} + \frac{1}{nK} \sum_{j=1}^J \|\mathbf{Y}_{:,1:j}^* \mathbf{L}_j\|_2^2 + \lambda_2 \sum_{j=1}^J p(\mathbf{L}_{j,\cdot}).$$

We can now decompose the optimization with respect to $\mathbf{L}$ into J separately optimization problems with respect to $\mathbf{L}_{1,\cdot}, \dots, \mathbf{L}_{J,\cdot}$ respectively. Specifically,

$$\begin{aligned} \widehat{L}_{1,1}^{(s+1)} &= \underset{L_{1,1} > 0}{\text{argmin}} \{-2K \log L_{1,1} + \frac{1}{n} \|\mathbf{Y}_{:,1}^* L_{1,1}\|_2^2\} = \{(\mathbf{Y}_{:,1}^*)^T \mathbf{Y}_{:,1}^* / nK\}^{-1/2}, \\ \widehat{\mathbf{L}}_j^{(s+1)} &= \underset{L_{j,j} > 0, \mathbf{L}_j \in \mathcal{R}^j}{\text{argmin}} \{-2 \log L_{j,j} + \frac{1}{nK} \|\mathbf{Y}_{:,1:j}^* \mathbf{L}_j\|_2^2 + \lambda_2 p(\mathbf{L}_j)\}, \quad j = 2, \dots, J. \end{aligned} \quad (2.5)$$

Note that $p(\mathbf{L}_{j,\cdot}) = p(\mathbf{L}_j)$. To solve each of the $J - 1$ optimization problems above, we adopt the alternating direction method of multipliers (ADMM) algorithm (Boyd et al. (2011)). To this end, to obtain $\widehat{\mathbf{L}}_j^{(s+1)}$, we introduce the constraints $\mathbf{L}_j = \boldsymbol{\Psi}_j$, and modify the objective function for $(\mathbf{L}_j, \boldsymbol{\Psi}_j)$ into

$$\begin{aligned} Q^*(\mathbf{L}_j, \lambda_2, \mathbf{U}_j, \rho, \boldsymbol{\Psi}_j) &= -2 \log L_{j,j} + \frac{1}{nK} \|\mathbf{Y}_{:,1:j}^* \mathbf{L}_j\|_2^2 \\ &\quad + \lambda_2 p(\boldsymbol{\Psi}_j) + \mathbf{U}_j^T (\mathbf{L}_j - \boldsymbol{\Psi}_j) + \frac{\rho}{2} \|\mathbf{L}_j - \boldsymbol{\Psi}_j\|_2^2. \end{aligned} \quad (2.6)$$

Given $\widehat{\Psi}_j^{(r)}$ and $\widehat{U}_j^{(r)}$, we compute the derivative of (2.6) with respect to $\mathbf{L}_j$ and obtain the estimation equation:

$$-2\frac{1}{L_{j,j}}\mathbf{e}_j + \frac{2}{nK}(\mathbf{Y}_{\cdot,1:j}^*)^T \mathbf{Y}_{\cdot,1:j}^* \mathbf{L}_j + \widehat{U}_j^{(r)} + \rho(\mathbf{L}_j - \widehat{\Psi}_j^{(r)}) = \mathbf{0},$$

where $\mathbf{e}_j$ denotes the j -dimensional indicator vector, with 1 on the j th element and the other elements 0. The above estimation equation can be written as

$$\begin{cases} \mathbf{L}_{j,1:j-1} \left\{ \frac{2}{nK}(\mathbf{Y}_{\cdot,1:j-1}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* + \rho \mathbf{I} \right\} + \frac{2}{nK}(\mathbf{Y}_{\cdot,j}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* L_{j,j} + \widehat{U}_{j,1:j-1}^{(r)} = \rho \widehat{\Psi}_{j,1:j-1}^{(r)}, \\ -\frac{2}{L_{j,j}} + \left\{ \frac{2}{nK}(\mathbf{Y}_{\cdot,j}^*)^T \mathbf{Y}_{\cdot,j}^* + \rho \right\} L_{j,j} + \frac{2}{nK}(\mathbf{Y}_{\cdot,j}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* \mathbf{L}_{j,1:j-1}^T + \widehat{U}_{j,j}^{(r)} = \rho \widehat{\Psi}_{j,j}^{(r)}. \end{cases}$$

The first equation leads to that $\mathbf{L}_{j,1:j-1} = -\{2(nK)^{-1}(\mathbf{Y}_{\cdot,j}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* L_{j,j} + \widehat{U}_{j,1:j-1}^{(r)} - \rho \widehat{\Psi}_{j,1:j-1}^{(r)}\} \{2(nK)^{-1}(\mathbf{Y}_{\cdot,1:j-1}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* + \rho \mathbf{I}\}^{-1}$. Inserting this into the second equation yields an equation of the form $AL_{j,j}^2 + BL_{j,j} + 2 = 0$, where

$$\begin{aligned} A &= \left\{ \frac{4}{nK}(\mathbf{Y}_{\cdot,j}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* \right\} \left\{ \frac{2}{nK}(\mathbf{Y}_{\cdot,1:j-1}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* + \rho \mathbf{I} \right\}^{-1} \\ &\quad \left\{ \frac{1}{nK}(\mathbf{Y}_{\cdot,1:j-1}^*)^T \mathbf{Y}_{\cdot,j}^* \right\} - \frac{2}{nK}(\mathbf{Y}_{\cdot,j}^*)^T \mathbf{Y}_{\cdot,j}^* - \rho, \\ B &= \left\{ \frac{2}{nK}(\mathbf{Y}_{\cdot,j}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* \right\} \left\{ \frac{2}{nK}(\mathbf{Y}_{\cdot,1:j-1}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* + \rho \mathbf{I} \right\}^{-1} \\ &\quad (\widehat{U}_{j,1:j-1}^{(r)} - \rho \widehat{\Psi}_{j,1:j-1}^{(r)}) - \widehat{U}_{j,j}^{(r)} + \rho \widehat{\Psi}_{j,j}^{(r)}. \end{aligned}$$

Note that $-1/A$ is the lower-right entry of the matrix $\{(\mathbf{Y}_{\cdot,1:j}^*)^T \mathbf{Y}_{\cdot,1:j}^* / nK + \rho \mathbf{I} / 2\}^{-1}$,

hence $A < 0$. Further $(B^2 - 8A)^{1/2} > |B|$. Thus, to satisfy the positive requirement of $L_{j,j}$, we get $\widehat{L}_{j,j}^{(r+1)} = -\{B + (B^2 - 8A)^{1/2}\}/2A$. This provides a closed-form solution for $\widehat{L}_{j,j}^{(r+1)}$, and subsequently a closed-form solution for $\widehat{\mathbf{L}}_{j,1:j-1}^{(r+1)}$. We next update Ψ_j based on $\widehat{\mathbf{L}}_j^{(r+1)}$ and $\widehat{\mathbf{U}}_j^{(r)}$ by minimizing the objective function

$$\frac{\rho}{2} \|\Psi_j - \widehat{\mathbf{L}}_j^{(r+1)} - \frac{1}{\rho} \widehat{\mathbf{U}}_j^{(r)}\|_2^2 + \lambda_2 p(\Psi_j), \quad (2.7)$$

which has the group lasso penalty. To minimize (2.7), we consider its dual problem.

Theorem 1. *A dual of (2.7) is given by*

$$\min_{\mathbf{A}} \|\mathbf{Z}_j - \frac{\lambda_2}{\rho} \sum_{l=1}^{j-1} \mathbf{A}_{\cdot,l}\|_2^2, \quad s.t. \quad \|\mathbf{A}_{1:l,l}\|_2 \leq 1, \mathbf{A}_{l+1:j,l} = \mathbf{0}, \text{ for } l = 1, \dots, j-1, \quad (2.8)$$

where $\mathbf{Z}_j = \widehat{\mathbf{L}}_j^{(r+1)} + \widehat{\mathbf{U}}_j^{(r)}/\rho$, $\mathbf{A}$ is a $j \times (j-1)$ matrix. Given $\widehat{\mathbf{A}}$, the optimizer of (2.7) is

$$\widehat{\Psi}_j^{(r+1)} = \mathbf{Z}_j - \frac{\lambda_2}{\rho} \sum_{l=1}^{j-1} \widehat{\mathbf{A}}_{\cdot,l}.$$

The proof of Theorem 1 is given in the Supplement Materials. Following Yu and Bien (2017), we use the blockwise coordinate descent (BCD) method to solve (2.8), where we sequentially perform elliptical projection to update each column of $\mathbf{A} \in \mathcal{R}^{j \times (j-1)}$. This strategy is developed in Bien et al. (2016) and Jenatton et al. (2011). It takes advantage of the upper triangle structure of $\mathbf{A}$, and only requires a single pass of BCD. Specifically, we first set $\widehat{\mathbf{A}} = \mathbf{0}$, then for $l = 1, \dots, j-1$, we sequentially update the l th column of

$\hat{\mathbf{A}}$. Following (2.8), the l th column of $\mathbf{A}$ is obtained by solving

$$\min_{\mathbf{A}_{\cdot,l}} \|\mathbf{\Gamma}_l - \frac{\lambda_2}{\rho} \mathbf{A}_{\cdot,l}\|_2^2, \quad s.t. \quad \|\mathbf{A}_{1:l,l}\|_2 \leq 1 \text{ and } \mathbf{A}_{l+1:j,l} = \mathbf{0},$$

where $\mathbf{\Gamma}_l \equiv \mathbf{Z}_j - \rho^{-1} \lambda_2 \sum_{q=1}^{l-1} \hat{\mathbf{A}}_{\cdot,q}$ is a j -dimensional vector. Obviously, if $\|(\mathbf{\Gamma}_l)_{1:l}\|_2 \leq \lambda_2/\rho$, then $\hat{\mathbf{A}}_{1:l,l} = \rho(\mathbf{\Gamma}_l)_{1:l}/\lambda_2$. Otherwise, $\hat{\mathbf{A}}_{1:l,l} = (\mathbf{\Gamma}_l)_{1:l}/\|(\mathbf{\Gamma}_l)_{1:l}\|_2$. Combining the two situations, we can write that $\hat{\mathbf{A}}_{1:l,l} = (\mathbf{\Gamma}_l)_{1:l}/\max\{\lambda_2/\rho, \|(\mathbf{\Gamma}_l)_{1:l}\|_2\}$. $\hat{\Psi}_j^{(r+1)} = \mathbf{Z}_j - \rho^{-1} \lambda_2 \sum_{l=1}^{j-1} \hat{\mathbf{A}}_{\cdot,l}$. Finally, based on $\hat{\mathbf{L}}_j^{(r+1)}$ and $\hat{\Psi}_j^{(r+1)}$, we follow the ADMM procedure to update Lagrange multiplier $\hat{\mathbf{U}}_j^{(r+1)}$ via $\hat{\mathbf{U}}_j^{(r+1)} = \hat{\mathbf{U}}_j^{(r)} + \rho(\hat{\mathbf{L}}_j^{(r+1)} - \hat{\Psi}_j^{(r+1)})$. The detailed process of solving the objective function (2.6) are provided in Algorithm 1. Algorithm 1 is applied to all $j = 1, \dots, J$ to yield $\hat{\mathbf{L}}^{(s+1)}$ defined in (2.4). We iteratively update the estimation for $\mathbf{W}^{-1}$ and $\mathbf{L}$ as described in Algorithm 2 to obtain $\widehat{\mathbf{W}}^{-1}$ and $\hat{\mathbf{L}}$, and form $\hat{\Sigma}^{-1} = \hat{\mathbf{L}}^T \hat{\mathbf{L}} \otimes \widehat{\mathbf{W}}^{-1}$ as our final estimator for the precision matrix. Although there is no guarantee that the algorithm converges to the global minimum, the algorithm converges to a local stationary point of (2.2).

Remark 2. The optimization process (Algorithm 2) requires an initial precision matrix for the time dimension, $\mathbf{R}^{-1}$. As the sample covariance matrix for the time dimension depends on $\mathbf{W}^{-1}$, we initialize $\mathbf{R}^{-1}$ as the identity matrix. For the variable dimension, $\mathbf{W}^{-1}$ is estimated using the glasso function, which defaults to the inverse of the sample covariance matrix. In Algorithm 1, there requires initial values of $\mathbf{L}$, Φ and $\mathbf{U}$. Given that $\mathbf{R}^{-1}$ is initialized as the identity matrix, we set the initial values as $\mathbf{L}^{(0)} = \mathbf{I}$, $\Phi^{(0)} = \mathbf{I}$,

Algorithm 1 ADMM algorithm to solve (2.6)

Input: Initial values $\widehat{\mathbf{L}}_j^{(0)}, \widehat{\Psi}_j^{(0)}, \widehat{\mathbf{U}}_j^{(0)}, \lambda_2, \rho > 0, r = 0$.

Main procedure:

Step 1. Update $\widehat{L}_{j,j}^{(r+1)} = \frac{-B - \sqrt{B^2 - 8A}}{2A}$ and

$$\widehat{\mathbf{L}}_{j,1:j-1}^{(r+1)} = \left\{ \rho \widehat{\Psi}_{j,1:j-1}^{(r)} - \frac{2}{nK} (\mathbf{Y}_{\cdot,j}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* \widehat{L}_{j,j}^{(r+1)} - \widehat{\mathbf{U}}_{j,1:j-1}^{(r)} \right\} \left\{ \frac{2}{nK} (\mathbf{Y}_{\cdot,1:j-1}^*)^T \mathbf{Y}_{\cdot,1:j-1}^* + \rho \mathbf{I} \right\}^{-1}$$

Step 2. Let $\mathbf{Z}_j = \widehat{\mathbf{L}}_j^{(r+1)} + \widehat{\mathbf{U}}_j^{(r)}/\rho$. For $l = 1, \dots, j-1$, let $\mathbf{\Gamma}_l = \mathbf{Z}_j - (\lambda_2/\rho) \sum_{q=1}^{l-1} \widehat{\mathbf{A}}_{\cdot,q}$, and $\widehat{\mathbf{A}}_{1:l,l} = (\mathbf{\Gamma}_l)_{1:l}/\max\{\frac{\lambda_2}{\rho}, \|(\mathbf{\Gamma}_l)_{1:l}\|_2\}$, $\widehat{\mathbf{A}}_{l+1:j,l} = \mathbf{0}$. Set $\widehat{\Psi}_j^{(r+1)} = \mathbf{Z}_j - \frac{\lambda_2}{\rho} \sum_{l=1}^{j-1} \widehat{\mathbf{A}}_{\cdot,l}$.

Step 3. Set $\widehat{\mathbf{U}}_j^{(r+1)} = \widehat{\mathbf{U}}_j^{(r)} + \rho(\widehat{\mathbf{L}}_j^{(r+1)} - \widehat{\Psi}_j^{(r+1)})$;

Step 4. Increase r by 1 and go back to Step 1 until $\|\widehat{\mathbf{L}}_j^{(r+1)} - \widehat{\Psi}_j^{(r+1)}\|_2 < \varepsilon_{\text{prime}}$, and $\|\rho(\widehat{\Psi}_j^{(r+1)} - \widehat{\Psi}_j^{(r)})\|_2 < \varepsilon_{\text{dual}}$, where $\varepsilon_{\text{prime}} = \sqrt{j}\varepsilon_{\text{abs}} + \varepsilon_{\text{rel}} \max\{\|\widehat{\mathbf{L}}_j^{(r+1)}\|_2, \|\widehat{\Psi}_j^{(r+1)}\|_2\}$, $\varepsilon_{\text{dual}} = \sqrt{j}\varepsilon_{\text{abs}} + \varepsilon_{\text{rel}}\|\widehat{\mathbf{U}}_j^{(r+1)}\|_2$, and $\varepsilon_{\text{abs}}, \varepsilon_{\text{rel}}$ are predetermined constants.

Output: $\widehat{\mathbf{L}}_j^{(r+1)}, \widehat{\Psi}_j^{(r+1)}$.

and $\mathbf{U}^{(0)} = \mathbf{0}$.

We divide the entire optimization process into two stages. First, given $\widehat{\mathbf{L}}^{(s)}$, we apply the graphical lasso algorithm to obtain $(\widehat{\mathbf{W}}^{-1})^{(s+1)}$. Since the objective function is convex and the initial $\widehat{\mathbf{L}}^{(s)}$ ensures that $(\widehat{\mathbf{R}}^{-1})^{(s)}$ remains positive definite, thus, the optimization process always converges. In the second stage, based on $(\widehat{\mathbf{W}}^{-1})^{(s+1)}$, we impose the condition $\Psi = \mathbf{L}$ and employ the ADMM algorithm to compute $\widehat{\mathbf{L}}^{(s+1)}$. As the objective function is convex with respect to $\mathbf{L}$ and Ψ , and given an appropriate tuning parameter λ_2 , this computation process also converges. Since both subroutines are convergent, we leverage the convergence properties of bi-convex functions, as discussed in Gorski et al. (2007). By iteratively alternating between these two stages, the optimization process converges to a locally optimal solution, $\{(\widehat{\mathbf{W}}^{(-1)})^{(s+1)}, \widehat{\mathbf{L}}^{(s+1)}\}$.

Algorithm 2 The complete algorithm to solve (2.2)

Input: Initial values $\widehat{\mathbf{L}}^{(0)}$, $(\widehat{\mathbf{W}}^{-1})^{(0)}$, $\lambda_1, \lambda_2, \rho > 0$, $s = 0$.

Main procedure:

Step 1. At the given $\widehat{\mathbf{L}}^{(s)}$, obtain $(\widehat{\mathbf{W}}^{-1})^{(s+1)}$ by applying the glasso method to solve (2.3).

Step 2. At the given $(\widehat{\mathbf{W}}^{-1})^{(s+1)}$, obtain $\widehat{\mathbf{L}}^{(s+1)}$ by solving (2.4). (2.4) is solved row-wise, where for $j = 1, \dots, J$, the nonzero part of the j th row, i.e., $\widehat{\mathbf{L}}_j^{(s+1)}$, is obtained through solving (2.5) via Algorithm 1.

Step 3. Increase s by 1 and go back to Step 1 until $\|(\widehat{\mathbf{W}}^{-1})^{(s+1)} - (\widehat{\mathbf{W}}^{-1})^{(s)}\|_F + \|\widehat{\mathbf{L}}^{(s+1)} - \widehat{\mathbf{L}}^{(s)}\|_F < \varepsilon$, where ε is a pre-determined constant.

Output: $\widehat{\mathbf{W}}^{-1} \equiv (\widehat{\mathbf{W}}^{-1})^{(s+1)}$ and $\widehat{\mathbf{L}} \equiv \widehat{\mathbf{L}}^{(s+1)}$.

3. Statistical Properties

We now study the statistical properties of BKS, through establishing the converge rates of $\widehat{\mathbf{L}}$ and $\widehat{\mathbf{W}}^{-1}$ respectively. All technical proofs are provided in the Supplement Materials.

Let d_j be the true bandwidth of the j th row in $\mathbf{L}$, we will show the consistence of the estimators $\widehat{d}_j$, $j = 2, \dots, J$. We now explain some notations that will be used throughout the paper. For a $n \times p$ real matrix $\mathbf{M} = (M_{ij})$, the l_1 norm is defined as $\|\mathbf{M}\|_1 = \sum_{i,j} |M_{ij}|$, and the Frobenius norm is $\|\mathbf{M}\|_F = (\sum_{i,j} M_{i,j}^2)^{1/2}$. We make the following assumptions.

(C1) The true lower triangular matrix $\mathbf{L} \in \mathcal{R}^{J \times J}$ has bandwidth d_j on row j for $j = 2, \dots, J$, and has positive diagonal elements. Therefore, $L_{j,q} = 0$ for $1 \leq q < j - d_j$

and $q > j$. $\mathbf{W}^{-1} \in \mathcal{R}^{K \times K}$ is a sparse positive definition matrix. Let $\max_{j=2, \dots, J} d_j = O(1)$. Let w and $v \equiv \sum_{j=2}^J d_j$ represent the total numbers of non-zero off-diagonal elements in $\mathbf{W}^{-1}$ and $\mathbf{L}$ respectively, and satisfy $w = O(K)$, $v = O(J)$.

(C2) There exist positive constants τ_1 and τ_2 such that

$$\begin{aligned} 0 < \tau_1^{-1} &\leq \sigma_{\min}(\mathbf{L}) \leq \sigma_{\max}(\mathbf{L}) \leq \tau_1 < \infty, \\ 0 < \tau_2^{-1} &\leq \sigma_{\min}(\mathbf{W}^{-1}) \leq \sigma_{\max}(\mathbf{W}^{-1}) \leq \tau_2 < \infty, \end{aligned}$$

where $\sigma_{\min}(\cdot)$ and $\sigma_{\max}(\cdot)$ denote the minimum and maximum singular values of a matrix.

(C3) Let the square of the j th component of $\text{vec}(\mathbf{Y}_i)$ have distribution function G_j , then

$$\max_{1 \leq j \leq KJ} \int_0^\infty \exp(\psi t) dG_j(t) < \infty, \text{ for all } \psi \in (0, \psi_0),$$

where $\psi_0 > 0$ is a constant.

For (C1), it means that the true precision matrix $\boldsymbol{\Sigma}^{-1} = \mathbf{R}^{-1} \otimes \mathbf{W}^{-1}$ has a sparse block banded structure. (C2) suggests that the singular value of $\mathbf{L}$ is bounded, which is equivalent to the bounded eigenvalue condition generally. In reference to (C3), which is defined similarly to the condition in Bickel and Levina (2008a), to accommodate the departure from normality, we establish that the maximum difference between $\mathbf{S}_{j,l}$ and $\boldsymbol{\Sigma}_{j,l}$, denoted as $\max_{1 \leq j, l \leq KJ} |\mathbf{S}_{j,l} - \boldsymbol{\Sigma}_{j,l}|$, always satisfies the inequality $O_p\{\log(KJ)/n\}$.

3.1 Precision matrix estimation consistency

Here, $\mathbf{S}$ represents the sample covariance matrix for the random samples $\{\mathbf{y}_i\}_{i=1}^n$.

3.1 Precision matrix estimation consistency

We now consider the convergence rate of the precision matrix estimator $\widehat{\Sigma}^{-1}$ when $m \equiv KJ$ and n both diverge to infinity. Denote $a \asymp b$ as $c_1 \leq |a/b| \leq c_2$, where c_1 and c_2 are positive constants.

Theorem 2. *Assuming that Assumptions (C1), (C2) and (C3) hold, and the tuning parameters satisfy $\lambda_1 \asymp (\log m/n)^{1/2}$ and $\lambda_2 \asymp (\log m/n)^{1/2}$. If $(K + J) \log m = o(n)$, then there exists a local minimizer of (2.2). Moreover, the estimators $\widehat{\mathbf{W}}^{-1}$, $\widehat{\mathbf{L}}_j$ and $\widehat{\Sigma}^{-1}$ converge in the sense that*

$$\|\widehat{\mathbf{W}}^{-1} - \mathbf{W}^{-1}\|_F = O_p\{(K \log m/n)^{1/2}\}, \quad \|\widehat{\mathbf{L}}_j - \mathbf{L}_j\|_2 = O_p\{(\log m/n)^{1/2}\},$$

and

$$\|\widehat{\Sigma}^{-1} - \Sigma^{-1}\|_F = O_p\{(m \log m/n)^{1/2}\}.$$

Furthermore, if $\text{vec}(\mathbf{Y}_i) \sim N(\mathbf{0}, \mathbf{R} \otimes \mathbf{W})$, then the estimators $\widehat{\mathbf{W}}^{-1}$, $\widehat{\mathbf{L}}_j$ and $\widehat{\Sigma}^{-1}$ will converge at a faster rate. Specifically, assume Assumptions (C1), (C2) and (C3) to hold and the tuning parameters to satisfy $\lambda_1 \asymp \{\log K/(nJ)\}^{1/2}$ and $\lambda_2 \asymp \{\log J/(nK)\}^{1/2}$. If $J \log J \asymp K \log K$, and $\log K = o(n)$, $\log J = o(n)$, then the estimators $\widehat{\mathbf{W}}^{-1}$, $\widehat{\mathbf{L}}_j$ and

3.2 Uniqueness of the banded estimator

$\widehat{\Sigma}^{-1}$ converge in the sense that

$$\|\widehat{\mathbf{W}}^{-1} - \mathbf{W}^{-1}\|_F = O_p[\{K \log K/(nJ)\}^{1/2}], \quad \|\widehat{\mathbf{L}}_j - \mathbf{L}_j\|_2 = O_p[\{\log J/(nK)\}^{1/2}],$$

and

$$\|\widehat{\Sigma}^{-1} - \Sigma^{-1}\|_F = O_p[\max\{(J \log J/n)^{1/2}, (K \log K/n)^{1/2}\}].$$

Theorem 2 establishes the convergence rate of the precision matrix estimation. This rate agrees with that of SMGM (Leng and Tang, 2012), and is better than that of ABR (Qian et al., 2020), which will be reflected in the simulation studies. As noted by Leng and Tang (2012), when multiple local minimizers exist, identifying the optimal solution in practice becomes challenging, and there does not seem to be an algorithm that can consistently find the optimal solution.

3.2 Uniqueness of the banded estimator

In handling the optimization problem in (2.4), we have decomposed it into J separate optimization problems. It is worth noting that without the banded requirement, the individual estimator $\widehat{\mathbf{L}}_j$ may not be unique. For example, when the dimension of $\mathbf{L}_j$ satisfies $j > nK$, the objective function in (2.5) may not be strictly convex as a function of $\mathbf{L}_j$ which leads to multiple minimizers. However, when λ_2 is sufficiently large, we will show that the additional banded requirement will lead to sufficient sparseness so that the optimizer $\widehat{\mathbf{L}}_j$ will be unique. In order to show this, we first establish two lemmas.

3.2 Uniqueness of the banded estimator

Lemma 1. For any given $\lambda_2 > 0$ and any given $\widehat{\mathbf{W}}^{-1}$, $\widehat{\mathbf{L}}_j$ is a solution to the objective function $\min_{L_{j,j} > 0, \mathbf{L}_j \in \mathcal{R}^j} f(\mathbf{L}_j)$, where

$$f(\mathbf{L}_j) \equiv -2 \log L_{j,j} + \frac{1}{nK} \|\mathbf{Y}_{:,1:j}^* \mathbf{L}_j\|_2^2 + \lambda_2 p(\mathbf{L}_j), \quad (3.1)$$

iff there exists $\widehat{\mathbf{A}} \in \mathcal{R}^{j \times (j-1)}$ such that

$$-\frac{2}{\widehat{L}_{j,j}} \mathbf{e}_j + \frac{2}{nK} \mathbf{Y}_{:,1:j}^{*\top} \mathbf{Y}_{:,1:j}^* \widehat{\mathbf{L}}_j + \lambda_2 \sum_{l=1}^{j-1} \widehat{\mathbf{A}}_{:,l} = \mathbf{0}, \quad (3.2)$$

where for $l = 1, \dots, j-1$,

$$\widehat{\mathbf{A}}_{l+1:j,l} = \mathbf{0}, \widehat{\mathbf{A}}_{1:l,l} = (\widehat{\mathbf{L}}_j)_{1:l} / \|(\widehat{\mathbf{L}}_j)_{1:l}\|_2, \text{ if } (\widehat{\mathbf{L}}_j)_{1:l} \neq \mathbf{0}, \text{ and } \|\widehat{\mathbf{A}}_{1:l,l}\|_2 \leq 1. \quad (3.3)$$

Further, if the tuning parameter $\lambda_2 = C(\log m/n)^{1/2}$, where C is a sufficiently large constant, then under the conditions in Theorem 2, the estimator $\widehat{\mathbf{L}}_j$ is sparse with bandwidth $\widehat{d}_j$, and $\|\widehat{\mathbf{A}}_{1:l,l}\|_2 < 1$ for $l = 1, \dots, j-1 - \widehat{d}_j$.

Lemma 2. Let $\widehat{\mathbf{L}}_j$ and $\widehat{\mathbf{A}}$ be as defined in Lemma 1. Assume that $\|\widehat{\mathbf{A}}_{1:l,l}\|_2 < 1$ for $l = 1, \dots, j - \widehat{d}_j - 1$. Then, any other solution $\widetilde{\mathbf{L}}_j$ has a bandwidth at most that of $\widehat{\mathbf{L}}_j$, i.e., $\widetilde{d}_j \leq \widehat{d}_j$.

Theorem 3. For any given $\lambda_2 > 0$ and $\widehat{\mathbf{W}}^{-1}$, let $\widehat{\mathbf{L}}_j$ be a solution to the objective function $\min_{L_{j,j} > 0, \mathbf{L}_j \in \mathcal{R}^j} f(\mathbf{L}_j)$ with bandwidth $\widehat{d}_j$, where $f(\mathbf{L}_j)$ is defined in Lemma 1. Let $\widehat{\mathbf{A}}$ be as defined in Lemma 1, and define the non-zero index set $\widehat{D} \equiv \{l : \widehat{L}_{j,l} \neq 0\}$. Let

$\lambda_2 = C(\log m/n)^{1/2}$, where C is a sufficiently large constant, and assume $\mathbf{Y}_{\cdot, \hat{D}}^*$ has full column rank, i.e., $\text{rank}(\mathbf{Y}_{\cdot, \hat{D}}^*) = \hat{d}_j + 1$. Then, $\hat{\mathbf{L}}_j$ is unique.

3.3 True bandwidth recovery

In this section, we show that our estimator $\hat{\mathbf{L}}$ can correctly recover the true bandwidth of each row uniformly with probability approaching 1 under mild conditions. To show this, following the primal-dual witness procedure in Yu and Bien (2017), we first construct the primal-dual witness solution pairs $(\tilde{\mathbf{A}}, \tilde{\mathbf{L}})$ for the optimal problem assuming the true bandwidth d_j of each row is known. We then prove that this solution is identical to the solution to (2.4), which in turn implies that the estimated bandwidths are identical to the true bandwidths.

Theorem 4. *Assume the conditions required in Lemmas 1, 2 and Theorems 2, 3 are satisfied, if the condition $\min_{j \in \{2, \dots, J\}} \min_{l \geq j-d_j} |L_{j,l}| > \lambda_2$ is satisfied, then*

$$pr(\sup_{j=2, \dots, J} |\hat{d}_j - d_j| = 0) \rightarrow 1.$$

Remark 3. Theorem 4 holds under the assumption that the nonzero entries in $\mathbf{L}$ are uniformly bounded below by $(\log m/n)^{1/2}$. This implies that the minimal signal strength of $\mathbf{L}$ that is detectable is determined by the relation between the matrix size K, J and the sample size n . A larger matrix requires stronger minimal signal.

4. Simulation Studies

We now conduct simulation studies to study the finite sample performance of BKS. For comparison, in addition to BKS, we also implement some competitive methods, including SMGM (Leng and Tang, 2012) and ABR (Qian et al., 2021), both are designed to handle matrix-valued data. In addition, we also implement VB and Unweighted VB (UVB) proposed in Yu and Bien (2017), which is suitable when the vector $\text{vec}(\mathbf{Y}_i)$ is ordered, i.e. nearby components of $\text{vec}(\mathbf{Y}_i)$ have larger correlations. To facilitate the comparison to VB, we introduce a weighted version of BKS, where the penalty in (2.1) is modified to $p_w(\mathbf{L}_{j\cdot}) \equiv \sum_{l=1}^{j-1} (\sum_{q=1}^l w_{jq}^2 L_{j,q}^2)^{1/2}$, where $w_{jq} = 1/(j-q+1)^2$. Please note that the weights w_{jq} are the same as that of VB. We name the corresponding method Weighted BKS (WBKS). Note that these figures in simulation studies both can be found in Supplement Material S6.

4.1 Multivariate normal distribution

We generate the precision matrix by setting $\Sigma^{-1} = (\mathbf{L}^T \mathbf{L}) \otimes \mathbf{W}^{-1}$, where $\mathbf{L}$ is a lower triangular matrix with ones on the diagonal and row specific bandwidth $d_j, j = 1, \dots, J$, and $\mathbf{W}^{-1}$ is a sparse positive definite matrix. To generate $\mathbf{L}$, we consider two cases.

- Case 1. $L_{j,l} = I(j-l=0) + 0.8I(j-l=1) + 0.6I(j-l=2) + 0.4I(j-l=3) + 0.2I(j-l=4)$, where $j = 1, \dots, J$, and $1 \leq l \leq j$.
- Case 2. $L_{j,l} = 0.7^{|j-l|}$, where $j = 1, \dots, J$, and $j-d_j \leq l \leq j$. Here, d_j is randomly generated from a discrete uniform distribution on $[1, j/2]$.

4.1 Multivariate normal distribution

We can see that $\mathbf{L}$ in Case 1 is a banded matrix with equal values 1, 0.8, 0.6, 0.4, 0.2 on the lower bands starting from the diagonal, while in Case 2, $\mathbf{L}$ has row-specific bandwidth, and in each row, the values of the nonzero elements are decreasing while moving away from the diagonal position. To generate $\mathbf{W}^{-1}$, we set $\mathbf{W}^{-1} = 0.5(\mathbf{B} + \mathbf{B}^T) + c\mathbf{I}_K$, where $\mathbf{B}$ is a strictly upper triangular matrix with its elements independently generated from a Bernoulli distribution with parameter 0.1, and c is chosen such that the condition number of $\mathbf{W}^{-1}$ is K . We illustrate the structure of the precision matrix in the two cases when $K = 20, J = 10$ in Figure 1.

We then proceed to generate the n independent longitudinal data $\{\text{vec}(\mathbf{Y}_i)\}_{i=1}^n$ from the $m = KJ$ dimensional multivariate normal distribution $N(\mathbf{0}, \Sigma)$. We consider sample sizes $n = 10, 50$ and 100 , and repeat 100 times under each sample size. We consider four combinations of (K, J) .

We report the estimation accuracy of the estimator $\hat{\Sigma}^{-1}$ in terms of two criterions, Frobenius norm (FN) and Kullback-Leibler (KL) loss, defined as

$$\Delta_{\text{FN}}(\hat{\Sigma}^{-1}, \Sigma^{-1}) \equiv \frac{1}{m} \|\hat{\Sigma}^{-1} - \Sigma^{-1}\|_F^2, \quad \Delta_{\text{KL}}(\hat{\Sigma}^{-1}, \Sigma^{-1}) \equiv \frac{1}{m} \{\text{tr}(\Sigma^{-1}\hat{\Sigma}) - \ln(\Sigma^{-1}\hat{\Sigma}) - m\}$$

respectively. We also report the performance of bandwidth recovery using the true negative rate (TNR) and the true positive rate (TPR) (Leng and Tang, 2012), that is,

$$\text{TNR} = \frac{\#\{\hat{\Sigma}_{ij} = 0 \ \& \ \Sigma_{ij} = 0\}}{\#\{\Sigma_{ij} = 0\}}, \quad \text{TPR} = \frac{\#\{\hat{\Sigma}_{ij} \neq 0 \ \& \ \Sigma_{ij} \neq 0\}}{\#\{\Sigma_{ij} \neq 0\}}.$$

4.1 Multivariate normal distribution

Table 1: Comparison in terms of FN and KL for different estimators of the precision matrix Σ^{-1} , in the form $\text{average}_{\text{standarderror}}$, over 100 replications in Case 1.

(K, J)		(10,10)	(10,20)	(20,10)	(20,20)	(10,10)	(10,20)	(20,10)	(20,20)
Methods		FN				KL			
10	WBKS	1.276 _{0.315}	3.256 _{0.352}	2.620 _{0.527}	4.782 _{0.499}	0.106 _{0.021}	0.120 _{0.014}	0.075 _{0.012}	0.078 _{0.008}
	BKS	1.499 _{0.349}	2.865 _{0.330}	2.412 _{0.458}	4.157 _{0.467}	0.096 _{0.017}	0.093 _{0.011}	0.073 _{0.012}	0.065 _{0.007}
	SMGM	7.714 _{3.912}	7.287 _{2.823}	12.39 _{4.107}	12.45 _{2.833}	0.764 _{0.538}	0.403 _{0.449}	0.365 _{0.325}	0.215 _{0.208}
50	WBKS	0.454 _{0.150}	0.539 _{0.114}	0.541 _{0.138}	0.881 _{0.135}	0.018 _{0.003}	0.022 _{0.003}	0.015 _{0.002}	0.014 _{0.001}
	BKS	0.328 _{0.113}	0.477 _{0.103}	0.634 _{0.158}	0.644 _{0.112}	0.017 _{0.003}	0.019 _{0.003}	0.014 _{0.002}	0.013 _{0.001}
	SMGM	0.404 _{0.094}	0.506 _{0.095}	1.252 _{0.176}	1.862 _{0.181}	0.023 _{0.003}	0.026 _{0.004}	0.027 _{0.002}	0.022 _{0.002}
	ABR	10.59 _{0.093}	12.75 _{0.070}	20.12 _{0.092}	—	0.386 _{0.009}	0.429 _{0.007}	0.481 _{0.007}	—
	VB	10.73 _{0.077}	12.96 _{0.055}	18.59 _{0.102}	22.64 _{0.061}	0.396 _{0.015}	0.468 _{0.012}	0.442 _{0.010}	0.520 _{0.007}
	UVB	11.03 _{0.067}	12.61 _{0.050}	20.30 _{0.071}	23.54 _{0.051}	0.475 _{0.012}	0.498 _{0.010}	0.568 _{0.008}	0.598 _{0.005}
100	WBKS	0.172 _{0.059}	0.327 _{0.076}	0.261 _{0.071}	0.308 _{0.068}	0.009 _{0.001}	0.011 _{0.001}	0.007 _{0.001}	0.007 _{0.001}
	BKS	0.152 _{0.052}	0.324 _{0.075}	0.215 _{0.049}	0.273 _{0.062}	0.009 _{0.001}	0.009 _{0.001}	0.007 _{0.001}	0.006 _{0.001}
	SMGM	0.299 _{0.078}	0.476 _{0.085}	1.078 _{0.145}	1.880 _{0.143}	0.011 _{0.002}	0.013 _{0.001}	0.014 _{0.001}	0.014 _{0.001}
	ABR	10.59 _{0.093}	11.78 _{0.046}	17.98 _{0.096}	—	0.386 _{0.009}	0.327 _{0.003}	0.364 _{0.004}	—
	VB	10.73 _{0.077}	12.96 _{0.055}	17.13 _{0.079}	20.72 _{0.052}	0.396 _{0.014}	0.468 _{0.012}	0.314 _{0.005}	0.366 _{0.004}
	UVB	11.03 _{0.067}	12.60 _{0.050}	18.21 _{0.057}	21.63 _{0.040}	0.475 _{0.012}	0.498 _{0.010}	0.441 _{0.006}	0.476 _{0.004}

During the estimation process, it is common for the estimated values of elements in Σ that should be 0 to be very small in absolute value but not exactly zero. In order to examine the accuracy of structure recovery, we assign a value of zero to all estimates below 0.01 in absolute value across all methods. In Table 1, we present the matrix estimation performance in Case 1. The results show that the average of FN and KL both decrease when the sample size increases for all estimators.

4.1 Multivariate normal distribution

Table 2: Comparison in terms of TNR and TPR for different estimators of the precision matrix Σ^{-1} , in the form $\text{average}_{\text{standarderror}}$, over 100 replications in Case 1.

(K, J)		(10,10)	(10,20)	(20,10)	(20,20)	(10,10)	(10,20)	(20,10)	(20,20)
Methods		TNR				TPR			
10	WBKS	70.89 _{0.054}	85.17 _{0.017}	86.56 _{0.023}	92.45 _{0.008}	98.34 _{0.018}	93.50 _{0.022}	97.13 _{0.020}	94.24 _{0.018}
	BKS	77.21 _{0.044}	89.20 _{0.018}	84.49 _{0.023}	93.78 _{0.008}	97.44 _{0.021}	94.73 _{0.019}	97.27 _{0.019}	94.79 _{0.018}
	SMGM	88.51 _{0.133}	92.69 _{0.038}	89.85 _{0.054}	93.96 _{0.011}	41.25 _{0.395}	62.98 _{0.290}	54.09 _{0.255}	67.83 _{0.099}
50	WBKS	85.88 _{0.030}	84.16 _{0.020}	87.31 _{0.016}	92.95 _{0.007}	99.61 _{0.009}	99.19 _{0.010}	99.94 _{0.004}	99.82 _{0.003}
	BKS	82.39 _{0.035}	86.32 _{0.021}	87.96 _{0.015}	93.40 _{0.007}	99.77 _{0.007}	99.37 _{0.008}	99.94 _{0.004}	99.91 _{0.002}
	SMGM	56.72 _{0.043}	77.27 _{0.021}	59.00 _{0.020}	82.08 _{0.010}	98.48 _{0.018}	96.70 _{0.020}	95.72 _{0.031}	91.36 _{0.025}
	ABR	81.96 _{0.009}	92.86 _{0.003}	89.75 _{0.007}	—	39.56 _{0.009}	33.33 _{0.005}	30.76 _{0.009}	—
	VB	93.51 _{0.005}	97.89 _{0.002}	94.74 _{0.003}	98.43 _{0.001}	33.25 _{0.010}	25.88 _{0.006}	30.62 _{0.008}	22.57 _{0.005}
	UVB	86.19 _{0.005}	92.63 _{0.001}	95.87 _{0.002}	98.12 _{0.001}	31.21 _{0.007}	29.51 _{0.004}	14.79 _{0.003}	12.81 _{0.002}
100	WBKS	76.73 _{0.034}	88.70 _{0.012}	86.30 _{0.015}	91.10 _{0.013}	99.94 _{0.003}	99.85 _{0.004}	100.00 _{0.00}	100.00 _{0.00}
	BKS	76.56 _{0.034}	90.05 _{0.014}	85.58 _{0.015}	93.05 _{0.012}	99.97 _{0.002}	99.89 _{0.003}	100.00 _{0.00}	100.00 _{0.00}
	SMGM	65.31 _{0.038}	82.97 _{0.019}	66.97 _{0.020}	87.41 _{0.013}	99.12 _{0.015}	98.06 _{0.015}	97.11 _{0.023}	92.01 _{0.019}
	ABR	74.90 _{0.009}	90.96 _{0.003}	82.30 _{0.004}	—	52.11 _{0.013}	39.84 _{0.008}	38.33 _{0.006}	—
	VB	84.39 _{0.006}	93.51 _{0.003}	92.01 _{0.003}	91.40 _{0.001}	48.44 _{0.012}	37.59 _{0.007}	38.66 _{0.008}	31.96 _{0.005}
	UVB	76.98 _{0.005}	87.19 _{0.002}	82.54 _{0.003}	92.25 _{0.001}	46.57 _{0.011}	47.53 _{0.007}	34.41 _{0.005}	29.47 _{0.003}

WBKS and BKS perform the best on average in every circumstance, because these methods fully take into account the sparsity of $\mathbf{W}^{-1}$, the bandedness of $\mathbf{L}$ and the Kronecker product structure. In contrast, SMGM does not take into account the banded matrix feature for $\mathbf{L}$, ABR does not utilize the Kronecker product nature of the precision matrix, and VB and UVB ignore sparsity and the Kronecker product structure. We also find that WBKS and BKS tend to have rather small variability. In fact, they have the smallest variability when sample size $n = 10$, reflecting superior estimation efficiency.

4.1 Multivariate normal distribution

Table 3: Comparison in terms of FN and KL for different estimators of the precision matrix Σ^{-1} , in the form average_{standarderror}, over 100 replications in Case 2.

(K, J)		(10,10)	(10,20)	(20,10)	(20,20)	(10,10)	(10,20)	(20,10)	(20,20)
Methods		FN				KL			
10	WBKS	1.651 _{0.220}	2.333 _{0.210}	1.708 _{0.220}	2.220 _{0.286}	0.104 _{0.018}	0.109 _{0.012}	0.078 _{0.012}	0.072 _{0.008}
	BKS	1.308 _{0.210}	1.687 _{0.199}	1.610 _{0.205}	1.934 _{0.254}	0.091 _{0.017}	0.084 _{0.010}	0.075 _{0.011}	0.060 _{0.007}
	SMGM	3.976 _{1.825}	5.042 _{2.019}	5.153 _{2.788}	6.521 _{0.778}	0.827 _{0.416}	0.883 _{0.435}	0.485 _{0.418}	0.138 _{0.086}
50	WBKS	0.195 _{0.064}	0.268 _{0.058}	0.313 _{0.076}	0.511 _{0.091}	0.015 _{0.003}	0.021 _{0.002}	0.015 _{0.002}	0.014 _{0.001}
	BKS	0.180 _{0.060}	0.242 _{0.054}	0.285 _{0.069}	0.390 _{0.077}	0.015 _{0.003}	0.018 _{0.002}	0.015 _{0.002}	0.013 _{0.001}
	SMGM	0.190 _{0.042}	0.243 _{0.044}	0.410 _{0.058}	0.935 _{0.109}	0.021 _{0.003}	0.024 _{0.002}	0.025 _{0.003}	0.020 _{0.002}
	ABR	4.797 _{0.062}	6.080 _{0.041}	8.882 _{0.051}	—	0.350 _{0.010}	0.351 _{0.006}	0.455 _{0.006}	—
	VB	4.635 _{0.058}	6.310 _{0.031}	7.982 _{0.052}	12.91 _{0.042}	0.340 _{0.016}	0.397 _{0.009}	0.421 _{0.012}	0.480 _{0.008}
	UVB	4.836 _{0.051}	6.351 _{0.034}	8.719 _{0.043}	13.61 _{0.041}	0.425 _{0.014}	0.447 _{0.007}	0.534 _{0.008}	0.562 _{0.006}
100	WBKS	0.124 _{0.037}	0.214 _{0.042}	0.167 _{0.038}	0.203 _{0.044}	0.008 _{0.001}	0.011 _{0.001}	0.007 _{0.001}	0.007 _{0.001}
	BKS	0.113 _{0.035}	0.177 _{0.038}	0.186 _{0.042}	0.187 _{0.042}	0.008 _{0.001}	0.009 _{0.001}	0.007 _{0.001}	0.006 _{0.000}
	SMGM	0.097 _{0.025}	0.190 _{0.035}	0.254 _{0.045}	0.881 _{0.087}	0.010 _{0.002}	0.013 _{0.001}	0.011 _{0.001}	0.011 _{0.001}
	ABR	3.719 _{0.056}	4.194 _{0.054}	7.831 _{0.037}	—	0.226 _{0.006}	0.257 _{0.005}	0.339 _{0.004}	—
	VB	3.821 _{0.047}	5.600 _{0.030}	6.804 _{0.046}	11.70 _{0.038}	0.216 _{0.007}	0.267 _{0.004}	0.272 _{0.005}	0.330 _{0.004}
	UVB	3.799 _{0.041}	5.188 _{0.028}	7.733 _{0.037}	12.36 _{0.025}	0.275 _{0.008}	0.299 _{0.005}	0.409 _{0.005}	0.428 _{0.004}

Between the weighted and unweighted versions of BKS, we would recommend BKS, due to its simplicity and its better performance under larger J . To demonstrate the bandwidth recovery performance of different estimators, we further report the averages and standard errors of TNR and TPR in Table 2. Once again, WBKS and BKS consistently demonstrate superior performance compared to their competitors, as evidenced by their large and balanced TNR and TPR values, and the superior ROC curves presented in Figure 4. Due to the similar performance from BKS and WBKS, we only choose to show

4.1 Multivariate normal distribution

Table 4: Comparison in terms of TNR and TPR for different estimators of the precision matrix Σ^{-1} , in the form $\text{average}_{\text{standarderror}}$, over 100 replications in Case 2.

(K, J)		(10,10)	(10,20)	(20,10)	(20,20)	(10,10)	(10,20)	(20,10)	(20,20)
Methods		TNR				TPR			
10	WBKS	82.17 _{0.052}	84.48 _{0.023}	80.40 _{0.035}	80.08 _{0.023}	98.02 _{0.028}	95.84 _{0.026}	99.37 _{0.013}	98.47 _{0.011}
	BKS	83.39 _{0.045}	87.71 _{0.021}	81.62 _{0.031}	82.70 _{0.019}	99.14 _{0.017}	97.69 _{0.020}	99.43 _{0.013}	98.73 _{0.010}
	SMGM	85.90 _{0.235}	94.16 _{0.100}	84.00 _{0.131}	90.74 _{0.014}	32.53 _{0.385}	26.63 _{0.352}	54.27 _{0.383}	65.10 _{0.075}
50	WBKS	75.38 _{0.041}	77.57 _{0.029}	74.08 _{0.029}	81.45 _{0.015}	100.00 _{0.00}	99.95 _{0.006}	100.00 _{0.00}	99.03 _{0.010}
	BKS	75.76 _{0.039}	76.82 _{0.028}	76.00 _{0.026}	81.71 _{0.015}	99.96 _{0.004}	99.48 _{0.007}	100.00 _{0.00}	99.32 _{0.007}
	SMGM	36.65 _{0.056}	57.35 _{0.045}	42.13 _{0.041}	65.51 _{0.021}	99.89 _{0.006}	97.27 _{0.015}	99.90 _{0.006}	94.02 _{0.019}
	ABR	66.55 _{0.023}	80.64 _{0.009}	62.07 _{0.013}	—	66.25 _{0.040}	45.89 _{0.020}	72.32 _{0.021}	—
	VB	75.86 _{0.015}	93.44 _{0.004}	85.03 _{0.007}	93.22 _{0.003}	58.98 _{0.025}	30.36 _{0.013}	48.30 _{0.015}	28.20 _{0.009}
	UVB	72.62 _{0.010}	90.62 _{0.003}	89.29 _{0.004}	94.93 _{0.001}	58.63 _{0.019}	27.09 _{0.007}	28.74 _{0.010}	15.55 _{0.003}
100	WBKS	72.52 _{0.038}	75.44 _{0.030}	79.28 _{0.020}	81.86 _{0.013}	100.00 _{0.00}	99.28 _{0.009}	100.00 _{0.00}	99.67 _{0.006}
	BKS	72.46 _{0.039}	76.45 _{0.028}	80.03 _{0.019}	82.27 _{0.012}	100.00 _{0.00}	99.42 _{0.007}	100.00 _{0.00}	99.65 _{0.006}
	SMGM	53.73 _{0.040}	70.07 _{0.033}	52.12 _{0.025}	75.65 _{0.013}	100.00 _{0.00}	96.80 _{0.015}	99.96 _{0.004}	94.51 _{0.019}
	ABR	36.57 _{0.018}	54.04 _{0.011}	84.28 _{0.007}	—	94.96 _{0.013}	86.75 _{0.013}	47.41 _{0.013}	—
	VB	57.40 _{0.015}	83.85 _{0.007}	70.25 _{0.008}	86.15 _{0.003}	77.54 _{0.018}	45.46 _{0.015}	66.50 _{0.013}	40.71 _{0.008}
	UVB	47.37 _{0.011}	74.52 _{0.005}	72.04 _{0.005}	85.67 _{0.002}	89.91 _{0.015}	59.48 _{0.011}	58.67 _{0.011}	35.33 _{0.005}

the ROC curve of BKS. We also experiment with the data generated from Case 2 under sample sizes $n = 10, 50$ and 100 , and present the corresponding results in Tables 3 and 4. A similar conclusion can be drawn as in Case 1. Besides, we provide the boxplot figures of FN and KL values, as well as the corresponding boxplots of TNR and TPR in two Cases in the Supplement Materials.

Further, we again find BKS tend to outperform WBKS, especially when J is large. This observation agrees with the relative performance of VB and UVB in Yu and Bien

4.2 Multivariate t distribution

Table 5: Average running time (seconds) of six precision matrix estimators under different combinations of (K, J) . The results are based on 100 replicates with sample size $n = 100$.

Method	$(K, J) = (10, 5)$	$(K, J) = (10, 10)$	$(K, J) = (20, 10)$	$(K, J) = (10, 20)$
WBKS	1.22	1.55	2.20	2.95
BKS	1.19	1.46	2.05	2.72
SMGM	1.41	1.91	3.03	3.44
ABR	8.13	118.86	1531.42	3044.67
VB	20.41	84.89	776.71	792.86
UVB	2.79	14.62	138.18	122.68

(2017). Intuitively, this is because the banded property is a special sparseness, where on each row, the elements farther away from the diagonal is more likely to be zero. $p(\mathbf{L}_j)$ incorporates this feature by repeated penalization, while $p_w(\mathbf{L}_j)$ downweights the penalty in each repetition, hence somewhat reduces the heavier penalty for elements farther away from the diagonal imposed by $p(\mathbf{L}_j)$. Such downweighting is especially harmful when J is large due to larger sparseness. Further, we provide the ROC curves in the right Figure 4.

Finally, we compare the computational complexity of these methods by examining their respective running times, measured in seconds. The results are presented in Table 5. It is evident that BKS is the fastest, and this advantage becomes particularly significant as J increases.

4.2 Multivariate t distribution

We will analyze the precision matrix used in Case 1 of the multivariate t distribution. We generate n independent longitudinal data $\{\text{vec}(\mathbf{Y}_i)\}_{i=1}^n$ from the $m = KJ$ dimensional

multivariate t distribution with $df = 4$ degrees of freedom. We consider sample sizes of $n = 10, 50$ and 100 . Similar to the multivariate normal distribution, we repeat the process 100 times for each sample size across four combinations of (K, J) .

We evaluate the accuracy of the Σ^{-1} estimator using FN and KL loss as metrics. Furthermore, we evaluate the performance of bandwidth recovery through the use of TNR and TPR measures. Table 6 presents the performance of matrix estimation accuracy in Case 1 for the multivariate t distribution. The results reveal that, for all estimators, the average values of FN and KL decrease as the sample size increases. Additionally, WBKS and BKS consistently exhibit superior performance across all scenarios, mirroring their performance with the multivariate normal distribution. To illustrate the bandwidth recovery performance of different estimators, we also report the averages and standard errors of TNR and TPR in Table 7. It is evident that WBKS and BKS consistently outperform the other methods, as reflected by their high and well-balanced TNR and TPR values.

5. Real Data Analysis

5.1 EEG data

We apply our method to analyze a public data set EEG from the UCI machine learning repository dataset (<http://archive.ics.uci.edu/ml/datasets/EEG+Database>). This data set contains $n = 122$ subjects, including $n_0 = 45$ alcoholic subjects ($z = 0$) and $n_1 = 77$ controls ($z = 1$). Each subject had 64 electrodes placed on his/her scalp,

5.1 EEG data

Table 6: Comparison in terms of FN and KL for different estimators of the precision matrix Σ^{-1} over 100 replications, in the form $\text{average}_{\text{standarderror}}$.

(K, J)		(10,10)	(10,20)	(20,10)	(20,20)	(10,10)	(10,20)	(20,10)	(20,20)
Methods		FN				KL			
10	WBKS	7.453 _{1.324}	13.376 _{0.65}	17.117 _{1.57}	14.598 _{2.80}	0.299 _{0.153}	0.521 _{0.177}	0.366 _{0.158}	0.276 _{0.176}
	BKS	7.154 _{1.424}	11.584 _{1.68}	16.944 _{1.67}	13.467 _{3.24}	0.284 _{0.153}	0.395 _{0.184}	0.360 _{0.159}	0.263 _{0.185}
	SMGM	10.482 _{0.54}	13.064 _{0.48}	18.866 _{0.99}	20.281 _{2.56}	0.510 _{0.316}	0.552 _{0.229}	0.483 _{0.225}	0.574 _{0.372}
50	WBKS	4.789 _{1.337}	5.424 _{1.183}	7.228 _{2.481}	9.638 _{2.551}	0.224 _{0.123}	0.195 _{0.103}	0.198 _{0.121}	0.217 _{0.126}
	BKS	4.748 _{1.361}	5.544 _{1.204}	7.646 _{2.484}	9.137 _{2.635}	0.223 _{0.124}	0.194 _{0.104}	0.201 _{0.122}	0.212 _{0.126}
	SMGM	6.003 _{3.175}	8.339 _{4.785}	12.862 _{7.49}	18.173 _{7.68}	0.451 _{0.472}	0.742 _{0.581}	0.566 _{0.428}	0.753 _{0.420}
	ABR	12.643 _{0.11}	14.021 _{0.17}	22.438 _{0.23}	—	0.638 _{0.034}	0.573 _{0.023}	0.728 _{0.030}	—
	VB	12.485 _{0.12}	14.155 _{0.11}	21.570 _{0.20}	25.040 _{0.23}	0.596 _{0.041}	0.586 _{0.023}	0.616 _{0.032}	0.672 _{0.028}
	UVB	12.576 _{0.15}	14.176 _{0.14}	22.107 _{0.29}	25.482 _{0.34}	0.633 _{0.047}	0.611 _{0.027}	0.705 _{0.057}	0.747 _{0.056}
100	WBKS	4.028 _{1.181}	4.884 _{1.172}	7.057 _{1.706}	8.644 _{2.069}	0.196 _{0.093}	0.189 _{0.082}	0.191 _{0.072}	0.200 _{0.087}
	BKS	4.120 _{0.171}	4.925 _{1.175}	7.001 _{1.704}	8.816 _{2.059}	0.196 _{0.093}	0.187 _{0.083}	0.190 _{0.072}	0.199 _{0.087}
	SMGM	4.412 _{2.728}	5.712 _{3.485}	8.219 _{4.455}	16.386 _{1.96}	0.307 _{0.365}	0.375 _{0.437}	0.269 _{0.261}	0.317 _{0.193}
	ABR	11.459 _{0.15}	13.346 _{0.12}	20.905 _{0.21}	—	0.455 _{0.028}	0.484 _{0.024}	0.532 _{0.021}	—
	VB	11.398 _{0.03}	13.675 _{0.12}	20.497 _{0.20}	24.212 _{0.16}	0.423 _{0.033}	0.492 _{0.029}	0.486 _{0.029}	0.553 _{0.024}
	UVB	11.479 _{0.14}	13.430 _{0.13}	22.065 _{0.23}	25.337 _{0.21}	0.452 _{0.030}	0.488 _{0.028}	0.660 _{0.042}	0.688 _{0.041}

and the measurements were taken at 256 Hz (3.9ms epoch) for 1 second. The electrode positions were located at standard sites, see Zhang et al. (1995) for specific names of these standard sites. In addition, each subject was exposed to two situations, either a single stimulus (S1) or two stimuli (S1 and S2), which were pictures of objects chosen from the 1980 Snodgrass and Vanderwart picture set. When two stimuli were shown, they were presented in either a matched condition where S1 was identical to S2 or in a non-matched condition, where S1 differed from S2. In this dataset, each subject completed 120

5.1 EEG data

Table 7: Comparison in terms of TNR and TPR for different estimators of the precision matrix Σ^{-1} over 100 replications, in the form $\text{average}_{\text{standarderror}}$.

(K, J)		(10,10)	(10,20)	(20,10)	(20,20)	(10,10)	(10,20)	(20,10)	(20,20)
Methods		TNR				TPR			
10	WBKS	90.64 _{0.027}	99.23 _{0.005}	95.61 _{0.020}	93.42 _{0.013}	86.47 _{0.043}	37.02 _{0.033}	57.77 _{0.053}	83.40 _{0.031}
	BKS	90.51 _{0.031}	98.71 _{0.006}	95.76 _{0.020}	96.06 _{0.010}	86.71 _{0.043}	56.53 _{0.052}	56.99 _{0.041}	87.00 _{0.028}
	SMGM	97.81 _{0.013}	98.80 _{0.007}	98.14 _{0.011}	98.67 _{0.011}	45.50 _{0.178}	35.28 _{0.093}	35.43 _{0.133}	36.83 _{0.263}
50	WBKS	81.25 _{0.051}	88.16 _{0.022}	83.03 _{0.046}	92.84 _{0.017}	97.64 _{0.022}	95.59 _{0.021}	98.36 _{0.019}	97.21 _{0.018}
	BKS	84.77 _{0.046}	90.83 _{0.024}	87.51 _{0.041}	93.46 _{0.018}	97.91 _{0.020}	95.37 _{0.021}	98.18 _{0.019}	97.97 _{0.017}
	SMGM	78.30 _{0.135}	89.04 _{0.116}	79.14 _{0.210}	94.03 _{0.090}	72.28 _{0.393}	49.20 _{0.455}	51.69 _{0.465}	32.73 _{0.419}
	ABR	89.41 _{0.030}	93.47 _{0.015}	93.57 _{0.038}	—	32.45 _{0.028}	31.98 _{0.025}	24.02 _{0.045}	—
	VB	95.98 _{0.017}	97.21 _{0.009}	97.04 _{0.012}	98.62 _{0.005}	25.49 _{0.021}	23.12 _{0.014}	21.71 _{0.018}	16.89 _{0.014}
	UVB	91.67 _{0.014}	95.23 _{0.006}	96.76 _{0.008}	98.68 _{0.003}	23.41 _{0.020}	23.50 _{0.013}	13.08 _{0.006}	11.00 _{0.004}
100	WBKS	84.51 _{0.045}	89.08 _{0.021}	88.65 _{0.038}	93.38 _{0.018}	99.11 _{0.015}	98.06 _{0.015}	99.61 _{0.011}	99.24 _{0.017}
	BKS	85.59 _{0.044}	91.47 _{0.023}	88.32 _{0.039}	94.58 _{0.017}	98.09 _{0.016}	99.42 _{0.015}	99.61 _{0.012}	99.17 _{0.011}
	SMGM	69.60 _{0.125}	84.99 _{0.073}	68.42 _{0.114}	98.30 _{0.006}	86.62 _{0.312}	81.10 _{0.355}	88.29 _{0.293}	78.04 _{0.186}
	ABR	80.21 _{0.029}	90.80 _{0.013}	85.78 _{0.024}	—	45.57 _{0.036}	40.91 _{0.032}	35.54 _{0.020}	—
	VB	87.19 _{0.025}	95.38 _{0.012}	93.79 _{0.014}	97.64 _{0.007}	42.45 _{0.022}	31.82 _{0.015}	33.00 _{0.013}	25.81 _{0.009}
	UVB	81.36 _{0.013}	91.70 _{0.006}	95.61 _{0.007}	97.92 _{0.004}	39.02 _{0.020}	33.61 _{0.017}	15.79 _{0.009}	13.80 _{0.009}

trials under each situation. Taking averages over 120 trials, each subject has a 64×256 measurement matrix. Following Qian et al. (2021), we take the average of every 32 measurements to reduce the time dimension from 256 to 8 and obtain a 64×8 matrix for each subject. This eventually leads to a data set $\mathbf{Y}_i \in \mathcal{R}^{K \times J} (i = 1, \dots, n)$ with $K = 64$, $J = 8$ and $n = 122$. In addition, we also have a class label $z_i \in \{0, 1\}$ for $i = 1, \dots, n$.

Similar to Qian et al. (2021), we aim to classify these subjects into two classes, alcoholic (class 0) and control (class 1), based on the information in $\mathbf{Y}_i$'s. We consider two

methods, linear discriminant analysis (LDA) and quadratic discriminant analysis (QDA) to perform the classification. LDA classifies subject i to class 0 if $\delta_{\text{LDA}}^{(0)}(\mathbf{y}_i) > \delta_{\text{LDA}}^{(1)}(\mathbf{y}_i)$, otherwise to class 1, where

$$\delta_{\text{LDA}}^{(c)}(\mathbf{y}_i) = \mathbf{y}_i^T \hat{\Sigma}^{-1} \hat{\boldsymbol{\mu}}^{(c)} - \frac{1}{2} (\hat{\boldsymbol{\mu}}^{(c)})^T \hat{\Sigma}^{-1} \hat{\boldsymbol{\mu}}^{(c)} + \log \hat{\pi}^{(c)}$$

for $c = 0, 1$ and $\hat{\pi}^{(c)}$ is the estimated proportion of group c . Here, $\hat{\Sigma}$ is the estimated overall covariance matrix, and $\hat{\boldsymbol{\mu}}^{(c)}$ is the estimated mean in group c . Similarly, QDA classifies subject i to class 0 if $\delta_{\text{QDA}}^{(0)}(\mathbf{y}_i) > \delta_{\text{QDA}}^{(1)}(\mathbf{y}_i)$, otherwise to class 1, where

$$\delta_{\text{QDA}}^{(c)}(\mathbf{y}_i) = \mathbf{y}_i^T (\hat{\Sigma}^{(c)})^{-1} \hat{\boldsymbol{\mu}}^{(c)} - \frac{1}{2} (\hat{\boldsymbol{\mu}}^{(c)})^T (\hat{\Sigma}^{(c)})^{-1} \hat{\boldsymbol{\mu}}^{(c)} + \log \hat{\pi}^{(c)}.$$

for $c = 0, 1$. Here, $\hat{\Sigma}^{(c)}$ is the estimated covariance matrix in group c . To implement these methods, we randomly sample 70% of the data to form a training set and use the remaining 30% as the testing data. We used sample proportions to form $\hat{\pi}^{(c)}$, sample averages to form $\hat{\boldsymbol{\mu}}^{(c)}$, ($c = 0, 1$), and used WBKS, BKS, SMGM, ABR, VB and UVB to estimate Σ^{-1} in LDA and $(\Sigma^{(c)})^{-1}$, ($c = 0, 1$) in QDA. To select the tuning parameters in these methods, we used a five-fold crossvalidation.

The upper row of Figure 5 in Supplement shows the estimated precision matrix $\hat{\Sigma}^{-1}$ obtained from WBKS, BKS and SMGM. The plots from ABR, VB, and UVB are excluded since they are already provided in Qian et al. (2021).

WBKS and BKS lead to a block-banded precision matrix, which agrees with the

Table 8: The average classification errors of different methods over 10 random train-test splits.

	WBKS	BKS	SMGM	ABR	VB	UVB
LDA	0.19	0.18	0.24	0.25	0.24	0.28
QDA	0.18	0.18	–	0.24	0.29	0.29

general conclusion from ABR (Qian et al., 2021). In contrast, SMGM exhibits a lack of time correlation during the first three time points, followed by inter-correlation within the remaining five time points. This pattern is counter-intuitive. Finally, UVB and VB lead to simple banded precision matrix estimation, which is also unrealistic due to the spacial correlation that tends to persist across time. We further plot the resulting $\hat{\mathbf{R}}^{-1}$ and $\hat{\mathbf{W}}^{-1}$ by BKS in the lower row of Figure 5. We can clearly see the banded feature of $\hat{\mathbf{R}}^{-1}$ and the sparseness of $\hat{\mathbf{W}}^{-1}$. To evaluate the performance of these methods, we compute the classification errors on the testing data. The results in Table 6 contain the average testing data classification errors over 10 random train-test splits. It is clear that WBKS and BKS outperform the other methods. Among the remaining methods, ABR is the winner, indicating that the true precision matrix is close to have the block-banded structure. However, ABR is inferior to WBKS and BKS, possibly because it contains too many parameters (Qian et al., 2021). Note that due to the relative small sample size, SMGM fails to produce a result when performing QDA.

5.2 ADHD data

In this section, we will analyze a dataset pertaining Attention Deficit Hyperactivity Disorder (ADHD). ADHD is a prevalent mental disorder observed in children and adolescents, characterized by symptoms such as distractibility, impulsivity, and restlessness. Functional Magnetic Resonance Imaging (fMRI) data at rest from the ADHD-200 sample dataset (http://www.nitrc.org/frs/?group_id=383) were collected by Oregon Health and Science University. The data were processed using the Automated Anatomical Labeling (AAL) software package and a dedicated digital atlas designed for the human brain.

The ADHD dataset consists of 42 subjects who belong to the typical developmental control groups. These children are used as a baseline for comparing with individuals diagnosed with Attention Deficit Hyperactivity Disorder (ADHD). Brain activity is measured by detecting changes in blood flow correlated with low-frequency Blood Oxygen Level Dependent (BOLD) signals. Tzourio-Mazoyer et al. (2002) provided a detailed description of brain region segmentation. Each individual's brain was monitored in 116 regions of interest (ROIs), and the signals from these 116 ROIs for each child were recorded over 74 scans. Consequently, we obtained multivariate longitudinal data $\mathbf{Y}_i \in \mathcal{R}^{K \times J}$ ($i = 1, \dots, n$), where $K = 116$, $J = 74$, $n = 42$. Leng and Pan (2018) assumed a Kronecker structure for the covariance matrix of this data and estimated it using large-dimensional random matrix theory. Their analysis revealed that the temporal covariance matrix exhibits a banded structure, where correlations are strongest near the main diagonal and gradually

decrease as they move away from it. However, the covariance matrix of the 116 brain regions demonstrates sparsity. Additionally, their method does not incorporate the banded structure information that exists in the longitudinal data.

To further explore the structural information among the K brain regions and J measurements in the ADHD data, we applied our proposed method to estimate the precision matrix of this dataset. Obtaining the ABR, UVB, and VB estimators for the data is challenging due to the high temporal and variable dimensions. Furthermore, we utilized the sparse matrix graphical model with Kronecker structure, as proposed by Leng and Tang (2012), to estimate the precision matrix for SMGM. Figure 6 shows the plotted structural information of the precision matrix estimator obtained using our method, while Figure 7 presents the structural information of the precision matrix estimator for SMGM. Due to the high dimensionality of $KJ = 8584$, which is conducive to detailed structural representation, we present the precision matrix separately for the temporal dimension (left) and the variable dimension (right).

From Figure 6, it is evident that the precision matrix in the temporal dimension exhibits an adaptive banded structure. The band width starts wider at the beginning and gradually narrows. The conditional correlations among the 116 brain regions in the variable dimension exhibit sparsity, with a distinct block structure observed in the top-left and bottom-right corners. However, the structural information in other regions appears relatively scattered, which can be attributed to the division of brain regions. Figure 7 displays the estimated precision matrix by SMGM, which is observed to be a diagonal

matrix, indicating a lack of captured structural information within the data. This finding aligns with the conclusions of Leng and Pan (2018). Additionally, Figure 8 presents the correlation matrix in the temporal and variable dimensions obtained by inverting the precision matrix. The results reveal a banded correlation structure among the 74 time points and a localized block-structured correlation pattern among the 116 brain regions. This can be attributed to the collective influence of specific local brain regions on certain human behaviors. In conclusion, the precision matrix estimator obtained from our proposed method effectively captures the conditional correlations among the 74 time points. Furthermore, we observe that the conditional correlations among the 116 brain regions exhibit sparsity, which aligns with the observed characteristics in reality.

6. Conclusions

We have proposed a new precision matrix estimator named BKS. BKS takes advantage of the fact that the precision matrix is the Kronecker product of a banded matrix and a sparse matrix. It incorporates the matrix bandedness by considering its Cholesky decomposition and imposing a new penalty which increasingly encourages zeros for elements farther away from the matrix diagonal. Matrix sparsity is enforced by applying the standard lasso penalty. BKS also guarantees a positive definite estimator of the precision matrix. BKS is easy to implement. It optimizes a biconvex objective function, and is achieved through an alternative optimization algorithm named ACS. ACS comprises of two repeating steps, with each step solving a convex optimization problem using the glas-

so and ADMM algorithms, respectively. This approach ensures computational efficiency. We show the algorithmic convergence of ACS and establish the statistical convergence rate of BKS. We find that BKS exhibits the same convergence rate as SMGM when the data is normally distributed. However, the application and advantageous characteristics of BKS extend beyond normality. We also demonstrate that BKS has the ability to recover the true bandwidths of the banded matrix with a probability close to 1. Both simulation studies and real data applications consistently indicate that BKS exhibits favorable performance overall.

Supplementary Material

The supplementary materials provides the proofs of the lemmas, Theorem 1, Theorem 2, Theorem 3 and Theorem 4, and these figures in simulation and real data studies.

Acknowledgments

We are very grateful to the Editor, Associate Editor and referees, as well as our financial sponsors for their insightful comments and suggestions that have improved the manuscript significantly. This work was supported by the Key Program of the National Natural Science Foundation of China [grant number 12431009].

References

- Aston, J. A. D., D. Pigoli, and S. Tavakoli (2017). Tests for separability in nonparametric covariance operators of random surfaces. *The Annals of Statistics* 45(4), 1431–1461.
- Banerjee, O., L. El Ghaoui, and A. d’Aspremont (2008). Model selection through sparse maximum likelihood estimation for multivariate gaussian or binary data. *The Journal of Machine Learning Research* 9(3), 485–516.
- Bickel, P. J. and E. Levina (2008a). Covariance regularization by thresholding. *Annals of Statistics* 36, 2577–2604.
- Bickel, P. J. and E. Levina (2008b). Regularized estimation of large covariance matrices. *Annals of Statistics* 36(1), 199–227.
- Bien, J., F. Bunea, and L. Xiao (2016). Convex banding of the covariance matrix. *Journal of the American Statistical Association* 111(514), 834–845.
- Boyd, S., N. Parikh, E. Chu, B. Peleato, and J. Eckstein (2011). Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends in Machine learning* 3(1), 1–122.
- Cai, T. T., C. H. Zhang, and H. H. Zhou (2010). Optimal rates of convergence for covariance matrix estimation. *Annals of Statistics* 38(1), 2118–2144.
- Dai, D., C. Hao, S. Jin, and Y. Liang (2023). Regularized estimation of kronecker structured covariance matrix using modified cholesky decomposition. *Journal of Statistical Computation and Simulation* 95(5), 1–26.
- Friedman, J., T. Hastie, and R. Tibshirani (2008). Sparse inverse covariance estimation with the graphical lasso. *Biostatistics* 9(3), 432–441.

REFERENCES

- Furrer, R. and T. Bengtsson (2007). Estimation of high-dimensional prior and posterior covariance matrices in kalman filter variants. *Journal of Multivariate Analysis* 98(1), 227–255.
- Gorski, J., F. Pfeuffer, and K. Klamroth (2007). Biconvex sets and optimization with biconvex functions: a survey and extensions. *Mathematical Methods of Operations Research* 66, 373408.
- Greenewald, K. and A. O. Hero (2015). Robust kronecker product pca for spatio-temporal covariance estimation. *IEEE Transactions on Signal Processing* 63(23), 6368–6378.
- Jenatton, R., J.-Y. Audibert, and F. Bach (2011). Structured variable selection with sparsity-inducing norms. *Journal of Machine Learning Research* 12, 2777–2824.
- Kim, C. and D. L. Zimmerman (2012). Unconstrained models for the covariance structure of multivariate longitudinal data. *Journal of Multivariate Analysis* 107, 104–118.
- Lee, K., H. Cho, M. Kwak, and E. Jang (2020). Estimation of covariance matrix of multivariate longitudinal data using modified cholesky and hypersphere decompositions. *Journal of Multivariate Analysis* 76(1), 75–86.
- Leng, C. and G. Pan (2018). Covariance estimation via sparse kronecker structures. *Bernoulli* 24(4B), 3833–3863.
- Leng, C. and C. Tang (2012). Sparse matrix graphical models. *Journal of the American Statistical Association* 107(499), 1187–1200.
- Qian, F., Y. Chen, and W. Zhang (2020). Regularized estimation of precision matrix for high-dimensional multivariate longitudinal data. *Journal of Multivariate Analysis* 176, 104580.
- Qian, F., W. Zhang, and Y. Chen (2021). Adaptive banding covariance estimation for high-dimensional multivariate longitudinal data. *Canadian Journal of Statistics* 49(3), 906–938.
- Rothman, A., E. Levina, and J. Zhu (2008). Sparse permutation invariant covariance estimation. *Electronic*

REFERENCES

- Journal of Statistics* 2, 494–515.
- Tsiligkaridis, T. and A. O. Hero (2013). Covariance estimation in high dimensions via kronecker product expansions. *IEEE Transactions on Signal Processing* 61(21), 5347–5360.
- Yu, G. and J. Bien (2017). Learning local dependence in ordered data. *Journal of Machine Learning Research* 18(42), 1–60.
- Yuan, M. and Y. Lin (2006). Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society: Statistical Methodology Series B* 68, 49–67.
- Yuan, M. and Y. Lin (2007). Model selection and estimation in the gaussian graphical model. *Biometrika* 94(1), 19–35.
- Zhang, X. L., H. Begleiter, B. Porjesz, W. Wang, and A. Litke (1995). Event related potentials during object recognition tasks. *Brain research bulletin* 6, 531–538.
- Zhang, Y., W. Shen, and K. Dehan (2023). Covariance estimation for matrix-valued data. *Journal of the American Statistical Association* 118(554), 2620–2631.
- KLAS and Department of Mathematics and Statistics, Northeast Normal University, 5268 Renmin Street, Changchun, China
E-mail: (liangch801@nenu.edu.cn)
- School of Mathematical Sciences, Capital Normal University, 105 North Xisanhuan Road, Beijing, China
E-mail: (wenqingma@cnu.edu.cn)
- Department of Statistics, The Pennsylvania State University, University Park, Pennsylvania 16802, U.S.A.
E-mail: (yanyuanma@gmail.com)