

SO FAR: Large-Scale Association Network Learning

Yoshimasa Uematsu^{ID}, Yingying Fan, Kun Chen, Jinchi Lv^{ID}, and Wei Lin^{ID}

Abstract—Many modern big data applications feature large scale in both numbers of responses and predictors. Better statistical efficiency and scientific insights can be enabled by understanding the large-scale response-predictor association network structures via layers of sparse latent factors ranked by importance. Yet sparsity and orthogonality have been two largely incompatible goals. To accommodate both features, in this paper, we suggest the method of sparse orthogonal factor regression (SO FAR) via the sparse singular value decomposition with orthogonality constrained optimization to learn the underlying association networks, with broad applications to both unsupervised and supervised learning tasks, such as biclustering with sparse singular value decomposition, sparse principal component analysis, sparse factor analysis, and sparse vector autoregression analysis. Exploiting the framework of convexity-assisted nonconvex optimization, we derive nonasymptotic error bounds for the suggested procedure characterizing the theoretical advantages. The statistical guarantees are powered by an efficient SO FAR algorithm with convergence property. Both computational and theoretical advantages of our procedure are demonstrated with several simulations and real data examples.

Index Terms—Big data, large-scale association network, simultaneous response and predictor selection, latent factors, sparse singular value decomposition, orthogonality constrained optimization, nonconvex statistical learning.

I. INTRODUCTION

THE genetics of gene expression variation may be complex due to the presence of both local and distant

genetic effects and shared genetic components across multiple genes [14], [18]. A useful statistical analysis in such studies is to simultaneously classify the genetic variants and gene expressions into groups that are associated. For example, in a yeast expression quantitative trait loci (eQTLs) mapping analysis, the goal is to understand how the eQTLs, which are regions of the genome containing DNA sequence variants, influence the expression level of genes in the yeast MAPK signaling pathways. Extensive genetic and biochemical analysis has revealed that there are a few functionally distinct signaling pathways of genes [14], [36], suggesting that the association structure between the eQTLs and the genes is of low rank. Each signaling pathway involves only a subset of genes, which are regulated by only a few genetic variants, suggesting that each association between the eQTLs and the genes is sparse in both the input and the output (or in both the responses and the predictors), and the pattern of sparsity should be pathway specific. Moreover, it is known that the yeast MAPK pathways regulate and interact with each other [36]. The complex genetic structures described above clearly call for a joint statistical analysis that can reveal multiple distinct associations between subsets of genes and subsets of genetic variants. If we treat the genetic variants and gene expressions as the predictors and responses, respectively, in a multivariate regression model, the task can then be carried out by seeking a sparse representation of the coefficient matrix and performing predictor and response selection simultaneously. The problem of large-scale response-predictor association network learning is indeed of fundamental importance in many modern big data applications featuring large scale in both numbers of responses and predictors.

Observing n independent pairs $(\mathbf{x}_i, \mathbf{y}_i)$, $i = 1, \dots, n$, with $\mathbf{x}_i \in \mathbb{R}^p$ the covariate vector and $\mathbf{y}_i \in \mathbb{R}^q$ the response vector, motivated from the above applications we consider the following multivariate regression model

$$\mathbf{Y} = \mathbf{X}\mathbf{C}^* + \mathbf{E}, \quad (1)$$

where $\mathbf{Y} = (\mathbf{y}_1, \dots, \mathbf{y}_n)^T \in \mathbb{R}^{n \times q}$ is the response matrix, $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_n)^T \in \mathbb{R}^{n \times p}$ is the predictor matrix, $\mathbf{C}^* \in \mathbb{R}^{p \times q}$ is the true regression coefficient matrix, and $\mathbf{E} = (\mathbf{e}_1, \dots, \mathbf{e}_n)^T$ is the error matrix. To model the sparse relationship between the responses and the predictors as in the yeast eQTLs mapping analysis, we exploit the following singular value decomposition (SVD) of the coefficient matrix

$$\mathbf{C}^* = \mathbf{U}^* \mathbf{D}^* \mathbf{V}^{*T} = \sum_{j=1}^r d_j^* \mathbf{u}_j^* \mathbf{v}_j^{*T}, \quad (2)$$

Manuscript received August 29, 2017; revised March 29, 2019; accepted March 30, 2019. Date of publication April 11, 2019; date of current version July 12, 2019. This work was supported in part by Grant-in-Aid for JSPS Fellows under Grant 26-1905, in part by NIH under Grant 1R01GM131407-01 and Grant U01 HL114494, in part by NSF CAREER Awards under Grant DMS-095531 and Grant DMS-1150318, in part by NSF under Grant DMS-1613295 and Grant IIS-1718798, in part by a grant from the Simons Foundation, in part by the Adobe Data Science Research Award, in part by NSFC under Grant 11671018 and Grant 71532001, and in part by the National Key R&D Program of China under Grant 2016YFC0207703.

Y. Uematsu was with USC Marshall, Los Angeles, CA 90089 USA. He is now with the Department of Economics and Management, Tohoku University, Sendai 980-8576, Japan (e-mail: yoshimasa.uematsu.e7@tohoku.ac.jp).

Y. Fan is with Data Sciences and Operations Department, Marshall School of Business, University of Southern California at Los Angeles, Los Angeles, CA 90089 USA (e-mail: fanyingy@marshall.usc.edu).

K. Chen is with the Department of Statistics, University of Connecticut, Storrs, CT 06269 USA (e-mail: kun.chen@uconn.edu).

J. Lv is with the Data Sciences and Operations Department, Marshall School of Business, University of Southern California at Los Angeles, Los Angeles, CA 90089 USA (e-mail: jinchilv@marshall.usc.edu).

W. Lin is with the School of Mathematical Sciences, Peking University, Beijing 100871, China, and also with the Center for Statistical Science, Peking University, Beijing 100871, China (e-mail: weilin@math.pku.edu.cn).

Communicated by K.-R. Müller, Associate Editor for Machine Learning.

This paper has supplementary downloadable material available at <http://ieeexplore.ieee.org>.

Color versions of one or more of the figures in this paper are available online at <http://ieeexplore.ieee.org>.

Digital Object Identifier 10.1109/TIT.2019.2909889

where $1 \leq r \leq \min(p, q)$ is the rank of matrix $\mathbf{C}^*$, $\mathbf{D}^* = \text{diag}(d_1^*, \dots, d_r^*)$ is a diagonal matrix of nonzero singular values, and $\mathbf{U}^* = (\mathbf{u}_1^*, \dots, \mathbf{u}_r^*) \in \mathbb{R}^{p \times r}$ and $\mathbf{V}^* = (\mathbf{v}_1^*, \dots, \mathbf{v}_r^*) \in \mathbb{R}^{q \times r}$ are the orthonormal matrices of left and right singular vectors, respectively. Here, we assume that $\mathbf{C}^*$ is low-rank with only r nonzero singular values, and the matrices $\mathbf{U}^*$ and $\mathbf{V}^*$ are sparse.

Under the sparse SVD structure (2), model (1) can be rewritten as

$$\tilde{\mathbf{Y}} = \tilde{\mathbf{X}}\mathbf{D}^* + \tilde{\mathbf{E}},$$

where $\tilde{\mathbf{Y}} = \mathbf{Y}\mathbf{V}^*$, $\tilde{\mathbf{X}} = \mathbf{X}\mathbf{U}^*$, and $\tilde{\mathbf{E}} = \mathbf{E}\mathbf{V}^* \in \mathbb{R}^{n \times r}$ are the matrices of latent responses, predictors, and random errors, respectively. The associations between the predictors and responses are thus diagonalized under the pairs of transformations specified by $\mathbf{U}^*$ and $\mathbf{V}^*$. When $\mathbf{C}^*$ is of low rank, this provides an appealing *low-dimensional latent model* interpretation for model (1). Further, note that the latent responses and predictors are linear combinations of the original responses and predictors, respectively. Thus, the interpretability of the SVD can be enhanced if we require that the left and right singular vectors be sparse so that each latent predictor/response involves only a small number of the original predictors/responses, thereby performing the task of variable selection among the predictors/responses, as needed in the yeast eQTLs analysis.

The above model (1) with low-rank coefficient matrix has been commonly adopted in the literature. In particular, the reduced rank regression [1], [39], [55] is an effective approach to dimension reduction by constraining the coefficient matrix $\mathbf{C}^*$ to be of low rank. Bunea *et al.* [15] proposed a rank selection criterion that can be viewed as an L_0 regularization on the singular values of $\mathbf{C}^*$. The popularity of L_1 regularization methods such as the Lasso [60] led to the development of nuclear norm regularization in multivariate regression [66]. Chen *et al.* [21] proposed an adaptive nuclear norm penalization approach to bridge the gap between L_0 and L_1 regularization methods and combine some of their advantages. With the additional SVD structure (2), Chen *et al.* [19] proposed a new estimation method with a correctly specified rank by imposing a weighted L_1 penalty on each rank-1 SVD layer for the classical setting of fixed dimensionality. Chen and Huang [22] and Bunea *et al.* [16] explored a low-rank representation of $\mathbf{C}^*$ in which the rows of $\mathbf{C}^*$ are sparse; however, their approaches do not impose sparsity on the right singular vectors and, hence, are inapplicable to settings with high-dimensional responses where response selection is highly desirable.

Recently, there have been some new developments in sparse and low-rank regression problems. Ma and Sun [50] studied the properties of row-sparse reduced-rank regression model with nonconvex sparsity-inducing penalties, and later Ma *et al.* [49] extended their work to two-way sparse reduced-rank regression. Chen and Huang [23] extended the row-sparse reduced-rank regression by incorporating covariance matrix estimation, and the authors mainly focused on computational issues. Lian *et al.* [46] proposed a semiparametric reduced-rank regression with a sparsity penalty on the coefficient

matrix itself. Goh *et al.* [33] studied the Bayesian counterpart of the row/column-sparse reduced-rank regression and established its posterior consistency. However, none of these works considered the possible entrywise sparsity in the SVD of the coefficient matrix. The sparse and low-rank regression models have also been applied in various fields to solve important scientific problems. To name a few, Chen *et al.* [20] applied a sparse and low-rank bi-linear model for the task of source-sink reconstruction in marine ecology, Zhu *et al.* [70] used a Bayesian low-rank model for associating neuroimaging phenotypes and genetic markers, and Ma *et al.* [48] used a threshold SVD regression model for learning regulatory relationships in genomics.

In view of the key role that the sparse SVD plays for simultaneous dimension reduction and variable selection in model (1), in this paper we suggest a unified regularization approach to estimating such a sparse SVD structure. Our proposal successfully meets three key methodological challenges that are posed by the complex structural constraints on the SVD. First, sparsity and orthogonality are two largely incompatible goals and would seem difficult to be accommodated within a single framework. For instance, a standard orthogonalization process such as QR factorization will generally destroy the sparsity pattern of a matrix. Previous methods either relaxed the orthogonality constraint to allow efficient search for sparsity patterns [19], or avoided imposing both sparsity and orthogonality requirements on the same factor matrix [16], [22]. To resolve this issue, we formulate our approach as an orthogonality constrained regularization problem, which yields *simultaneously sparse and orthogonal* factor matrices in the SVD. Second, we employ the nuclear norm penalty to encourage sparsity among the singular values and achieve rank reduction. As a result, our method produces a continuous solution path, which facilitates rank parameter tuning and distinguishes it from the L_0 regularization method adopted by Bunea *et al.* [16]. Third, unlike rank-constrained estimation, the nuclear norm penalization approach makes the estimation of singular vectors more intricate, since one does not know a priori which singular values will vanish and, hence, which pairs of left and right singular vectors are unidentifiable. Noting that the degree of identifiability of the singular vectors increases with the singular value, we propose to penalize the singular vectors weighted by singular values, which proves to be meaningful and effective. Combining these aspects, we introduce *sparse orthogonal factor regression* (SOFAR), a novel regularization framework for high-dimensional multivariate regression. While respecting the orthogonality constraint, we allow the sparsity-inducing penalties to take a general, flexible form, which includes special cases that adapt to the entrywise and rowwise sparsity of the singular vector matrices, resulting in a nonconvex objective function for the SOFAR method.

In addition to the aforementioned three methodological challenges, the nonconvexity of the SOFAR objective function also poses important algorithmic and theoretical challenges in obtaining and characterizing the SOFAR estimator. To address these challenges, we suggest a two-step approach exploiting the framework of convexity-assisted nonconvex

optimization (CANO) to obtain the SOFAR estimator. More specifically, in the first step we minimize the L_1 -penalized squared loss for the multivariate regression (1) to obtain an initial estimator. Then in the second step, we minimize the SOFAR objective function in an asymptotically shrinking neighborhood of the initial estimator. Thanks to the convexity of its objective function, the initial estimator can be obtained effectively and efficiently. Yet since the finer sparsity structure imposed through the sparse SVD (2) is completely ignored in the first step, the initial estimator meets none of the aforementioned three methodological challenges. Nevertheless, since it is theoretically guaranteed that the initial estimator is not far away from the true coefficient matrix $\mathbf{C}^*$ with asymptotic probability one, searching in an asymptotically shrinking neighborhood of the initial estimator significantly alleviates the nonconvexity issue of the SOFAR objective function. In fact, under the framework of CANO we derive nonasymptotic bounds for the prediction, estimation, and variable selection errors of the SOFAR estimator characterizing the theoretical advantages. In implementation, to disentangle the sparsity and orthogonality constraints we develop an efficient SOFAR algorithm and establish its convergence properties.

Our suggested SOFAR method for large-scale association network learning is in fact connected to a variety of statistical methods in both unsupervised and supervised multivariate analysis. For example, the sparse SVD and sparse principal component analysis (PCA) for a high-dimensional data matrix can be viewed as unsupervised versions of our general method. Other prominent examples include sparse factor models, sparse canonical correlation analysis [63], and sparse vector autoregressive (VAR) models for high-dimensional time series. See Section II-B for more details on these applications and connections.

The rest of the paper is organized as follows. Section II introduces the SOFAR method and discusses its applications to several unsupervised and supervised learning tasks. We present the nonasymptotic properties of the method in Section III. Section IV develops an efficient optimization algorithm and discusses its convergence and tuning parameter selection. We provide several simulation and real data examples in Section V. All the proofs of main results and technical details are detailed in the Supplementary Material. An associated R package implementing the suggested method is available at <https://cran.r-project.org/package=rrpack>.

II. LARGE-SCALE ASSOCIATION NETWORK LEARNING VIA SOFAR

A. Sparse Orthogonal Factor Regression

To estimate the sparse SVD of the true regression coefficient matrix $\mathbf{C}^*$ in model (1), we start by considering an estimator of the form $\mathbf{UDV}^T$, where $\mathbf{D} = \text{diag}(d_1, \dots, d_m) \in \mathbb{R}^{m \times m}$ with $d_1 \geq \dots \geq d_m \geq 0$ and $1 \leq m \leq \min\{p, q\}$ is a diagonal matrix of singular values, and $\mathbf{U} = (\mathbf{u}_1, \dots, \mathbf{u}_m) \in \mathbb{R}^{p \times m}$ and $\mathbf{V} = (\mathbf{v}_1, \dots, \mathbf{v}_m) \in \mathbb{R}^{q \times m}$ are orthonormal matrices of left and right singular vectors, respectively. Although it is always possible to take $m = \min(p, q)$ without prior knowledge of the rank r , it is often sufficient in practice to take a small m that is slightly larger than the expected rank (estimated by some

procedure such as in Bunea *et al.* [15]), which can dramatically reduce computation time and space. Throughout the paper, for any matrix $\mathbf{M} = (m_{ij})$ we denote by $\|\mathbf{M}\|_F$, $\|\mathbf{M}\|_1$, $\|\mathbf{M}\|_\infty$, and $\|\mathbf{M}\|_{2,1}$ the Frobenius norm, entrywise L_1 -norm, entrywise L_∞ -norm, and rowwise $(2, 1)$ -norm defined, respectively, as $\|\mathbf{M}\|_F = (\sum_{i,j} m_{ij}^2)^{1/2}$, $\|\mathbf{M}\|_1 = \sum_{i,j} |m_{ij}|$, $\|\mathbf{M}\|_\infty = \max_{i,j} |m_{ij}|$, and $\|\mathbf{M}\|_{2,1} = \sum_i (\sum_j m_{ij}^2)^{1/2}$. We also denote by $\|\cdot\|_2$ the induced matrix norm (operator norm).

As mentioned in the Introduction, we employ the nuclear norm penalty to encourage sparsity among the singular values, which is exactly the entrywise L_1 penalty on $\mathbf{D}$. Penalization directly on $\mathbf{U}$ and $\mathbf{V}$, however, is inappropriate since the singular vectors are not equally identifiable and should not be subject to the same amount of regularization. Singular vectors corresponding to larger singular values can be estimated more accurately and should contribute more to the regularization, whereas those corresponding to vanishing singular values are unidentifiable and should play no role in the regularization. Therefore, we propose an *importance weighting* by the singular values and place sparsity-inducing penalties on the weighted versions of singular vector matrices, $\mathbf{UD}$ and $\mathbf{VD}$. Note also that our goal is to estimate not only the low-rank matrix $\mathbf{C}^*$ but also the factor matrices $\mathbf{D}^*$, $\mathbf{U}^*$, and $\mathbf{V}^*$. Taking into account these points, we consider the orthogonality constrained optimization problem

$$\begin{aligned} & (\hat{\mathbf{D}}, \hat{\mathbf{U}}, \hat{\mathbf{V}}) \\ &= \arg \min_{\mathbf{D}, \mathbf{U}, \mathbf{V}} \left\{ \frac{1}{2} \|\mathbf{Y} - \mathbf{XUDV}^T\|_F^2 \right. \\ & \quad \left. + \lambda_d \|\mathbf{D}\|_1 + \lambda_a \rho_a(\mathbf{UD}) + \lambda_b \rho_b(\mathbf{VD}) \right\} \\ & \text{subject to } \mathbf{U}^T \mathbf{U} = \mathbf{I}_m, \quad \mathbf{V}^T \mathbf{V} = \mathbf{I}_m, \end{aligned} \quad (3)$$

where $\rho_a(\cdot)$ and $\rho_b(\cdot)$ are penalty functions to be clarified later, and $\lambda_d, \lambda_a, \lambda_b \geq 0$ are tuning parameters that control the strengths of regularization. We call this regularization method *sparse orthogonal factor regression* (SOFAR) and the regularized estimator $(\hat{\mathbf{D}}, \hat{\mathbf{U}}, \hat{\mathbf{V}})$ the SOFAR estimator. Note that $\rho_a(\cdot)$ and $\rho_b(\cdot)$ can be equal or distinct, depending on the scientific question and the goals of variable selection. Letting $\lambda_d = \lambda_b = 0$ while setting $\rho_a(\cdot) = \|\cdot\|_{2,1}$ reduces the SOFAR estimator to the sparse reduced-rank estimator of Chen and Huang [22]. In view of our choices of $\rho_a(\cdot)$ and $\rho_b(\cdot)$, although $\mathbf{D}$ appears in all three penalty terms, rank reduction is achieved mainly through the first term, while variable selection is achieved through the last two terms under necessary scalings by $\mathbf{D}$.

We note that one major advantage of SOFAR is that the final estimates satisfy the orthogonality constraints in (3). This is also the major distinction of our method from many existing ones. The orthogonality constraints are motivated from a combination of practical, methodological, and theoretical considerations. On the practical side, the orthogonality constraints maximize the separation of different latent layers, ensure that the importance of these layers can be measured by the magnitudes of diagonals in $\mathbf{D}^*$, and thus enhance the interpretation. On the methodological side, they are a natural, convenient way to ensure the identifiability of the

factor matrices $\mathbf{D}^*$, $\mathbf{U}^*$, and $\mathbf{V}^*$ [7]. On the theoretical side, they allow us to establish rigorous error bound inequalities for the estimates $\hat{\mathbf{D}}$, $\hat{\mathbf{A}}$, and $\hat{\mathbf{B}}$. Nevertheless the orthogonality condition among the sparse latent factors may not hold exactly in certain real applications. Thus in Section V, we will investigate the robustness of our method through simulation studies where the orthogonality condition in the model is violated. It would be interesting to formally study the scenario when the orthogonality condition may hold approximately, which is beyond the scope of the current paper and we leave it for future research.

Note that for simplicity we do not explicitly state the ordering constraint $d_1 \geq \dots \geq d_m \geq 0$ in optimization problem (3). In fact, when $\rho_a(\cdot)$ and $\rho_b(\cdot)$ are matrix norms that satisfy certain invariance properties, such as the entrywise L_1 -norm and rowwise $(2, 1)$ -norm, this constraint can be easily enforced by simultaneously permuting and/or changing the signs of the singular values and the corresponding singular vectors. The orthogonality constraints are, however, essential to the optimization problem in that a solution cannot be simply obtained through solving the unconstrained regularization problem followed by an orthogonalization process. The interplay between sparse regularization and orthogonality constraints is crucial for achieving important theoretical and practical advantages, which distinguishes our SOFAR method from most previous procedures.

B. Applications of SOFAR

The SOFAR method provides a unified framework for a variety of statistical problems in multivariate analysis. We give four such examples, and in each example, briefly review existing techniques and suggest new methods.

1) *Biclustering With Sparse SVD*: The biclustering problem of a data matrix, which can be traced back to Hartigan [37], aims to simultaneously cluster the rows (samples) and columns (features) of a data matrix into statistically related subgroups. A variety of biclustering techniques, which differ in the criteria used to relate clusters of samples and clusters of features and in whether overlapping of clusters is allowed, have been suggested as useful tools in the exploratory analysis of high-dimensional genomic and text data. See, for example, Busygin *et al.* [17] for a survey. One way of formulating the biclustering problem is through the mean model

$$\mathbf{X} = \mathbf{C}^* + \mathbf{E}, \quad (4)$$

where the mean matrix $\mathbf{C}^*$ admits a sparse SVD (2) and the sparsity patterns in the left (or right) singular vectors serve as indicators for the samples (or features) to be clustered. Lee *et al.* [44] proposed to estimate the first sparse SVD layer by solving the optimization problem

$$\begin{aligned} &(\hat{\mathbf{d}}, \hat{\mathbf{u}}, \hat{\mathbf{v}}) \\ &= \arg \min_{\mathbf{d}, \mathbf{u}, \mathbf{v}} \left\{ \frac{1}{2} \|\mathbf{X} - \mathbf{d}\mathbf{u}\mathbf{v}^T\|_F^2 + \lambda_a \rho_a(\mathbf{d}\mathbf{u}) + \lambda_b \rho_b(\mathbf{d}\mathbf{v}) \right\} \\ &\text{subject to } \|\mathbf{u}\|_2 = 1, \quad \|\mathbf{v}\|_2 = 1, \end{aligned} \quad (5)$$

and obtain the next sparse SVD layer by applying the same procedure to the residual matrix $\mathbf{X} - \hat{\mathbf{d}}\hat{\mathbf{u}}\hat{\mathbf{v}}^T$. Clearly, problem (5) is a specific example of the SOFAR problem (3) with

$m = 1$ and $\lambda_d = 0$; however, the orthogonality constraints are not maintained during the layer-by-layer extraction process. The orthogonality issue also exists in most previous proposals, for example, Zhang *et al.* [68].

The multivariate linear model (1) with a sparse SVD (2) can be viewed as a supervised version of the above biclustering problem, which extends the mean model (4) to a general design matrix and can be used to identify interpretable clusters of predictors and clusters of responses that are significantly associated. Applying the SOFAR method to model (4) yields the new estimator

$$\begin{aligned} (\hat{\mathbf{D}}, \hat{\mathbf{U}}, \hat{\mathbf{V}}) &= \arg \min_{\mathbf{D}, \mathbf{U}, \mathbf{V}} \left\{ \frac{1}{2} \|\mathbf{X} - \mathbf{UDV}^T\|_F^2 \right. \\ &\quad \left. + \lambda_d \|\mathbf{D}\|_1 + \lambda_a \rho_a(\mathbf{UD}) + \lambda_b \rho_b(\mathbf{VD}) \right\} \\ &\text{subject to } \mathbf{U}^T \mathbf{U} = \mathbf{I}_m, \quad \mathbf{V}^T \mathbf{V} = \mathbf{I}_m, \end{aligned} \quad (6)$$

which estimates all sparse SVD layers simultaneously while determining the rank by nuclear norm penalization and preserving the orthogonality constraints.

2) *Sparse PCA*: A useful technique closely related to sparse SVD is sparse principal component analysis (PCA), which enhances the convergence and improves the interpretability of PCA by introducing sparsity in the loadings of principal components. There has been a fast growing literature on sparse PCA due to its importance in dimension reduction for high-dimensional data. Various formulations coupled with efficient algorithms, notably through L_0 regularization and its L_1 and semidefinite relaxations, have been proposed by Zou *et al.* [72], d'Aspremont *et al.* [25], Shen and Huang [57], Johnstone and Lu [40], and Guo *et al.* [35], among others. Recently, Benidis *et al.* [9] developed a new method to estimate sparse eigenvectors without trading off their orthogonality based on the eigenvalue decomposition rather than the SVD using the Procrustes reformulation.

We are interested in two different ways of casting sparse PCA in our sparse SVD framework. The first approach bears a resemblance to the proposal of Zou *et al.* [72], which formulates sparse PCA as a regularized multivariate regression problem with the data matrix $\mathbf{X}$ treated as both the responses and the predictors. Specifically, they proposed to solve the optimization problem

$$\begin{aligned} (\hat{\mathbf{A}}, \hat{\mathbf{V}}) &= \arg \min_{\mathbf{A}, \mathbf{V}} \left\{ \frac{1}{2} \|\mathbf{X} - \mathbf{XAV}^T\|_F^2 + \lambda_a \rho_a(\mathbf{A}) \right\} \\ &\text{subject to } \mathbf{V}^T \mathbf{V} = \mathbf{I}_m, \end{aligned} \quad (7)$$

and the loading vectors are given by the normalized columns of $\hat{\mathbf{A}}$, $\hat{\mathbf{a}}_j / \|\hat{\mathbf{a}}_j\|_2$, $j = 1, \dots, m$. However, the orthogonality of the loading vectors, a desirable property enjoyed by the standard PCA, is not enforced by problem (7). Similarly applying the SOFAR method leads to the estimator

$$\begin{aligned} (\hat{\mathbf{D}}, \hat{\mathbf{U}}, \hat{\mathbf{V}}) &= \arg \min_{\mathbf{D}, \mathbf{U}, \mathbf{V}} \left\{ \frac{1}{2} \|\mathbf{X} - \mathbf{XUDV}^T\|_F^2 \right. \\ &\quad \left. + \lambda_d \|\mathbf{D}\|_1 + \lambda_a \rho_a(\mathbf{UD}) \right\} \\ &\text{subject to } \mathbf{U}^T \mathbf{U} = \mathbf{I}_m, \quad \mathbf{V}^T \mathbf{V} = \mathbf{I}_m, \end{aligned}$$

which explicitly imposes orthogonality among the loading vectors (the columns of $\widehat{\mathbf{U}}$). One can optionally ignore the nuclear norm penalty and determine the number of principal components by some well-established criterion.

The second approach exploits the connection of sparse PCA with regularized SVD suggested by Shen and Huang [57]. They proposed to solve the rank-1 matrix approximation problem

$$\begin{aligned} (\widehat{\mathbf{u}}, \widehat{\mathbf{b}}) = \arg \min_{\mathbf{u}, \mathbf{b}} & \left\{ \frac{1}{2} \|\mathbf{X} - \mathbf{u}\mathbf{b}^T\|_F^2 + \lambda_b \rho_b(\mathbf{b}) \right\} \\ \text{subject to } & \|\mathbf{u}\|_2 = 1, \end{aligned} \quad (8)$$

and obtain the first loading vector $\widehat{\mathbf{b}}/\|\widehat{\mathbf{b}}\|_2$. Applying the SOFAR method similarly to the rank- m matrix approximation problem yields the estimator

$$\begin{aligned} (\widehat{\mathbf{D}}, \widehat{\mathbf{U}}, \widehat{\mathbf{V}}) = \arg \min_{\mathbf{D}, \mathbf{U}, \mathbf{V}} & \left\{ \frac{1}{2} \|\mathbf{X} - \mathbf{U}\mathbf{D}\mathbf{V}^T\|_F^2 \right. \\ & \left. + \lambda_d \|\mathbf{D}\|_1 + \lambda_b \rho_b(\mathbf{V}\mathbf{D}) \right\} \\ \text{subject to } & \mathbf{U}^T \mathbf{U} = \mathbf{I}_m, \quad \mathbf{V}^T \mathbf{V} = \mathbf{I}_m, \end{aligned}$$

which constitutes a multivariate generalization of problem (8), with the desirable orthogonality constraint imposed on the loading vectors (the columns of $\widehat{\mathbf{V}}$) and the optional nuclear norm penalty useful for determining the number of principal components.

3) *Sparse Factor Analysis*: Factor analysis plays an important role in dimension reduction and feature extraction for high-dimensional time series. A low-dimensional factor structure is appealing from both theoretical and practical angles, and can be conveniently incorporated into many other statistical tasks, such as forecasting with factor-augmented regression [59] and covariance matrix estimation [28]. See, for example, Bai and Ng [6] for an overview.

Let $\mathbf{x}_t \in \mathbb{R}^p$ be a vector of observed time series. Consider the factor model

$$\mathbf{x}_t = \mathbf{A}\mathbf{f}_t + \mathbf{e}_t, \quad t = 1, \dots, T, \quad (9)$$

where $\mathbf{f}_t \in \mathbb{R}^m$ is a vector of latent factors, $\mathbf{A} \in \mathbb{R}^{p \times m}$ is the factor loading matrix, and $\mathbf{e}_t$ is the idiosyncratic error. Most existing methods for high-dimensional factor models rely on classical PCA [2], [5] or maximum likelihood to estimate the factors and factor loadings [3], [4]; as a result, the estimated factors and loadings are generally nonzero. However, in order to assign economic meanings to the factors and loadings and to further mitigate the curse of dimensionality, it would be desirable to introduce sparsity in the factors and loadings. Writing model (9) in the matrix form

$$\mathbf{X} = \mathbf{F}\mathbf{A}^T + \mathbf{E}$$

with $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_T)^T$, $\mathbf{F} = (\mathbf{f}_1, \dots, \mathbf{f}_T)^T$, and $\mathbf{E} = (\mathbf{e}_1, \dots, \mathbf{e}_T)^T$ reveals its equivalence to model (4). Therefore, under the usual normalization restrictions that $\mathbf{F}^T \mathbf{F} / T = \mathbf{I}_m$ and $\mathbf{A}^T \mathbf{A}$ is diagonal, we can solve for $(\widehat{\mathbf{D}}, \widehat{\mathbf{U}}, \widehat{\mathbf{V}})$ in problem (6) and estimate the sparse factors and loadings by $\widehat{\mathbf{F}} = \sqrt{T}\widehat{\mathbf{U}}$ and $\widehat{\mathbf{A}} = \widehat{\mathbf{V}}\widehat{\mathbf{D}}/\sqrt{T}$.

4) *Sparse VAR Analysis*: Vector autoregressive (VAR) models have been widely used to analyze the joint dynamics of multivariate time series; see, for example, Stock and Watson [58]. Classical VAR analysis suffers greatly from the large number of free parameters in a VAR model, which grows quadratically with the dimensionality. Early attempts in reducing the impact of dimensionality have explored reduced rank methods such as canonical analysis and reduced rank regression [12], [62]. Regularization methods such as the Lasso have recently been adapted to VAR analysis for variable selection [8], [38], [42], [52].

We present an example in which our parsimonious model setup is most appropriate. Suppose we observe the data $(\mathbf{y}_t, \mathbf{x}_t)$, where $\mathbf{y}_t \in \mathbb{R}^q$ is a low-dimensional vector of time series whose dynamics are of primary interest, and $\mathbf{x}_t \in \mathbb{R}^p$ is a high-dimensional vector of informational time series. We assume that $\mathbf{x}_t$ are generated by the VAR equation

$$\mathbf{x}_t = \mathbf{C}^* \mathbf{x}_{t-1} + \mathbf{e}_t,$$

where $\mathbf{C}$ has a sparse SVD (2). This implies a low-dimensional latent model of the form

$$\mathbf{g}_t = \mathbf{D}^* \mathbf{f}_{t-1} + \tilde{\mathbf{e}}_t,$$

where $\mathbf{f}_t = \mathbf{U}^{*T} \mathbf{x}_t$, $\mathbf{g}_t = \mathbf{V}^{*T} \mathbf{x}_t$, and $\tilde{\mathbf{e}}_t = \mathbf{V}^{*T} \mathbf{e}_t$. Following the factor-augmented VAR (FAVAR) approach of Bernanke *et al.* [10], we augment the latent factors $\mathbf{f}_t$ and $\mathbf{g}_t$ to the dynamic equation of $\mathbf{y}_t$ and consider the joint model

$$\begin{pmatrix} \mathbf{y}_t \\ \mathbf{g}_t \end{pmatrix} = \begin{pmatrix} \mathbf{A}^T & \mathbf{B}^T \\ \mathbf{0} & \mathbf{D}^* \end{pmatrix} \begin{pmatrix} \mathbf{y}_{t-1} \\ \mathbf{f}_{t-1} \end{pmatrix} + \begin{pmatrix} \mathbf{e}_t \\ \tilde{\mathbf{e}}_t \end{pmatrix}.$$

We can estimate the parameters $\mathbf{A}$, $\mathbf{B}$, and $\mathbf{D}^*$ by a two-step method: first apply the SOFAR method to obtain estimates of $\mathbf{D}^*$ and $\mathbf{f}_t$, and then estimate $\mathbf{A}$ and $\mathbf{B}$ by a usual VAR since both $\mathbf{y}_t$ and $\mathbf{f}_t$ are of low dimensionality. Our approach differs from previous methods in that we enforce sparse factor loadings; hence, it would allow the factors to be given economic interpretations and would be useful for uncovering the structural relationships underlying the joint dynamics of $(\mathbf{y}_t, \mathbf{x}_t)$.

III. THEORETICAL PROPERTIES

We now investigate the theoretical properties of the SOFAR estimator (3) for model (1) under the sparse SVD structure (2). Our results concern nonasymptotic error bounds, where both response dimensionality q and predictor dimensionality p can diverge simultaneously with sample size n . The major theoretical challenges stem from the nonconvexity issues of our optimization problem which are prevalent in nonconvex statistical learning.

A. Technical Conditions

We begin with specifying a few assumptions that facilitate our technical analysis. To simplify the technical presentation, we focus on the scenario of $p \geq q$ and our proofs can be adapted easily to the case of $p < q$ with the only difference that the rates of convergence in Theorems 1 and 2 will be modified correspondingly. Assume that each column of $\mathbf{X}$,

$\tilde{\mathbf{x}}_j$ with $j = 1, \dots, p$, has been rescaled such that $\|\tilde{\mathbf{x}}_j\|_2^2 = n$. The SOFAR method minimizes the objective function in (3). Since the true rank r is unknown and we cannot expect that one can choose m to perfectly match r , the SOFAR estimates $\hat{\mathbf{U}}$, $\hat{\mathbf{V}}$, and $\hat{\mathbf{D}}$ are generally of different sizes than $\mathbf{U}^*$, $\mathbf{V}^*$, and $\mathbf{D}^*$, respectively. To ease the presentation, we expand the dimensions of matrices $\mathbf{U}^*$, $\mathbf{V}^*$, and $\mathbf{D}^*$ by simply adding columns and rows of zeros to the right and to the bottom of each of the matrices to make them of sizes $p \times q$, $q \times q$, and $q \times q$, respectively. We also expand the matrices $\hat{\mathbf{D}}$, $\hat{\mathbf{U}}$, and $\hat{\mathbf{V}}$ similarly to match the sizes of $\mathbf{D}^*$, $\mathbf{U}^*$, and $\mathbf{V}^*$, respectively. Define $\mathbf{A}^* = \mathbf{U}^* \mathbf{D}^*$ and $\mathbf{B}^* = \mathbf{V}^* \mathbf{D}^*$, and correspondingly $\hat{\mathbf{A}} = \hat{\mathbf{U}} \hat{\mathbf{D}}$ and $\hat{\mathbf{B}} = \hat{\mathbf{V}} \hat{\mathbf{D}}$ using the SOFAR estimates $(\hat{\mathbf{U}}, \hat{\mathbf{V}}, \hat{\mathbf{D}})$.

Definition 1 (Robust Spark). *The robust spark κ_c of the $n \times p$ design matrix $\mathbf{X}$ is defined as the smallest possible positive integer such that there exists an $n \times \kappa_c$ submatrix of $n^{-1/2} \mathbf{X}$ having a singular value less than a given positive constant c .*

Condition 1. (Parameter Space) *The true parameters $(\mathbf{C}^*, \mathbf{D}^*, \mathbf{A}^*, \mathbf{B}^*)$ lie in $\mathcal{C} \times \mathcal{D} \times \mathcal{A} \times \mathcal{B}$, where $\mathcal{C} = \{\mathbf{C} \in \mathbb{R}^{p \times q} : \|\mathbf{C}\|_0 < \kappa_{c_2}/2\}$, $\mathcal{D} = \{\mathbf{D} = \text{diag}\{d_j\} \in \mathbb{R}^{q \times q} : d_j = 0 \text{ or } |d_j| \geq \tau\}$, $\mathcal{A} = \{\mathbf{A} = (a_{ij}) \in \mathbb{R}^{p \times q} : a_{ij} = 0 \text{ or } |a_{ij}| \geq \tau\}$, and $\mathcal{B} = \{\mathbf{B} = (b_{ij}) \in \mathbb{R}^{q \times q} : b_{ij} = 0 \text{ or } |b_{ij}| \geq \tau\}$ with κ_{c_2} the robust spark of $\mathbf{X}$, $c_2 > 0$ some constant, and $\tau > 0$ asymptotically vanishing.*

Condition 2. (Constrained Eigenvalue) *It holds that $\max_{\|\mathbf{u}\|_0 < \kappa_{c_2}/2, \|\mathbf{u}\|_2=1} \|\mathbf{X}\mathbf{u}\|_2^2 \leq c_3 n$ and $\max_{1 \leq j \leq r} \|\mathbf{X}\mathbf{u}_j^*\|_2^2 \leq c_3 n$ for some constant $c_3 > 0$, where $\mathbf{u}_j^*$ is the left singular vector of $\mathbf{C}^*$ corresponding to singular value d_j^* .*

Condition 3. (Error Term) *The error term $\mathbf{E} \in \mathbb{R}^{n \times q} \sim N(\mathbf{0}, \mathbf{I}_n \otimes \Sigma)$ with the maximum eigenvalue $\alpha_{\max}$ of Σ bounded from above and diagonal entries of Σ being σ_j^2 's.*

Condition 4. (Penalty Functions) *For matrices $\mathbf{M}$ and $\mathbf{M}^*$ of the same size, the penalty functions ρ_h with $h \in \{a, b\}$ satisfy $|\rho_h(\mathbf{M}) - \rho_h(\mathbf{M}^*)| \leq \|\mathbf{M} - \mathbf{M}^*\|_1$.*

Condition 5. (Relative Spectral Gap) *The nonzero singular values of $\mathbf{C}^*$ satisfy that $d_{j-1}^{*2} - d_j^{*2} \geq \delta^{1/2} d_{j-1}^{*2}$ for $2 \leq j \leq r$ with $\delta > 0$ some constant, and r and $\sum_{j=1}^r (d_1^*/d_j^*)^2$ can diverge as $n \rightarrow \infty$.*

The concept of robust spark in Definition 1 was introduced initially in [30] and [69], where the thresholded parameter space was exploited to characterize the global optimum for regularization methods with general penalties. Similarly, the thresholded parameter space and the constrained eigenvalue condition which builds on the robust spark condition of the design matrix in Conditions 1 and 2 are essential for investigating the computable solution to the nonconvex SOFAR optimization problem in (3). By [30, Proposition 1], the robust spark κ_{c_2} can be at least of order $O\{n/(\log p)\}$ with asymptotic probability one when the rows of $\mathbf{X}$ are independently sampled from multivariate Gaussian distributions with dependency. Although Condition 3 assumes Gaussianity, our theory can in principle carry over to the case of sub-Gaussian errors, provided that the concentration inequalities for Gaussian random variables used in our proofs are replaced by those for sub-Gaussian random variables.

Condition 4 includes many kinds of penalty functions that bring about sparse estimates. Important examples include the entrywise L_1 -norm and rowwise $(2, 1)$ -norm, where the former encourages sparsity among the predictor/response effects specific to each rank-1 SVD layer, while the latter promotes predictor/response-wise sparsity regardless of the specific layer. To see why the rowwise $(2, 1)$ -norm satisfies Condition 4, observe that

$$\begin{aligned} \|\mathbf{M}\|_1 &\equiv \sum_i \sum_j |m_{ij}| = \sum_i \left(\sum_{j,k} |m_{ij}| |m_{ik}| \right)^{1/2} \\ &\geq \sum_i \left(\sum_j m_{ij}^2 \right)^{1/2} \equiv \|\mathbf{M}\|_{2,1}, \end{aligned}$$

which along with the triangle inequality entails that Condition 4 is indeed satisfied. Moreover, Condition 4 allows us to use concave penalties such as SCAD [29] and MCP [67]; see, for instance, the proof of Lemma 1 in [30].

Intuitively, Condition 5 rules out the nonidentifiable case where some nonzero singular values are tied with each other and the associated singular vectors in matrices $\mathbf{U}^*$ and $\mathbf{V}^*$ are identifiable only up to some orthogonal transformation. In particular, Condition 5 enables us to establish the key Lemma 3 in Section F of Supplementary Material, where the matrix perturbation theory can be invoked.

B. Main Results

Since the objective function of the SOFAR method (3) is nonconvex, solving this optimization problem is highly challenging. To overcome the difficulties, as mentioned in the Introduction we exploit the framework of CANO and suggest a two-step approach, where in the first step we solve the following L_1 -penalized squared loss minimization problem

$$\tilde{\mathbf{C}} = \arg \min_{\mathbf{C} \in \mathbb{R}^{p \times q}} \left\{ (2n)^{-1} \|\mathbf{Y} - \mathbf{X}\mathbf{C}\|_F^2 + \lambda_0 \|\mathbf{C}\|_1 \right\} \quad (10)$$

to construct an initial estimator $\tilde{\mathbf{C}}$ with $\lambda_0 \geq 0$ some regularization parameter. If $\tilde{\mathbf{C}} = \mathbf{0}$, then we set the final SOFAR estimator as $\hat{\mathbf{C}} = \mathbf{0}$; otherwise, in the second step we do a refined search and minimize the SOFAR objective function (3) in an asymptotically shrinking neighborhood of $\tilde{\mathbf{C}}$ to obtain the final SOFAR estimator $\hat{\mathbf{C}}$. In the case of $\tilde{\mathbf{C}} = \mathbf{0}$, our two-step procedure reduces to a one-step procedure. Since Theorem 1 below establishes that $\tilde{\mathbf{C}}$ can be close to $\mathbf{C}^*$ with asymptotic probability one, having $\tilde{\mathbf{C}} = \mathbf{0}$ is a good indicator that the true $\mathbf{C}^* = \mathbf{0}$.

Thanks to its convexity, the objective function in (10) in the first step can be solved easily and efficiently. In fact, since the objective function in (10) is separable it follows that the j th column of $\tilde{\mathbf{C}}$ can be obtained by solving the univariate response Lasso regression

$$\min_{\boldsymbol{\beta} \in \mathbb{R}^p} \left\{ (2n)^{-1} \|\mathbf{Y}\mathbf{e}_j - \mathbf{X}\boldsymbol{\beta}\|_2^2 + \lambda_0 \|\boldsymbol{\beta}\|_1 \right\},$$

where $\mathbf{e}_j$ is a q -dimensional vector with j th component 1 and all other components 0. The above univariate response Lasso regression has been studied extensively and well understood, and many efficient algorithms have been proposed

for solving it. Denote by $(\tilde{\mathbf{D}}, \tilde{\mathbf{U}}, \tilde{\mathbf{V}})$ the initial estimator of $(\mathbf{D}^*, \mathbf{U}^*, \mathbf{V}^*)$ obtained from the SVD of $\tilde{\mathbf{C}}$, and let $\tilde{\mathbf{A}} = \tilde{\mathbf{U}}\tilde{\mathbf{D}}$ and $\tilde{\mathbf{B}} = \tilde{\mathbf{V}}\tilde{\mathbf{D}}$. Since the bounds for the SVD are key to the analysis of SOFAR estimator in the second step, for completeness we present the nonasymptotic bounds on estimation errors of the initial estimator in the following theorem.

Theorem 1 (Error Bounds for Initial Estimator). *Assume that Conditions 1–3 hold and let $\lambda_0 = c_0 \sigma_{\max} (n^{-1} \log(pq))^{1/2}$ with $\sigma_{\max} = \max_{1 \leq j \leq q} \sigma_j$ and $c_0 > \sqrt{2}$ some constant. Then with probability at least $1 - 2(pq)^{1-c_0^2/2}$, the estimation error is bounded as*

$$\|\tilde{\mathbf{C}} - \mathbf{C}^*\|_F \leq R_n \equiv c(n^{-1}s \log(pq))^{1/2} \quad (11)$$

with $s = \|\mathbf{C}^*\|_0$ and $c > 0$ some constant. Under additional Condition 5, with the same probability bound the following estimation error bounds hold simultaneously

$$\|\tilde{\mathbf{D}} - \mathbf{D}^*\|_F \leq c(n^{-1}s \log(pq))^{1/2}, \quad (12)$$

$$\|\tilde{\mathbf{A}} - \mathbf{A}^*\|_F + \|\tilde{\mathbf{B}} - \mathbf{B}^*\|_F \leq c\eta_n(n^{-1}s \log(pq))^{1/2}, \quad (13)$$

where $\eta_n = 1 + \delta^{-1/2}(\sum_{j=1}^r (d_1^*/d_j^*)^2)^{1/2}$.

For the case of $q = 1$, the estimation error bound (11) is consistent with the well-known oracle inequality for Lasso [11]. The additional estimation error bounds (12) and (13) for the SVD in Theorem 1 are, however, new to the literature. It is worth mentioning that Condition 5 and the latest results in [65] play a crucial role in establishing these additional error bounds.

After obtaining the initial estimator $\tilde{\mathbf{C}}$ from the first step, we can solve the SOFAR optimization problem in an asymptotically shrinking neighborhood of $\tilde{\mathbf{C}}$. More specifically, we define $\tilde{\mathcal{P}}_n = \{\mathbf{C} : \|\mathbf{C} - \tilde{\mathbf{C}}\|_F \leq 2R_n\}$ with R_n the upper bound in (11). Then it is seen from Theorem 1 that the true coefficient matrix $\mathbf{C}^*$ is contained in $\tilde{\mathcal{P}}_n$ with probability at least $1 - 2(pq)^{1-c_0^2/2}$. Further define

$$\mathcal{P}_n = \tilde{\mathcal{P}}_n \cap (\mathcal{C} \times \mathcal{D} \times \mathcal{A} \times \mathcal{B}), \quad (14)$$

where sets $\mathcal{C}$, $\mathcal{D}$, $\mathcal{A}$, and $\mathcal{B}$ are defined in Condition 1. Then with probability at least $1 - 2(pq)^{1-c_0^2/2}$, the set $\mathcal{P}_n$ defined in (14) is nonempty with at least one element $\mathbf{C}^*$ by Condition 1. We minimize the SOFAR objective function (3) by searching in the shrinking neighborhood $\mathcal{P}_n$ and denote by $\hat{\mathbf{C}}$ the resulting SOFAR estimator. Then it follows that with probability at least $1 - 2(pq)^{1-c_0^2/2}$,

$$\|\hat{\mathbf{C}} - \mathbf{C}^*\|_F \leq \|\hat{\mathbf{C}} - \tilde{\mathbf{C}}\|_F + \|\tilde{\mathbf{C}} - \mathbf{C}^*\|_F \leq 3R_n,$$

where the first inequality is by the triangle inequality and the second one is by the construction of set $\mathcal{P}_n$ and Theorem 1. Therefore, we see that the SOFAR estimator given by our two-step procedure is guaranteed to have convergence rate at least $O(R_n)$.

Since the initial estimator investigated in Theorem 1 completely ignores the finer sparse SVD structure of the coefficient matrix $\mathbf{C}^*$, intuitively the second step of SOFAR estimation can lead to improved error bounds. Indeed we show in Theorem 2 below that with the second step of refinement,

up to some columnwise sign changes the SOFAR estimator can admit estimator error bounds in terms of parameters r , s_a , and s_b with $r = \|\mathbf{D}^*\|_0$, $s_a = \|\mathbf{A}^*\|_0$, and $s_b = \|\mathbf{B}^*\|_0$. When r , s_a , and s_b are drastically smaller than s , these new upper bounds can have better rates of convergence.

Theorem 2 (Error Bounds for SOFAR Estimator). *Assume that Conditions 1–5 hold, $\lambda_{\max} \equiv \max(\lambda_d, \lambda_a, \lambda_b) = c_1 (n^{-1} \log(pr))^{1/2}$ with $c_1 > 0$ some large constant, $\log p = O(n^\alpha)$, $q = O(n^{\beta/2})$, $s = O(n^\gamma)$, and $\eta_n^2 = o(\min\{\lambda_{\max}^{-1} \tau, n^{1-\alpha-\beta-\gamma} \tau^2\})$ with $\alpha, \beta, \gamma \geq 0$, $\alpha + \beta + \gamma < 1$, and η_n as given in Theorem 1. Then with probability at least*

$$1 - \left\{ 2(pq)^{1-c_0^2/2} + 2(pr)^{-\tilde{c}_2} + 2pr \exp\left(-\tilde{c}_3 n^{1-\beta-\gamma} \tau^2 \eta_n^{-2}\right) \right\}, \quad (15)$$

the SOFAR estimator satisfies the following error bounds simultaneously:

$$(a) \|\hat{\mathbf{C}} - \mathbf{C}^*\|_F \leq c \min\{s, (r + s_a + s_b)\eta_n^2\}^{1/2} \{n^{-1} \log(pq)\}^{1/2}, \quad (16)$$

$$(b) \|\hat{\mathbf{D}} - \mathbf{D}^*\|_F + \|\hat{\mathbf{A}} - \mathbf{A}^*\|_F + \|\hat{\mathbf{B}} - \mathbf{B}^*\|_F \leq c \min\{s, (r + s_a + s_b)\eta_n^2\}^{1/2} \eta_n \{n^{-1} \log(pq)\}^{1/2}, \quad (17)$$

$$(c) \|\hat{\mathbf{D}} - \mathbf{D}^*\|_0 + \|\hat{\mathbf{A}} - \mathbf{A}^*\|_0 + \|\hat{\mathbf{B}} - \mathbf{B}^*\|_0 \leq (r + s_a + s_b)[1 + o(1)], \quad (18)$$

$$(d) \|\hat{\mathbf{D}} - \mathbf{D}^*\|_1 + \|\hat{\mathbf{A}} - \mathbf{A}^*\|_1 + \|\hat{\mathbf{B}} - \mathbf{B}^*\|_1 \leq c(r + s_a + s_b)\eta_n^2 \lambda_{\max}, \quad (19)$$

$$(e) n^{-1} \|\mathbf{X}(\hat{\mathbf{C}} - \mathbf{C}^*)\|_F^2 \leq c(r + s_a + s_b)\eta_n^2 \lambda_{\max}^2, \quad (20)$$

where $c_0 > \sqrt{2}$ and $c, \tilde{c}_2, \tilde{c}_3$ are some positive constants.

We see from Theorem 2 that the upper bounds in (16) and (17) are the minimum of two rates, one involving $r + s_a + s_b$ (the total sparsity of $\mathbf{D}^*$, $\mathbf{A}^*$, and $\mathbf{B}^*$) and the other one involving s (the sparsity of matrix $\mathbf{C}^*$). The rate involving s is from the first step of Lasso estimation, while the rate involving $r + s_a + s_b$ is from the second step of SOFAR refinement. For the case of $s > (r + s_a + s_b)\eta_n^2$, our two-step procedure leads to enhanced error rates under the Frobenius norm. Moreover, the error rates in (18)–(20) are new to the literature and not shared by the initial Lasso estimator, showing again the advantages of having the second step of refinement. It is seen that our two-step SOFAR estimator is capable of recovering the sparsity structure of $\mathbf{D}^*$, $\mathbf{A}^*$, and $\mathbf{B}^*$ very well.

Let us gain more insights into these new error bounds. In the case of univariate response with $q = 1$, we have $\eta_n = 1 + \delta$, $r = 1$, $s_a = s$, and $s_b = 1$. Then the upper bounds in (16)–(20) reduce to $c\{sn^{-1} \log p\}^{1/2}$, $c\{sn^{-1} \log p\}^{1/2}$, cs , $cs\{n^{-1} \log p\}^{1/2}$, and $cn^{-1}s \log p$, respectively, which are indeed within a logarithmic factor of the oracle rates for the case of high-dimensional univariate response regression. Furthermore, in the rank-one case of $r = 1$ we have $\eta_n = 1 + \delta^{-1/2}$ and $s = s_a s_b$. Correspondingly, the upper bounds in (11)–(13) for the initial Lasso estimator all become $c\{n^{-1}s_a s_b \log(pq)\}^{1/2}$, while the upper

bounds in (16)–(20) for the SOFAR estimator become $c\{(s_a + s_b)n^{-1} \log(pq)\}^{1/2}$, $c\{(s_a + s_b)n^{-1} \log(pq)\}^{1/2}$, $c(s_a + s_b)$, $c(s_a + s_b)\{n^{-1} \log(pq)\}^{1/2}$, and $cn^{-1}(s_a + s_b) \log(pq)$, respectively. In particular, we see that the SOFAR estimator can have much improved rates of convergence even in the setting of $r = 1$.

IV. IMPLEMENTATION OF SOFAR

The interplay between sparse regularization and orthogonality constraints creates substantial algorithmic challenges for solving the SOFAR optimization problem (3), for which many existing algorithms can become either inefficient or inapplicable. For example, coordinate descent methods that are popular for solving large-scale sparse regularization problems [32] are not directly applicable because the penalty terms in problem (3) are not separable under the orthogonality constraints. Also, the general framework for algorithms involving orthogonality constraints [26] does not take sparsity into account and hence does not lead to efficient algorithms in our context. Recently, Benidis *et al.* [9] focused on the unsupervised learning setting and introduced a new algorithm for estimating sparse eigenvectors without trading off their orthogonality based on the eigenvalue decomposition rather than the SVD. To obtain sparse orthogonal eigenvectors, they applied the minorization-maximization framework on the sparse PCA problem, which results in solving a sequence of rectangular Procrustes problems. Inspired by a recently revived interest in the augmented Lagrangian method (ALM) and its variants for large-scale optimization in statistics and machine learning [13], in this section we develop an efficient algorithm for solving problem (3).

A. SOFAR Algorithm With ALM-BCD

The architecture of the proposed SOFAR algorithm is based on the ALM coupled with block coordinate descent (BCD). The first construction step is to utilize variable splitting to separate the orthogonality constraints and sparsity-inducing penalties into different subproblems, which then enables efficient optimization in a block coordinate descent fashion. To this end, we introduce two new variables $\mathbf{A}$ and $\mathbf{B}$, and express problem (3) in the equivalent form

$$\begin{aligned} (\hat{\Theta}, \hat{\Omega}) = \arg \min_{\Theta, \Omega} & \left\{ \frac{1}{2} \|\mathbf{Y} - \mathbf{XUDV}^T\|_F^2 \right. \\ & \left. + \lambda_d \|\mathbf{D}\|_1 + \lambda_a \rho_a(\mathbf{A}) + \lambda_b \rho_b(\mathbf{B}) \right\} \\ \text{subject to } & \mathbf{U}^T \mathbf{U} = \mathbf{I}_m, \quad \mathbf{V}^T \mathbf{V} = \mathbf{I}_m, \\ & \mathbf{UD} = \mathbf{A}, \quad \mathbf{VD} = \mathbf{B}, \end{aligned} \quad (21)$$

where $\Theta = (\mathbf{D}, \mathbf{U}, \mathbf{V})$ and $\Omega = (\mathbf{A}, \mathbf{B})$. We form the augmented Lagrangian for problem (21) as

$$\begin{aligned} L_\mu(\Theta, \Omega, \Gamma) = & \frac{1}{2} \|\mathbf{Y} - \mathbf{XUDV}^T\|_F^2 \\ & + \lambda_d \|\mathbf{D}\|_1 + \lambda_a \rho_a(\mathbf{A}) + \lambda_b \rho_b(\mathbf{B}) + \langle \Gamma_a, \mathbf{UD} - \mathbf{A} \rangle \\ & + \langle \Gamma_b, \mathbf{VD} - \mathbf{B} \rangle + \frac{\mu}{2} \|\mathbf{UD} - \mathbf{A}\|_F^2 + \frac{\mu}{2} \|\mathbf{VD} - \mathbf{B}\|_F^2, \end{aligned}$$

TABLE I
SOFAR ALGORITHM WITH ALM-BCD

Parameters: $\lambda_d, \lambda_a, \lambda_b$, and $\gamma > 1$
Initialize $\mathbf{U}^0, \mathbf{V}^0, \mathbf{D}^0, \mathbf{A}^0, \mathbf{B}^0, \Gamma_a^0, \Gamma_b^0$, and μ^0
For $k = 0, 1, \dots$ do
 update $\mathbf{U}, \mathbf{V}, \mathbf{D}, \mathbf{A}$, and $\mathbf{B}$:
 (a) $\mathbf{U}^{k+1} \leftarrow \arg \min_{\mathbf{U}^T \mathbf{U} = \mathbf{I}_m} \left\{ \frac{1}{2} \|\mathbf{Y} - \mathbf{XUD}^k(\mathbf{V}^k)^T\|_F^2 + \frac{\mu^k}{2} \|\mathbf{UD}^k - \mathbf{A}^k + \Gamma_a^k / \mu^k\|_F^2 \right\}$
 (b) $\mathbf{V}^{k+1} \leftarrow \arg \min_{\mathbf{V}^T \mathbf{V} = \mathbf{I}_m} \left\{ \frac{1}{2} \|\mathbf{Y} - \mathbf{XU}^{k+1}\mathbf{D}^k\mathbf{V}^T\|_F^2 + \frac{\mu^k}{2} \|\mathbf{VD}^k - \mathbf{B}^k + \Gamma_b^k / \mu^k\|_F^2 \right\}$
 (c) $\mathbf{D}^{k+1} \leftarrow \arg \min_{\mathbf{D} \geq 0} \left\{ \frac{1}{2} \|\mathbf{Y} - \mathbf{XU}^{k+1}\mathbf{D}(\mathbf{V}^{k+1})^T\|_F^2 + \frac{\mu^k}{2} \|\mathbf{U}^{k+1}\mathbf{D} - \mathbf{A}^k + \Gamma_a^k / \mu^k\|_F^2 + \frac{\mu^k}{2} \|\mathbf{V}^{k+1}\mathbf{D} - \mathbf{B}^k + \Gamma_b^k / \mu^k\|_F^2 + \lambda_d \|\mathbf{D}\|_1 \right\}$
 (d) $\mathbf{A}^{k+1} \leftarrow \arg \min_{\mathbf{A}} \left\{ \frac{\mu^k}{2} \|\mathbf{U}^{k+1}\mathbf{D}^{k+1} - \mathbf{A}\|_F^2 + \Gamma_a^k / \mu^k\|_F^2 + \lambda_a \rho_a(\mathbf{A}) \right\}$
 (e) $\mathbf{B}^{k+1} \leftarrow \arg \min_{\mathbf{B}} \text{big} \left\{ \frac{\mu^k}{2} \|\mathbf{V}^{k+1}\mathbf{D}^{k+1} - \mathbf{B}\|_F^2 + \Gamma_b^k / \mu^k\|_F^2 + \lambda_b \rho_b(\mathbf{B}) \right\}$
 (f) optionally, repeat (a)–(e) until convergence
 update Γ_a and Γ_b :
 (a) $\Gamma_a^{k+1} \leftarrow \Gamma_a^k + \mu^k (\mathbf{U}^{k+1}\mathbf{D}^{k+1} - \mathbf{A}^{k+1})$
 (b) $\Gamma_b^{k+1} \leftarrow \Gamma_b^k + \mu^k (\mathbf{V}^{k+1}\mathbf{D}^{k+1} - \mathbf{B}^{k+1})$
 update μ by $\mu^{k+1} \leftarrow \gamma \mu^k$
end

where $\Gamma = (\Gamma_a, \Gamma_b)$ is the set of Lagrangian multipliers and $\mu > 0$ is a penalty parameter. Based on ALM, the proposed algorithm consists of the following iterations:

- 1) (Θ, Ω) -step: $(\Theta^{k+1}, \Omega^{k+1}) \leftarrow \arg \min_{\Theta, \Omega} \text{big} \left\{ \frac{1}{2} \|\mathbf{Y} - \mathbf{XUDV}^T\|_F^2 + \lambda_d \|\mathbf{D}\|_1 + \lambda_a \rho_a(\mathbf{A}) + \lambda_b \rho_b(\mathbf{B}) \right\}$
- 2) Γ -step: $\Gamma_a^{k+1} \leftarrow \Gamma_a^k + \mu (\mathbf{U}^{k+1}\mathbf{D}^{k+1} - \mathbf{A}^{k+1})$ and $\Gamma_b^{k+1} \leftarrow \Gamma_b^k + \mu (\mathbf{V}^{k+1}\mathbf{D}^{k+1} - \mathbf{B}^{k+1})$.

The (Θ, Ω) -step can be solved by a block coordinate descent method [61] cycling through the blocks $\mathbf{U}, \mathbf{V}, \mathbf{D}, \mathbf{A}$, and $\mathbf{B}$. Note that the orthogonality constraints and the sparsity-inducing penalties are now separated into subproblems with respect to Θ and Ω , respectively. To achieve convergence of the SOFAR algorithm in practice, an inexact minimization with a few block coordinate descent iterations is often sufficient. Moreover, to enhance the convergence of the algorithm to a feasible solution we optionally increase the penalty parameter μ by a ratio $\gamma > 1$ at the end of each iteration. This leads to the SOFAR algorithm with ALM-BCD described in Table I.

We still need to solve the subproblems in algorithm I. The $\mathbf{U}$ -update is similar to the weighted orthogonal Procrustes problem considered by Koschat and Swayne [43]. By expanding the squares and omitting terms not involving $\mathbf{U}$, this subproblem is equivalent to minimizing

$$\begin{aligned} & \frac{1}{2} \|\mathbf{XUD}^k\|_F^2 \\ & - \text{tr}(\mathbf{U}^T \mathbf{X}^T \mathbf{Y} \mathbf{V}^k \mathbf{D}^k) - \text{tr}(\mathbf{U}^T (\mu^k \mathbf{A}^k - \Gamma_a^k) \mathbf{D}^k) \end{aligned}$$

subject to $\mathbf{U}^T \mathbf{U} = \mathbf{I}_m$. Taking a matrix $\mathbf{Z}$ such that $\mathbf{Z}^T \mathbf{Z} = \rho^2 \mathbf{I}_p - \mathbf{X}^T \mathbf{X}$, where ρ^2 is the largest eigenvalue of $\mathbf{X}^T \mathbf{X}$,

we can follow the argument of Koschat and Swayne [43] to obtain the iterative algorithm: for $j = 0, 1, \dots$, form the $p \times m$ matrix $\mathbf{C}_1 = (\mathbf{X}^T \mathbf{Y} \mathbf{V}^k + \mu^k \mathbf{A}^k - \mathbf{\Gamma}_a^k + \mathbf{Z}^T \mathbf{Z} \mathbf{U}^j \mathbf{D}^k) \mathbf{D}^k$, compute the SVD $\mathbf{U}_1 \mathbf{\Sigma}_1 \mathbf{V}_1^T = \mathbf{C}_1$, and update $\mathbf{U}^{j+1} = \mathbf{U}_1 \mathbf{V}_1^T$. Note that $\mathbf{C}_1$ depends on $\mathbf{Z}^T \mathbf{Z}$ only, and hence the explicit computation of $\mathbf{Z}$ is not needed. The $\mathbf{V}$ -update is similar to a standard orthogonal Procrustes problem and amounts to maximizing

$$\text{tr}(\mathbf{V}^T \mathbf{Y}^T \mathbf{X} \mathbf{U}^{k+1} \mathbf{D}^k) + \text{tr}(\mathbf{V}^T (\mu^k \mathbf{B}^k - \mathbf{\Gamma}_b^k) \mathbf{D}^k)$$

subject to $\mathbf{V}^T \mathbf{V} = \mathbf{I}_m$. A direct method for this problem [34, pp. 327–328] gives the algorithm: form the $q \times m$ matrix $\mathbf{C}_2 = (\mathbf{Y}^T \mathbf{X} \mathbf{U}^{k+1} + \mu^k \mathbf{B}^k - \mathbf{\Gamma}_b^k) \mathbf{D}^k$, compute the SVD $\mathbf{U}_2 \mathbf{\Sigma}_2 \mathbf{V}_2^T = \mathbf{C}_2$, and set $\mathbf{V} = \mathbf{U}_2 \mathbf{V}_2^T$. Since m is usually small, the SVD computations in the $\mathbf{U}$ - and $\mathbf{V}$ -updates are cheap. The Lasso problem in the $\mathbf{D}$ -update reduces to a standard quadratic program with the nonnegativity constraint, which can be readily solved by efficient algorithms; see, for example, Sha *et al.* [56]. Note that the $\mathbf{D}$ -update may set some singular values to exactly zero; hence, a greedy strategy can be taken to further bring down the computational complexity, by removing the zero singular values and reducing the sizes of the relevant matrices accordingly in subsequent computations. The updates of $\mathbf{A}$ and $\mathbf{B}$ are free of orthogonality constraints and therefore easy to solve. With the popular choices of $\|\cdot\|_1$ and $\|\cdot\|_{2,1}$ as the penalty functions, the updates can be performed by entrywise and rowwise soft-thresholding, respectively.

Following the theoretical analysis for the SOFAR method in Section III, we employ the SVD of the cross-validated L_1 -penalized estimator $\tilde{\mathbf{C}}$ in (10) to initialize $\mathbf{U}$, $\mathbf{V}$, $\mathbf{D}$, $\mathbf{A}$, and $\mathbf{B}$; the $\mathbf{\Gamma}_a$ and $\mathbf{\Gamma}_b$ are initialized as zero matrices. In practice, for large-scale problems we can further scale up the SOFAR method by performing feature screening with the initial estimator $\tilde{\mathbf{C}}$, that is, the response variables corresponding to zero columns in $\tilde{\mathbf{C}}$ and the predictors corresponding to zero rows in $\tilde{\mathbf{C}}$ could be removed prior to the finer SOFAR analysis.

B. Convergence Analysis and Tuning Parameter Selection

For general nonconvex problems, an ALM algorithm needs not to converge, and even if it converges, it needs not to converge to an optimal solution. We have the following convergence results regarding the proposed SOFAR algorithm with ALM-BCD.

Theorem 3 (Convergence of SOFAR Algorithm). *Assume that $\sum_{k=1}^{\infty} \{[\Delta L_{\mu}(\mathbf{U}^k)]^{1/2} + [\Delta L_{\mu}(\mathbf{V}^k)]^{1/2} + [\Delta L_{\mu}(\mathbf{D}^k)]^{1/2}\} < \infty$ and the penalty functions $\rho_a(\cdot)$ and $\rho_b(\cdot)$ are convex, where $\Delta L_{\mu}(\cdot)$ denotes the decrease in $L_{\mu}(\cdot)$ by a block update. Then the sequence generated by the SOFAR algorithm converges to a local solution of the augmented Lagrangian for problem (21).*

Note that without the above assumption on $(\mathbf{U}^k)$, $(\mathbf{V}^k)$, and $(\mathbf{D}^k)$, we can only show that the differences between two consecutive $\mathbf{U}$ -, $\mathbf{V}$ -, and $\mathbf{D}$ -updates converge to zero by the convergence of the sequence $(L_{\mu}(\cdot))$, but the sequences $(\mathbf{U}^k)$, $(\mathbf{V}^k)$, and $(\mathbf{D}^k)$ may not necessarily converge. Although Theorem 3 does not ensure the convergence of algorithm I to an optimal solution, numerical evidence suggests that the

algorithm has strong convergence properties and the produced solutions perform well in numerical studies.

The above SOFAR algorithm is presented for a fixed triple of tuning parameters $(\lambda_d, \lambda_a, \lambda_b)$. One may apply a fine grid search with K -fold cross-validation or an information criterion such as BIC and its high-dimensional extensions including GIC [31] to choose an optimal triple of tuning parameters and hence a best model. In either case, a full search over a three-dimensional grid would be prohibitively expensive, especially for large-scale problems. Theorem 2, however, suggests that the parameter tuning can be effectively reduced to one or two dimensions. Hence, we adopt a search strategy which is computationally affordable and still provides reasonable and robust performance. To this end, we first estimate an upper bound on each of the tuning parameters by considering the *marginal null model*, where two of the three tuning parameters are fixed at zero and the other is set to the minimum value leading to a null model. We denote the upper bounds thus obtained by $(\lambda_d^*, \lambda_a^*, \lambda_b^*)$, and conduct a search over a one-dimensional grid of values between $(\lambda_d^*, \lambda_a^*, \lambda_b^*)$ and $(\varepsilon \lambda_d^*, \varepsilon \lambda_a^*, \varepsilon \lambda_b^*)$, with $\varepsilon > 0$ sufficiently small (e.g., 10^{-3}) to ensure the coverage of a full spectrum of reasonable solutions. Our numerical experience suggests that this simple search strategy works well in practice while reducing the computational cost dramatically. More flexibility can be gained by adjusting the ratios between λ_d , λ_a , and λ_b if additional information about the relative sparsity levels of $\mathbf{D}$, $\mathbf{A}$, and $\mathbf{B}$ is available.

V. NUMERICAL STUDIES

A. Simulation Examples

Our Condition 4 in Section III-A accommodates a large group of penalty functions including concave ones such as SCAD and MCP. As demonstrated in [27] and [73], nonconvex regularization problems can be solved using the idea of local linear approximation, which essentially reduces the original problem to the weighted L_1 -regularization with the weights chosen adaptively based on some initial solution. For this reason, in the simulation study we focus on the entrywise L_1 -norm $\|\cdot\|_1$ and the rowwise $(2, 1)$ -norm $\|\cdot\|_{2,1}$, as well as their adaptive extensions. The use of adaptively weighted penalties has also been explored in the contexts of reduced rank regression [21] and sparse PCA [45]. We next provide more details on the adaptive penalties used in our simulation study. To simplify the presentation, we use the entrywise L_1 -norm as an example.

Incorporating adaptive weighting into the penalty terms in problem (21) leads to the adaptive SOFAR estimator

$$\begin{aligned} (\hat{\mathbf{\Theta}}, \hat{\mathbf{\Omega}}) = \arg \min_{\mathbf{\Theta}, \mathbf{\Omega}} & \left\{ \frac{1}{2} \|\mathbf{Y} - \mathbf{X} \mathbf{U} \mathbf{D} \mathbf{V}^T\|_F^2 \right. \\ & \left. + \lambda_d \|\mathbf{W}_d \circ \mathbf{D}\|_1 + \lambda_a \|\mathbf{W}_a \circ \mathbf{A}\|_1 + \lambda_b \|\mathbf{W}_b \circ \mathbf{B}\|_1 \right\} \\ \text{subject to } & \mathbf{U}^T \mathbf{U} = \mathbf{I}_m, \quad \mathbf{V}^T \mathbf{V} = \mathbf{I}_m, \\ & \mathbf{U} \mathbf{D} = \mathbf{A}, \quad \mathbf{V} \mathbf{D} = \mathbf{B}, \end{aligned}$$

where $\mathbf{W}_d \in \mathbb{R}^{m \times m}$, $\mathbf{W}_a \in \mathbb{R}^{p \times m}$, and $\mathbf{W}_b \in \mathbb{R}^{q \times m}$ are weighting matrices that depend on the initial estimates $\tilde{\mathbf{D}}$, $\tilde{\mathbf{A}}$, and $\tilde{\mathbf{B}}$, respectively, and $\circ$ is the Hadamard or entrywise product. The weighting matrices are chosen to reflect

the intuition that singular values and singular vectors of larger magnitude should be less penalized in order to reduce bias and improve efficiency in estimation. As suggested in [73], if one is interested in using some nonconvex penalty functions $\rho_a(\cdot)$ and $\rho_b(\cdot)$ then the weight matrices can be constructed by using the first order derivatives of the penalty functions and the initial solution $(\tilde{\mathbf{A}}, \tilde{\mathbf{B}}, \tilde{\mathbf{D}})$. In our implementation, for simplification we adopt the alternative popular choice of $\mathbf{W}_d = \text{diag}(\tilde{d}_1^{-1}, \dots, \tilde{d}_m^{-1})$ with $\tilde{d}_j$ the j th diagonal entry of $\tilde{\mathbf{D}}$, as suggested in Zou [71]. Similarly, we set $\mathbf{W}_a = (\tilde{a}_{ij}^{-1})$ and $\mathbf{W}_b = (\tilde{b}_{ij}^{-1})$ with $\tilde{a}_{ij}$ and $\tilde{b}_{ij}$ the (i, j) th entries of $\tilde{\mathbf{A}}$ and $\tilde{\mathbf{B}}$, respectively. Extension of the SOFAR algorithm with ALM-BCD in Section IV-A is also straightforward, with the $\mathbf{D}$ -update becoming an adaptive Lasso problem and the updates of $\mathbf{A}$ and $\mathbf{B}$ now performed by adaptive soft-thresholding. A further way of improving the estimation efficiency is to exploit regularization methods in the thresholded parameter space [30] or thresholded regression [69], which we do not pursue in this paper.

We compare the SOFAR estimator with the entrywise L_1 -norm (Lasso) penalty (SOFAR-L) or the rowwise $(2, 1)$ -norm (group Lasso) penalty (SOFAR-GL) with five alternative methods, including three classical methods, namely, the ordinary least squares (OLS), separate adaptive Lasso regressions (Lasso), and reduced rank regression (RRR), and two recent sparse and low rank methods, namely, reduced rank regression with sparse SVD (RSSVD) proposed by Chen *et al.* [19] and sparse reduced rank regression (SRRR) considered by Chen and Huang [22] (see also the rank constrained group Lasso estimator in Bunea *et al.* 16). Both Chen *et al.* [19] and Chen and Huang [22] used adaptive weighted penalization. We thus consider both nonadaptive and adaptive versions of the SOFAR-L, SOFAR-GL, RSSVD, and SRRR methods.

1) *Simulation Setups*: We consider several simulation settings with various model dimensions and sparse SVD patterns in the coefficient matrix $\mathbf{C}^*$. In all settings, we took the sample size $n = 200$ and the true rank $r = 3$. Models 1 and 2 concern the entrywise sparse SVD structure in $\mathbf{C}^*$. The design matrix $\mathbf{X}$ was generated with i.i.d. rows from $N_p(\mathbf{0}, \Sigma_x)$, where $\Sigma_x = (0.5^{|i-j|})$. In model 1, we set $p = 100$ and $q = 40$, and let $\mathbf{C}^* = \sum_{j=1}^3 d_j^* \mathbf{u}_j^* \mathbf{v}_j^{*T}$ with $d_1^* = 20$, $d_2^* = 15$, $d_3^* = 10$, and

$$\begin{aligned} \tilde{\mathbf{u}}_1 &= (\text{unif}(S_u, 5), \text{rep}(0, 20))^T, \\ \tilde{\mathbf{u}}_2 &= (\text{rep}(0, 3), -\tilde{u}_{1,4}, \tilde{u}_{1,5}, \text{unif}(S_u, 3), \text{rep}(0, 17))^T, \\ \tilde{\mathbf{u}}_3 &= (\text{rep}(0, 8), \text{unif}(S_u, 2), \text{rep}(0, 15))^T, \\ \mathbf{u}_j^* &= \tilde{\mathbf{u}}_j / \|\tilde{\mathbf{u}}_j\|_2, \quad j = 1, 2, 3, \\ \tilde{\mathbf{v}}_1 &= (\text{unif}(S_v, 5), \text{rep}(0, 10))^T, \\ \tilde{\mathbf{v}}_2 &= (\text{rep}(0, 5), \text{unif}(S_v, 5), \text{rep}(0, 5))^T, \\ \tilde{\mathbf{v}}_3 &= (\text{rep}(0, 10), \text{unif}(S_v, 5))^T, \\ \mathbf{v}_j^* &= \tilde{\mathbf{v}}_j / \|\tilde{\mathbf{v}}_j\|_2, \quad j = 1, 2, 3, \end{aligned}$$

where $\text{unif}(S, k)$ denotes a k -vector with i.i.d. entries from the uniform distribution on the set S , $S_u = \{-1, 1\}$, $S_v = [-1, -0.5] \cup [0.5, 1]$, $\text{rep}(\alpha, k)$ denotes a k -vector replicating

the value α , and $\tilde{u}_{j,k}$ is the k th entry of $\tilde{\mathbf{u}}_j$. Model 2 is similar to Model 1 except with higher model dimensions, where we set $p = 400$, $q = 120$, and appended 300 and 80 zeros to each $\mathbf{u}_j^*$ and $\mathbf{v}_j^*$ defined above, respectively.

Models 3 and 4 pertain to the rowwise/columnwise sparse SVD structure in $\mathbf{C}^*$. Also, we intend to study the case of approximate low-rankness/sparsity, by not requiring the signals be bounded away from zero. We generated $\mathbf{X}$ with i.i.d. rows from $N_p(\mathbf{0}, \Sigma_x)$, where Σ_x has diagonal entries 1 and off-diagonal entries 0.5. The rowwise sparsity patterns were generated in a similar way to the setup in Chen and Huang [22] except that we allow also the matrix of right singular vectors to be rowwise sparse, so that response selection may also be necessary. Specifically, we let $\mathbf{C}^* = \mathbf{C}_1 \mathbf{C}_2^T$, where $\mathbf{C}_1 \in \mathbb{R}^{p \times r}$ with i.i.d. entries in its first p_0 rows from $N(0, 1)$ and the rest set to zero, and $\mathbf{C}_2 \in \mathbb{R}^{q \times r}$ with i.i.d. entries in its first q_0 rows from $N(0, 1)$ and the rest set to zero. We set $p = 100$, $p_0 = 10$, $q = q_0 = 10$ in Model 3, and $p = 400$, $p_0 = 10$, $q = 200$, and $q_0 = 10$ in Model 4. We also investigate models with even higher dimensions. In Model 5, we experimented with increasing the dimensions of Model 2 to $p = 1000$ and $q = 400$, by adding more noise variables, i.e., appending zeros to the $\mathbf{u}_j^*$ and $\mathbf{v}_j^*$ vectors.

Finally, we consider Model 6 where the orthogonality among the sparse factors is violated. Specifically, Model 6 is similar to Model 1, except that we modify the true values of $\mathbf{U}^*$ and $\mathbf{V}^*$ as follows,

$$\begin{aligned} \tilde{\mathbf{u}}_1 &= (\text{unif}(S_u, 5), \text{rep}(0, 20))^T, \\ \tilde{\mathbf{u}}_2 &= (\text{rep}(0, 3), \text{unif}(S_u, 5), \text{rep}(0, 17))^T, \\ \tilde{\mathbf{u}}_3 &= (\text{rep}(0, 8), \text{unif}(S_u, 2), \text{rep}(0, 15))^T, \\ \mathbf{u}_j^* &= \tilde{\mathbf{u}}_j / \|\tilde{\mathbf{u}}_j\|_2, \quad j = 1, 2, 3, \\ \tilde{\mathbf{v}}_1 &= (\text{unif}(S_v, 5), \text{rep}(0, 10))^T, \\ \tilde{\mathbf{v}}_2 &= (\text{rep}(0, 4), \text{unif}(S_v, 5), \text{rep}(0, 6))^T, \\ \tilde{\mathbf{v}}_3 &= (\text{rep}(0, 8), \text{unif}(S_v, 5), \text{rep}(0, 2))^T, \\ \mathbf{v}_j^* &= \tilde{\mathbf{v}}_j / \|\tilde{\mathbf{v}}_j\|_2, \quad j = 1, 2, 3. \end{aligned}$$

Model 7 is similar to Model 6 except with higher model dimensionality, where we set $p = 400$, $q = 120$, and appended 300 and 80 zeros to each $\mathbf{u}_j^*$ and $\mathbf{v}_j^*$ defined above, respectively. We would like to point out that when the sparse factors are not exactly orthogonal, the model in fact can be regarded as close to a two-way row-sparse SOFAR model (similar to Models 3 and 4 where both $\mathbf{U}^*$ and $\mathbf{V}^*$ are orthogonal and row-sparse); this is because if we compute the SVD of the true coefficient matrix, the resulting orthogonal factors will still have sparsity corresponding to the completely irrelevant responses and predictors.

In all the seven settings, we generated the data $\mathbf{Y}$ from the model $\mathbf{Y} = \mathbf{X} \mathbf{C}^* + \mathbf{E}$, where the error matrix $\mathbf{E}$ has i.i.d. rows from $N_q(\mathbf{0}, \sigma^2 \Sigma)$ with $\Sigma = (0.5^{|i-j|})$. In each simulation, σ^2 is computed to control the signal to noise ratio, defined as $\|\mathbf{d}_r^* \mathbf{X} \mathbf{u}_r^* \mathbf{v}_r^{*T}\|_F / \|\mathbf{E}\|_F$, to be exactly 1. The simulation was replicated 300 times in each setting.

All methods under comparison except OLS require selection of tuning parameters, which include the rank parameter in RRR, RSSVD, and SRRR and the regularization parameters

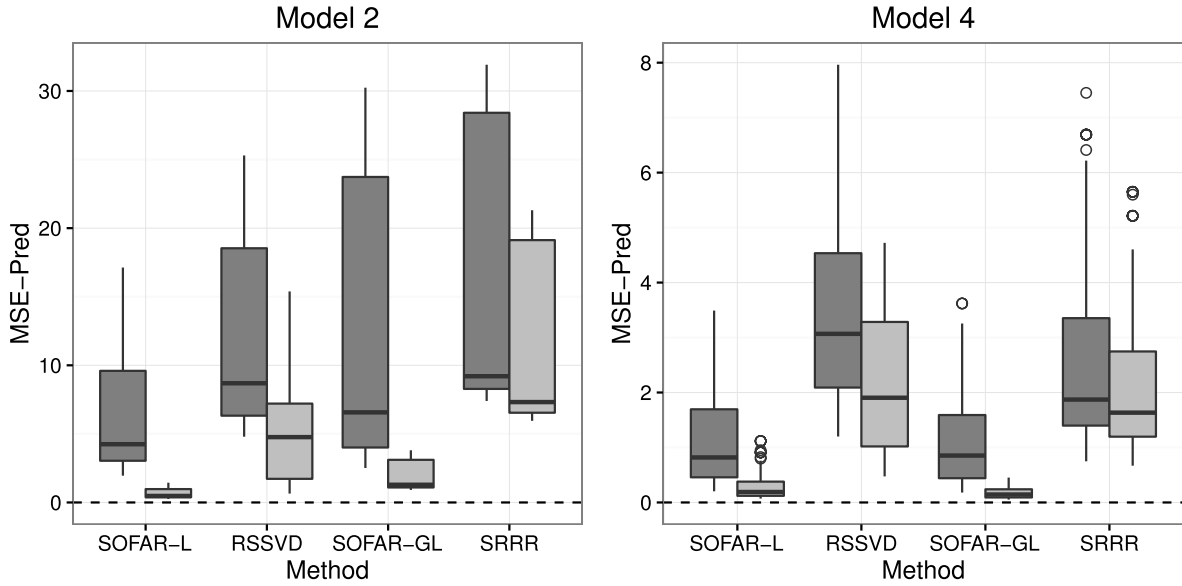


Fig. 1. Boxplots of $MSE\text{-}Pred$ for Models 2 and 4 with nonadaptive (dark gray) and adaptive (light gray) versions of various methods.

TABLE II
SIMULATION RESULTS FOR MODELS 1–2 WITH VARIOUS METHODS¹

Model	Method	$MSE\text{-}Est$	$MSE\text{-}Pred$	FPR (%)	FNR (%)	Rank	Rank (%)	Orth
1	OLS	250.7 (129.2)	753.8 (392.2)	100	0			
	Lasso	12.7 (5.9)	80.8 (34.1)	3.8	0			
	RRR	14.7 (6.8)	58.6 (29.3)	100	0	3	100	0
	SOFAR-L	0.4 (0.1)	2.8 (1.3)	0	0	3	100	0
	RSSVD	0.5 (0.3)	3.8 (2.3)	0.2	0	3	99.7	1.9
	SOFAR-GL	1.2 (0.5)	8.2 (4.1)	9.8	0	3	100	0
	SRRR	3.2 (1.0)	25.2 (12.6)	35.5	0	3	100	5.1
2	OLS	1013.0 (117.0)	765.6 (407.2)	100	0			
	Lasso	21.3 (7.0)	59.0 (18.1)	1.3	0			
	RRR	756.4 (56.8)	30.2 (15.9)	100	0	3	0	0
	SOFAR-L	0.2 (0.1)	0.7 (0.3)	0	0	3	0	0
	RSSVD	2.5 (2.4)	5.3 (4.1)	1	0.1	3	0	28.4
	SOFAR-GL	0.7 (0.4)	2.0 (1.0)	2.7	0	3	0	0
	SRRR	3.8 (1.5)	12.0 (6.3)	19.8	0	3	0	40.2

¹Adaptive versions of Lasso, SOFAR-L, RSSVD, SOFAR-GL, and SRRR were applied. Means of performance measures with standard deviations in parentheses over 300 replicates are reported. $MSE\text{-}Est$ values are scaled by multiplying 10^4 in Model 1 and 10^5 in Model 2, and $MSE\text{-}Pred$ values are scaled by multiplying 10^3 .

in SOFAR-L, SOFAR-GL, RSSVD, and SRRR. To reveal the full potential of each method, we chose the tuning parameters based on the predictive accuracy evaluated on a large, independently generated validation set of size 2000. The results with tuning parameters chosen by cross-validation or GIC [31] were similar to those based on a large validation set, and hence are not reported.

The model accuracy of each method is measured by the mean squared error $\|\hat{\mathbf{C}} - \mathbf{C}^*\|_F^2 / (pq)$ for estimation ($MSE\text{-}Est$) and $\|\mathbf{X}(\hat{\mathbf{C}} - \mathbf{C}^*)\|_F^2 / (nq)$ for prediction ($MSE\text{-}Pred$). The variable selection performance is characterized by the false positive rate (FPR%) and false negative rate (FNR%) in recovering the sparsity patterns of the SVD, that is, $FPR = FP / (TN + FP)$ and $FNR = FN / (TP + FN)$, where TP, FP, TN, and FN are the numbers of true nonzeros, false nonzeros, true zeros, and false zeros, respectively. The rank selection performance is evaluated by average estimated rank (Rank) and the percentage of

correct rank identification (Rank%). Finally, for the SOFAR-L, SOFAR-GL, and RSSVD methods which explicitly produce an SVD, the orthogonality of estimated factor matrices is measured by $100(\|\hat{\mathbf{U}}^T \hat{\mathbf{U}}\|_1 + \|\hat{\mathbf{V}}^T \hat{\mathbf{V}}\|_1 - 2r)$ (Orth), which is minimized at zero when exact orthogonality is achieved.

2) *Simulation Results:* We first compare the performance of nonadaptive and adaptive versions of the four sparse regularization methods. Because of the space constraint, only the results in terms of $MSE\text{-}Pred$ in high-dimensional models 2 and 4 are presented. The comparisons in other model settings are similar and thus omitted. From Fig. 1, we observe that adaptive weighting generally improves the empirical performance of each method. For this reason, we only consider the adaptive versions of these regularization methods in other comparisons.

The comparison results with adaptive penalty for Models 1 and 2 are summarized in Table II. The entrywise sparse SVD

TABLE III
SIMULATION RESULTS FOR MODELS 3–4 WITH VARIOUS METHODS¹

<i>Model</i>	<i>Method</i>	<i>MSE-Est</i>	<i>MSE-Pred</i>	<i>FPR (%)</i>	<i>FNR (%)</i>	<i>Rank</i>	<i>Rank (%)</i>	<i>Orth</i>
3	OLS	599.2 (339.2)	1530.1 (870.8)	100	0			
	Lasso	97.6 (50.0)	472.8 (242.7)	15.5	0.6			
	RRR	102.6 (70.2)	291.9 (191.8)	100	0	3	100	0
	SOFAR-L	24.8 (15.3)	129.5 (83.2)	0.3	7.4	3.7	30.3	0
	RSSVD	17.3 (11.3)	96.6 (66.4)	0.6	11	3	100	29
	SOFAR-GL	16.6 (11.4)	94.4 (67.5)	0.4	1.1	3.6	41.7	0
	SRRR	11.0 (6.7)	63.1 (40.2)	0.6	0.3	3	100	14.8
4	OLS	252.3 (78)	126.5 (65.4)	100	0			
	Lasso	37.4 (11.8)	73.2 (24.1)	0.8	2.5			
	RRR	186.6 (51.6)	6.1 (3.9)	100	0	3	99	0
	SOFAR-L	0.1 (0.1)	0.3 (0.2)	0.1	4.8	3	92.7	0.1
	RSSVD	1.0 (0.7)	2.2 (1.3)	0.3	11.5	3	100	40.1
	SOFAR-GL	0.1 (0.0)	0.2 (0.1)	0	0.1	3	100	0
	SRRR	0.8 (0.5)	2.0 (1.2)	24.9	0.2	3	100	31.3

¹Adaptive versions of Lasso, SOFAR-L, RSSVD, SOFAR-GL, and SRRR were applied. Means of performance measures with standard deviations in parentheses over 300 replicates are reported. *MSE-Est* values are scaled by multiplying 10^4 in Model 3 and 10^5 in Model 4, and *MSE-Pred* values are scaled by multiplying 10^3 .

TABLE IV
SIMULATION RESULTS FOR MODEL 5. WE USE MODEL 2 WITH INCREASED DIMENSIONS $p = 1000$, $q = 400$ BY ADDING NOISE VARIABLES¹

<i>Model</i>	<i>Method</i>	<i>MSE-Est</i>	<i>MSE-Pred</i>	<i>FPR (%)</i>	<i>FNR (%)</i>	<i>Rank</i>	<i>Rank (%)</i>	<i>Orth</i>
5	OLS	151.5 (5.7)	230.1 (122.9)	100	0			
	Lasso	3.9 (1.8)	29.3 (11.8)	0.6	0			
	RRR	146.8 (7.7)	61.5 (77.1)	100	0	2.6	57.7	0
	SOFAR-L	0.1 (0.0)	0.1 (0.0)	0	0	3	100	0
	RSSVD	6.6 (14.4)	2.8 (2.7)	3.1	1	3	99	49.1
	SOFAR-GL	0.1 (0.0)	0.2 (0.1)	0.8	0	3	100	0
	SRRR	0.5 (0.2)	3.6 (1.8)	19.7	0	3	100	55.5

¹Adaptive versions of Lasso, SOFAR-L, RSSVD, SOFAR-GL, and SRRR were applied. Means of performance measures with standard deviations in parentheses over 300 replicates are reported. *MSE-Est* values are scaled by 10^5 and *MSE-Pred* values are scaled by multiplying 10^3 .

structure is exactly what the SOFAR-L and RSSVD methods aim to recover. We observe that SOFAR-L performs the best among all methods in terms of both model accuracy and sparsity recovery. Although RSSVD performs only second to SOFAR-L in Model I, it has substantially worse performance in Model 2 in terms of model accuracy. This is largely because the RSSVD method does not impose any form of orthogonality constraints, which tends to cause nonidentifiability issues and compromise its performance in high dimensions. We note further that SOFAR-GL and SRRR perform worse than SOFAR-L, since they are not intended for entrywise sparsity recovery. However, these two methods still provide remarkable improvements over the OLS and RRR methods due to their ability to eliminate irrelevant variables, and over the Lasso method due to the advantages of imposing a low-rank structure. Compared to SRRR, the SOFAR-GL method results in fewer false positives and shows a clear advantage due to response selection.

The simulation results for Models 3 and 4 are reported in Table III. For the rowwise sparse SVD structure in these two models, SOFAR-GL and SRRR are more suitable than the other methods. All sparse regularization methods result in higher false negative rates than in Models 1 and 2 because of

the presence of some very weak signals. In Model 3, where the matrix of right singular vectors is not sparse and the dimensionality is moderate, SOFAR-GL has a slightly worse performance compared to SRRR since response selection is unnecessary. The advantages of SOFAR are clearly seen in Model 4, where the dimension is high and many irrelevant predictors and responses coexist; SOFAR-GL performs slightly better than SOFAR-L, and both methods substantially outperform the other methods. In both models, SOFAR-L and RSSVD result in higher false negative rates, since they introduce more parsimony than necessary by encouraging entrywise sparsity in $\mathbf{U}$ and $\mathbf{V}$. Table IV shows that the SOFAR methods still greatly outperform the others in both estimation and sparse recovery. In contrast, RSSVD becomes unstable and inaccurate; this again shows the effectiveness of enforcing the orthogonality in high-dimensional sparse SVD recovery.

Finally, Table V summarizes the results for Models 6 and 7. As expected, in Model 6 where the model dimensionality is low, SOFAR-GL performs the best, followed by RSSVD and SOFAR-L, whose performance is comparable to each other. In Model 7 where the model dimensionality is much higher, SOFAR-GL and SOFAR-L perform much better than RSSVD, as the latter becomes less stable. In both models, SRRR is

TABLE V
SIMULATION RESULTS FOR MODELS 6–7 WHEN THE SPARSE FACTORS ARE NOT EXACTLY ORTHOGONAL¹

Model	Method	MSE-Est	MSE-Pred	FPR (%)	FNR (%)	Rank	Rank (%)	Orth
6	OLS	291.7 (133.9)	868.5 (403.1)	100	0			
	Lasso	11.8 (4.7)	71.9 (28.4)	10.3	0			
	RRR	17.6 (7.3)	69.5 (31.2)	100	0	3	100	0
	SOFAR-L	2.2 (0.9)	14.1 (6.3)	8	0.1	3	100	0
	RSSVD	2.0 (0.7)	13.6 (6.1)	7.3	0.1	3	100	22.7
	SOFAR-GL	1.2 (0.5)	8.9 (3.9)	9.5	0	3	100	0
	SRRR	3.5 (1.1)	28.7 (13.1)	36.2	0	3	100	5.6
7	OLS	1045.1 (122.4)	886.5 (403.4)	100	0			
	Lasso	33.2 (12.7)	79.9 (24.4)	3.2	0			
	RRR	754.0 (67.8)	35.2 (15.4)	100	0	3	99.7	0
	SOFAR-L	1.2 (0.5)	3.1 (1.5)	2.4	0.2	3	100	0
	RSSVD	4.8 (4.1)	8.5 (5.1)	2.4	1.7	3	99.7	64.4
	SOFAR-GL	1.0 (0.4)	2.6 (1.1)	3.6	0	3	100	0
	SRRR	4.3 (1.5)	13.9 (6.2)	20	0	3	100	45

¹Adaptive versions of Lasso, SOFAR-L, RSSVD, SOFAR-GL, and SRRR were applied. Means of performance measures with standard deviations in parentheses over 300 replicates are reported. *MSE-Est* values are scaled by multiplying 10^4 in Model 6 and 10^5 in Model 7, and *MSE-Pred* values are scaled by multiplying 10^3 .

outperformed by SOFAR since the former does not pursue sparsity in $\mathbf{V}$. We have also experimented with other modified settings, and all the results are consistent with what have been reported. These results confirm that it is still preferable to apply SOFAR even when the underlying sparse factors are not exactly orthogonal, especially so in high-dimensional problems.

B. Real Data Analysis

In genetical genomics experiments, gene expression levels are treated as quantitative traits in order to identify expression quantitative trait loci (eQTLs) that contribute to phenotypic variation in gene expression. The task can be regarded as a multivariate regression problem with the gene expression levels as responses and the genetic variants as predictors, where both responses and predictors are often of high dimensionality. Most existing methods for eQTL data analysis exploit entrywise or rowwise sparsity of the coefficient matrix to identify individual genetic effects or master regulators [54], which not only tends to suffer from low detection power for multiple eQTLs that combine to affect a subset of gene expression traits, but also may offer little information about the functional grouping structure of the genetic variants and gene expressions. By exploiting a sparse SVD structure, the SOFAR method is particularly appealing for such applications, and may provide new insights into the complex genetics of gene expression variation. Here the orthogonality can be roughly interpreted as maximum separability, so that different association layers are more likely to reflect different functional pathways.

We illustrate our approach by the analysis of a yeast eQTL data set described by Brem and Kruglyak [14], where $n = 112$ segregants were grown from a cross between two budding yeast strains, BY4716 and RM11-1a. For each of the segregants, gene expression was profiled on microarrays containing 6216 genes, and genotyping was performed at 2957 markers. Similar to Yin and Li [64], we combined the markers into

blocks such that markers with the same block differed by at most one sample, and one representative marker was chosen from each block; a marginal gene–marker association analysis was then performed to identify markers that are associated with the expression levels of at least two genes with a p -value less than 0.05, resulting in a total of $p = 605$ markers.

Owing to the small sample size and weak genetic perturbations, we focused our analysis on $q = 54$ genes in the yeast MAPK signaling pathways [41]. We then applied the proposed SOFAR methods with adaptive weighting. Both SOFAR-L and SOFAR-GL methods resulted in a model of rank 3, indicating that dimension reduction is very effective for the data set. Also, the SVD layers estimated by the SOFAR methods are indeed sparse. The SOFAR-L estimates include 140 nonzeros in $\hat{\mathbf{U}}$, which involve only 112 markers, and 40 nonzeros in $\hat{\mathbf{V}}$, which involve only 27 genes. The sparse SVD produced by SOFAR-GL involves only 34 markers and 15 genes. The SOFAR-GL method is more conservative since it tends to identify markers that regulate all selected genes rather than a subset of genes involved in a specific SVD layer. We compare the original gene expression matrix $\mathbf{Y}$ and its estimates $\mathbf{X}\hat{\mathbf{C}}$ by the RRR, SOFAR-L and SOFAR-GL methods using heat maps in Fig. 2. It is seen that the SOFAR methods achieve both low-rankness and sparsity, while still capturing main patterns in the original matrix.

Fig. 3 shows the scatterplots of the latent responses $\mathbf{Y}\hat{\mathbf{v}}_j$ versus the latent predictors $\mathbf{X}\hat{\mathbf{u}}_j$ for $j = 1, 2, 3$, where $\hat{\mathbf{u}}_j$ and $\hat{\mathbf{v}}_j$ are the j th columns of $\hat{\mathbf{U}}$ and $\hat{\mathbf{V}}$, respectively. The plots demonstrate a strong association between each pair of latent variables, with the association strength descending from layer 1 to layer 3. A closer look at the SVD layers reveals further information about clustered samples and genes. The plot for layer 1 indicates that the yeast samples form two clusters, suggesting that our method may be useful for classification based on the latent variables. Also, examining the nonzero entries in $\hat{\mathbf{v}}_1$ shows that this layer is dominated by four genes, namely, STE3 (−0.66), STE2 (0.59), MFA2 (0.40),

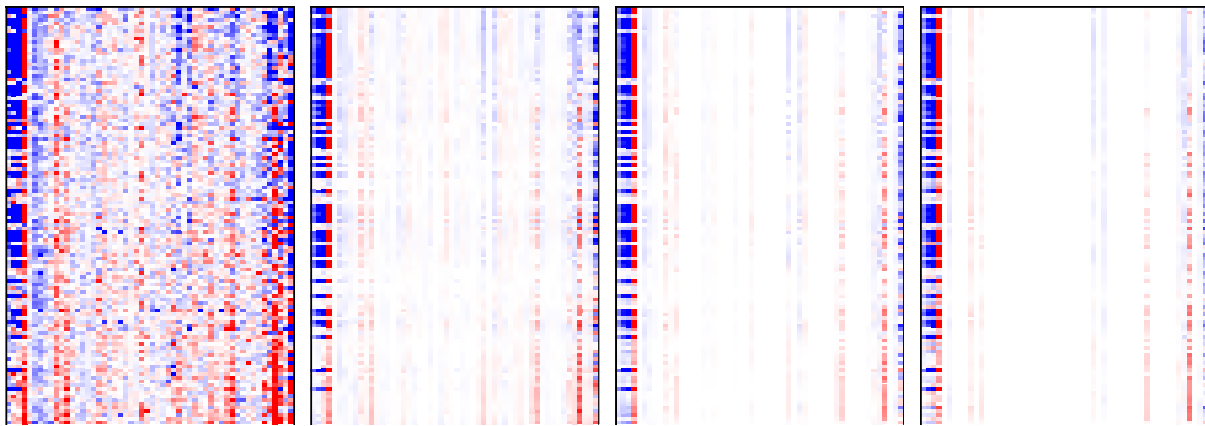


Fig. 2. Heat maps of $\mathbf{Y}$ and its estimates by RRR, SOFAR-L, and SOFAR-GL (from left to right).

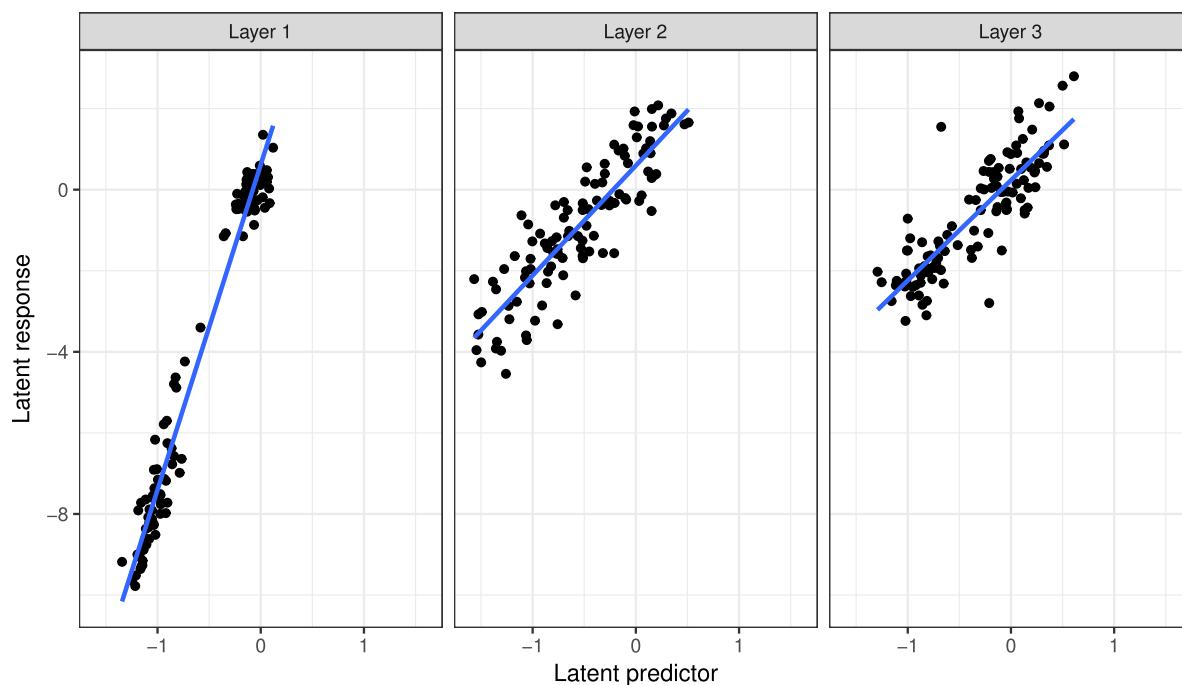


Fig. 3. Scatterplots of the latent responses versus the latent predictors in three SVD layers for the yeast data estimated by the SOFAR-L method.

and MFA1 (0.22). All four genes are upstream in the pheromone response pathway, where MFA2 and MFA1 are genes encoding mating pheromones and STE3 and STE2 are genes encoding pheromone receptors [24]. The second layer is mainly dominated by CTT1 (−0.93), and other leading genes include SLN1 (0.16), SLT2 (−0.14), MSN4 (−0.14), and GLO1 (−0.13). Interestingly, CTT1, MSN4, and GLO1 are all downstream genes linked to the upstream gene SLN1 in the high osmolarity/glycerol pathway required for survival in response to hyperosmotic stress. Finally, layer 3 includes the leading genes FUS1 (0.81), FAR1 (0.32), STE2 (0.25), STE3 (0.24), GPA1 (0.22), FUS3 (0.18), and STE12 (0.11). These genes consist of two major groups that are downstream (FUS1, FAR1, FUS3, and STE12) and upstream (STE2, STE3, and GPA1) in the pheromone response pathway. Overall, our results suggest that there are common genetic components shared by the expression traits of the clustered genes and

clearly reveal strong associations between the upstream and downstream genes on several signaling pathways, which are consistent with the current functional understanding of the MAPK signaling pathways.

To examine the predictive performance of SOFAR and other competing methods, we randomly split the data into a training set of size 92 and a test set of size 20. The model was fitted using the training set and the predictive accuracy was evaluated on the test set based on the prediction error $\|\mathbf{Y} - \mathbf{X}\hat{\mathbf{C}}\|_F^2/(nq)$. The splitting process was repeated 50 times. The scaled prediction errors for the RRR, SOFAR-L, SOFAR-GL, RSSVD, and SRRR methods are 3.4 (0.3), 2.6 (0.2), 2.5 (0.2), 2.9 (0.3), and 2.6 (0.2), respectively. The comparison shows the advantages of sparse and low-rank estimation. RSSVD yields higher prediction error and is less stable than the SOFAR methods. Although the SRRR method yielded similar predictive accuracy compared to SOFAR methods on this

data set, it resulted in a less parsimonious model and cannot be used for gene selection or clustering.

ACKNOWLEDGMENT

Part of this work was completed while Y. Fan and J. Lv visited the Departments of Statistics at University of California, Berkeley and Stanford University. These authors sincerely thank both departments for their hospitality.

REFERENCES

- [1] T. W. Anderson, "Estimating linear restrictions on regression coefficients for multivariate normal distributions," *Ann. Math. Statist.*, vol. 22, no. 3, pp. 327–351, 1951.
- [2] J. Bai, "Inferential theory for factor models of large dimensions," *Econometrica*, vol. 71, no. 1, pp. 135–171, 2003.
- [3] J. Bai and K. Li, "Statistical analysis of factor models of high dimension," *Ann. Statist.*, vol. 40, no. 1, pp. 436–465, 2012.
- [4] J. Bai and K. Li, "Maximum likelihood estimation and inference for approximate factor models of high dimension," *Rev. Econ. Statist.*, vol. 98, no. 2, pp. 298–309, 2016.
- [5] J. Bai and S. Ng, "Determining the number of factors in approximate factor models," *Econometrica*, vol. 70, no. 1, pp. 191–221, 2002.
- [6] J. Bai and S. Ng, "Large dimensional factor analysis," *Found. Trends Econ.*, vol. 3, no. 2, pp. 89–163, 2008.
- [7] J. Bai and S. Ng, "Principal components estimation and identification of static factors," *J. Econ.*, vol. 176, no. 1, pp. 18–29, 2013. [Online]. Available: <https://ideas.repec.org/a/eee/econom/v176y2013i1p18-29.html>
- [8] S. Basu and G. Michailidis, "Regularized estimation in sparse high-dimensional time series models," *Ann. Statist.*, vol. 43, no. 4, pp. 1535–1567, 2015.
- [9] K. Benidis, Y. Sun, P. Babu, and D. P. Palomar, "Orthogonal sparse pca and covariance estimation via Procrustes reformulation," *IEEE Trans. Signal Process.*, vol. 64, no. 23, pp. 6211–6226, Dec. 2016.
- [10] B. S. Bernanke, J. Boivin, and P. Elias, "Measuring the effects of monetary policy: A factor-augmented vector autoregressive (FAVAR) approach," *Quart. J. Econ.*, vol. 120, no. 1, pp. 387–422, 2005.
- [11] P. J. Bickel, Y. Ritov, and A. B. Tsybakov, "Simultaneous analysis of Lasso and Dantzig selector," *Ann. Statist.*, vol. 37, no. 4, pp. 1705–1732, 2009.
- [12] G. E. P. Box and G. C. Tiao, "A canonical analysis of multiple time series," *Biometrika*, vol. 64, no. 2, pp. 355–365, 1977.
- [13] S. Boyd, N. Parikh, E. Chu, B. Peleato, and J. Eckstein, "Distributed optimization and statistical learning via the alternating direction method of multipliers," *Found. Trends Mach. Learn.*, vol. 3, no. 1, pp. 1–122, Jan. 2011.
- [14] R. B. Brem and L. Kruglyak, "The landscape of genetic complexity across 5,700 gene expression traits in yeast," *Proc. Nat. Acad. Sci. USA*, vol. 102, no. 5, pp. 1572–1577, 2005.
- [15] F. Bunea, Y. She, and M. Wegkamp, "Optimal selection of reduced rank estimators of high-dimensional matrices," *Ann. Statist.*, vol. 39, no. 2, pp. 1282–1309, 2011.
- [16] F. Bunea, Y. She, and M. H. Wegkamp, "Joint variable and rank selection for parsimonious estimation of high-dimensional matrices," *Ann. Statist.*, vol. 40, no. 5, pp. 2359–2388, 2012.
- [17] S. Busygin, O. Prokopyev, and P. M. Pardalos, "Biclustering in data mining," *Comput. Oper. Res.*, vol. 35, no. 9, pp. 2964–2987, 2008.
- [18] T. T. Cai, H. Li, W. Liu, and J. Xie, "Covariate-adjusted precision matrix estimation with an application in genetical genomics," *Biometrika*, vol. 100, no. 1, pp. 139–156, 2013.
- [19] K. Chen, K.-S. Chan, and N. C. Stenseth, "Reduced rank stochastic regression with a sparse singular value decomposition," *J. Roy. Stat. Soc. B, Stat. Methodol.*, vol. 74, pp. 203–221, 2012.
- [20] K. Chen, K.-S. Chan, and N. C. Stenseth, "Source-sink reconstruction through regularized multicomponent regression analysis—With application to assessing whether north sea cod larvae contributed to local fjord cod in Skagerrak," *J. Amer. Stat. Assoc.*, vol. 109, no. 506, pp. 560–573, 2014.
- [21] K. Chen, H. Dong, and K.-S. Chan, "Reduced rank regression via adaptive nuclear norm penalization," *Biometrika*, vol. 100, no. 4, pp. 901–920, 2012.
- [22] L. Chen and J. Z. Huang, "Sparse reduced-rank regression for simultaneous dimension reduction and variable selection," *J. Amer. Statist. Assoc.*, vol. 107, no. 500, pp. 1533–1545, 2012.
- [23] L. Chen and J. Z. Huang, "Sparse reduced-rank regression with covariance estimation," *Statist. Comput.*, vol. 26, nos. 1–2, pp. 461–470, 2016.
- [24] R. E. Chen and J. Thorner, "Function and regulation in MAPK signaling pathways: Lessons learned from the yeast *Saccharomyces cerevisiae*," *Biochimica Biophys. Acta (BBA)-Mol. Cell Res.*, vol. 1773, no. 8, pp. 1311–1340, 2007.
- [25] A. d'Aspremont, L. El Ghaoui, M. I. Jordan, and G. R. G. Lanckriet, "A direct formulation for sparse PCA using semidefinite programming," *SIAM Rev.*, vol. 49, no. 3, pp. 434–448, 2007.
- [26] A. Edelman, T. A. Arias, and S. T. Smith, "The geometry of algorithms with orthogonality constraints," *SIAM J. Matrix Anal. Appl.*, vol. 20, no. 2, pp. 303–353, 1998.
- [27] J. Fan, Y. Fan, and E. Barut, "Adaptive robust variable selection," *Ann. Statist.*, vol. 42, no. 1, pp. 324–351, 2014.
- [28] J. Fan, Y. Fan, and J. Lv, "High dimensional covariance matrix estimation using a factor model," *J. Econ.*, vol. 147, no. 1, pp. 186–197, 2008.
- [29] J. Fan and R. Li, "Variable selection via nonconcave penalized likelihood and its oracle properties," *J. Amer. Stat. Assoc.*, vol. 96, no. 456, pp. 1348–1360, 2001.
- [30] Y. Fan and J. Lv, "Asymptotic equivalence of regularization methods in thresholded parameter space," *J. Amer. Stat. Assoc.*, vol. 108, no. 503, pp. 1044–1061, 2013.
- [31] Y. Fan and C. Y. Tang, "Tuning parameter selection in high dimensional penalized likelihood," *J. Roy. Stat. Soc. B, Stat. Methodol.*, vol. 75, no. 3, pp. 531–552, 2013.
- [32] J. Friedman, T. Hastie, H. Höfling, and R. Tibshirani, "Pathwise coordinate optimization," *Ann. Appl. Statist.*, vol. 1, no. 2, pp. 302–332, 2007.
- [33] G. Goh, D. K. Dey, and K. Chen, "Bayesian sparse reduced rank multivariate regression," *J. Multivariate Anal.*, vol. 157, pp. 14–28, May 2017.
- [34] G. H. Golub and C. F. Van Loan, *Matrix Computations*, 4th ed. Baltimore, MD, USA: The Johns Hopkins Univ. Press, 2013.
- [35] J. Guo, G. James, E. Levina, G. Michailidis, and J. Zhu, "Principal component analysis with sparse fused loadings," *J. Comput. Graph. Statist.*, vol. 19, no. 4, pp. 930–946, Sep. 2010.
- [36] M. C. Gustin, J. Albertyn, M. Alexander, and K. Davenport, "MAP kinase pathways in the yeast *Saccharomyces cerevisiae*," *Microbiol. Mol. Biol. Rev.*, vol. 62, no. 4, pp. 1264–1300, Dec. 1998.
- [37] J. A. Hartigan, "Direct clustering of a data matrix," *J. Amer. Statist. Assoc.*, vol. 67, no. 337, pp. 123–129, 1972.
- [38] N.-J. Hsu, H.-L. Hung, and Y.-M. Chang, "Subset selection for vector autoregressive processes using Lasso," *Comput. Statist. Data Anal.*, vol. 52, no. 7, pp. 3645–3657, 2008.
- [39] A. J. Izenman, "Reduced-rank regression for the multivariate linear model," *J. Multivariate Anal.*, vol. 5, no. 2, pp. 248–264, 1975.
- [40] I. M. Johnstone and A. Y. Lu, "On consistency and sparsity for principal components analysis in high dimensions," *J. Amer. Stat. Assoc.*, vol. 104, no. 486, pp. 682–693, 2009.
- [41] M. Kanehisa, S. Goto, Y. Sato, M. Kawashima, M. Furumichi, and M. Tanabe, "Data, information, knowledge and principle: Back to metabolism in KEGG," *Nucleic Acids Res.*, vol. 42, no. D1, pp. D199–D205, 2014.
- [42] A. Kock and L. Callot, "Oracle inequalities for high dimensional vector autoregressions," *J. Econ.*, vol. 186, no. 2, pp. 325–344, 2015.
- [43] M. A. Koschat and D. F. Swayne, "A weighted Procrustes criterion," *Psychometrika*, vol. 56, no. 2, pp. 229–239, 1991.
- [44] M. Lee, H. Shen, J. Z. Huang, and J. S. Marron, "Biclustering via sparse singular value decomposition," *Biometrics*, vol. 66, pp. 1087–1095, Dec. 2010.
- [45] C. Leng and H. Wang, "On general adaptive sparse principal component analysis," *J. Comput. Graph. Statist.*, vol. 18, no. 1, pp. 201–215, 2009.
- [46] H. Lian, S. Feng, and K. Zhao, "Parametric and semiparametric reduced-rank regression with flexible sparsity," *J. Multivariate Anal.*, vol. 136, pp. 163–174, Apr. 2015.
- [47] J. Lv, "Impacts of high dimensionality in finite samples," *Ann. Statist.*, vol. 41, no. 4, pp. 2236–2262, 2013.
- [48] X. Ma, L. Xiao, and W. H. Wong, "Learning regulatory programs by threshold SVD regression," *Proc. Nat. Acad. Sci. USA*, vol. 111, no. 44, pp. 15675–15680, 2014.
- [49] Z. Ma, Z. Ma, and T. Sun. (2014). "Adaptive estimation in two-way sparse reduced-rank regression." [Online]. Available: <https://arxiv.org/abs/1403.1922>
- [50] Z. Ma, Z. Ma, and T. Sun. (2014). "Adaptive estimation in two-way sparse reduced-rank regression." [Online]. Available: <https://arxiv.org/abs/1403.1922>

- [51] L. Mirsky, "Symmetric gauge functions and unitarily invariant norms," *Quart. J. Math.*, vol. 11, no. 1, pp. 50–59, 1960.
- [52] Y. Nardi and A. Rinaldo, "Autoregressive process modeling via the Lasso procedure," *J. Multivariate Anal.*, vol. 102, no. 3, pp. 528–549, 2011.
- [53] S. N. Negahban, P. Ravikumar, M. J. Wainwright, and B. Yu, "A unified framework for high-dimensional analysis of M -estimators with decomposable regularizers," *Stat. Sci.*, vol. 27, no. 4, pp. 538–557, 2012.
- [54] J. Peng *et al.*, "Regularized multivariate regression for identifying master predictors with application to integrative genomics study of breast cancer," *Ann. Appl. Statist.*, vol. 4, no. 1, pp. 53–77, 2010.
- [55] G. C. Reinsel and R. P. Velu, *Multivariate Reduced-Rank Regression: Theory and Applications*. New York, NY, USA: Springer, 1998.
- [56] F. Sha, Y. Lin, L. K. Saul, and D. D. Lee, "Multiplicative updates for nonnegative quadratic programming," *Neural Comput.*, vol. 19, no. 8, pp. 2004–2031, 2007.
- [57] H. Shen and J. Z. Huang, "Sparse principal component analysis via regularized low rank matrix approximation," *J. Multivariate Anal.*, vol. 99, no. 6, pp. 1015–1034, Jul. 2008.
- [58] J. H. Stock and M. W. Watson, "Vector autoregressions," *J. Econ. Perspect.*, vol. 15, no. 4, pp. 101–115, 2001.
- [59] J. H. Stock and M. W. Watson, "Forecasting using principal components from a large number of predictors," *J. Amer. Stat. Assoc.*, vol. 97, no. 460, pp. 1167–1179, 2002.
- [60] R. Tibshirani, "Regression shrinkage and selection via the lasso," *J. Roy. Statist. Soc. B, Methodol.*, vol. 58, no. 1, pp. 267–288, 1996.
- [61] P. Tseng, "Convergence of a block coordinate descent method for nondifferentiable minimization," *J. Optim. Theory Appl.*, vol. 109, no. 3, pp. 475–494, 2001.
- [62] R. P. Velu, G. C. Reinsel, and D. W. Wichern, "Reduced rank models for multiple time series," *Biometrika*, vol. 73, no. 1, pp. 105–118, 1986.
- [63] D. M. Witten, R. Tibshirani, and T. Hastie, "A penalized matrix decomposition, with applications to sparse principal components and canonical correlation analysis," *Biostatistics*, vol. 10, no. 3, pp. 515–534, Apr. 2009.
- [64] J. Yin and H. Li, "A sparse conditional Gaussian graphical model for analysis of genetical genomics data," *Ann. Appl. Statist.*, vol. 5, no. 4, pp. 2630–2650, 2011.
- [65] Y. Yu, T. Wang, and R. J. Samworth, "A useful variant of the Davis–Kahan theorem for statisticians," *Biometrika*, vol. 102, no. 2, pp. 315–323, 2015.
- [66] M. Yuan, A. Ekici, Z. Lu, and R. Monteiro, "Dimension reduction and coefficient estimation in multivariate linear regression," *J. Roy. Statist. Soc.*, vol. 69, no. 3, pp. 329–346, 2007.
- [67] C.-H. Zhang, "Nearly unbiased variable selection under minimax concave penalty," *Ann. Statist.*, vol. 38, no. 2, pp. 894–942, 2010.
- [68] Z. Zhang, H. Zha, and H. Simon, "Low-rank approximations with sparse factors I: Basic algorithms and error analysis," *SIAM J. Matrix Anal. Appl.*, vol. 23, no. 3, pp. 706–727, 2002.
- [69] Z. Zheng, Y. Fan, and J. Lv, "High dimensional thresholded regression and shrinkage effect," *J. Roy. Stat. Soc. B, Stat. Methodol.*, vol. 76, no. 3, pp. 627–649, 2014.
- [70] H. Zhu, Z. Khondker, Z. Lu, and J. G. Ibrahim, "Bayesian generalized low rank regression models for neuroimaging phenotypes and genetic markers," *J. Amer. Stat. Assoc.*, vol. 109, no. 507, pp. 990–997, 2014.
- [71] H. Zou, "The adaptive lasso and its oracle properties," *J. Amer. Stat. Assoc.*, vol. 101, no. 476, pp. 1418–1429, Dec. 2006.
- [72] H. Zou, T. Hastie, and R. Tibshirani, "Sparse principal component analysis," *J. Comput. Graph. Statist.*, vol. 15, no. 2, pp. 265–286, Jun. 2006.
- [73] H. Zou and R. Li, "One-step sparse estimates in nonconcave penalized likelihood models," *Ann. Statist.*, vol. 36, no. 4, pp. 1509–1533, 2008.

Yoshimasa Uematsu is Assistant Professor in Department of Economics and Management at Tohoku University, Sendai, Japan. He received his Ph.D. in Economics from Hitotsubashi University, Tokyo, in 2013. He was JSPS Research Fellow and Postdoctoral Scholar in Research Center for Statistical Machine Learning of The Institute of Statistical Mathematics, Tokyo, and in Data Sciences and Operations Department of the Marshall School of Business at the University of Southern California, Los Angeles. His research interests include high-dimensional statistics and econometrics, time series analysis, and forecasting.

Yingying Fan is Dean's Associate Professor in Business Administration in Data Sciences and Operations Department of the Marshall School of Business at the University of Southern California, Associate Professor in Departments of Economics and Computer Science at USC, and an Associate Fellow of USC Dornsife Institute for New Economic Thinking. She received her Ph.D. in Operations Research and Financial Engineering from Princeton University in 2007. She was Lecturer in the Department of Statistics at Harvard University from 2007–2008. Her research interests include statistics, data science, machine learning, economics, and big data and business applications.

Her papers have been published in journals in statistics, economics, computer science, and information theory. She is the recipient of the Royal Statistical Society Guy Medal in Bronze, the American Statistical Association Noether Young Scholar Award, NSF Faculty Early Career Development (CAREER) Award, NIH R01 Grant, USC Marshall Dean's Award for Research Excellence, the USC Marshall Inaugural Dr. Douglas Basil Award for Junior Business Faculty, Zumberge Individual Award from USC's James H. Zumberge Faculty Research and Innovation Fund, and USC Marshall Dean's Award for Research Excellence, as well as a Plenary Speaker at the 2011 Institute of Mathematical Statistics Workshop on Finance, Probability, and Statistics held at Columbia University. She has served as an associate editor of *Journal of the American Statistical Association*, *Journal of Econometrics*, *Journal of Business & Economic Statistics*, *The Econometrics Journal*, and *Journal of Multivariate Analysis*.

Kun Chen is Associate Professor in the Department of Statistics at the University of Connecticut, and Research Fellow at the Center for Population Health at the University of Connecticut Health Center. He received his B.Econ. in Finance and Dual B.S. in Computer Science & Technology from the University of Science and Technology of China in 2003, his M.S. in Statistics from the University of Alaska Fairbanks in 2007, and his Ph.D. in Statistics from the University of Iowa in 2011. His research mainly focuses on multivariate statistical learning, dimension reduction, high-dimensional statistics, and healthcare analytics with large-scale heterogeneous data.

Jinchi Lv is Kenneth King Stonier Chair in Business Administration and Professor in Data Sciences and Operations Department of the Marshall School of Business at the University of Southern California, Professor in Department of Mathematics at USC, and an Associate Fellow of USC Dornsife Institute for New Economic Thinking. He received his Ph.D. in Mathematics from Princeton University in 2007. He was McAlister Associate Professor in Business Administration at USC from 2016–2019. His research interests include statistics, machine learning, data science, and business applications.

His papers have been published in journals in statistics, economics, computer science, information theory, and biology, and one of them was published as a Discussion Paper in Journal of the Royal Statistical Society Series B. He is the recipient of the Royal Statistical Society Guy Medal in Bronze, NSF Faculty Early Career Development (CAREER) Award, USC Marshall Dean's Award for Research Impact, Adobe Data Science Research Award, USC Marshall Dean's Award for Research Excellence, and Zumberge Individual Award from USC's James H. Zumberge Faculty Research and Innovation Fund. He has served as an associate editor of the *Annals of Statistics*, *Journal of Business & Economic Statistics*, and *Statistica Sinica*.

Wei Lin is currently Assistant Professor in the School of Mathematical Sciences and the Center for Statistical Science at Peking University, Beijing, China. He received his Ph.D. degree in Applied Mathematics from the University of Southern California, Los Angeles, CA, in 2011. His research interests include high-dimensional statistics, machine learning, causal inference, compositional data analysis, survival analysis, and statistical genetics and genomics.