

Are Latent Factor Regression and Sparse Regression Adequate?

Jianqing Fan Zhipeng Lou Mengxin Yu

Abstract

We propose the **F**actor **A**ugmented (sparse linear) **R**egression **M**odel (FARM) that not only admits both the latent factor regression and sparse linear regression as special cases but also bridges dimension reduction and sparse regression together. We provide theoretical guarantees for the estimation of our model under the existence of sub-Gaussian and heavy-tailed noises (with bounded $(1 + \vartheta)$ -th moment, for all $\vartheta > 0$) respectively. In addition, the existing works on supervised learning often assume the latent factor regression or sparse linear regression is the true underlying model without justifying its adequacy. To fill in such an important gap on high-dimensional inference, we also leverage our model as the alternative model to test the sufficiency of the latent factor regression and the sparse linear regression models. To accomplish these goals, we propose the **F**actor-**A**ddjusted **d**e**B**iased **T**est (FabTest) and a two-stage ANOVA type test respectively. We also conduct large-scale numerical experiments including both synthetic and FRED macroeconomics data to corroborate the theoretical properties of our methods. Numerical results illustrate the robustness and effectiveness of our model against latent factor regression and sparse linear regression models.

Keyword: Factor model, High-dimensional Inference, Hypothesis, Sparse linear regression, Robustness

¹Jianqing Fan is Frederick L. Moore '18 Professor of Finance, Professor of Statistics, and Professor of Operations Research and Financial Engineering at the Princeton University. Zhipeng Lou is a Postdoctoral Researcher at Department of Operations Research and Financial Engineering, Princeton University. Mengxin Yu is a Ph.D. student at Department of Operations Research and Financial Engineering, Princeton University, Princeton, NJ 08544, USA. Emails: {jqfan, zlou, mengxiny}@princeton.edu. The research is supported in part by the ONR grant N00014-22-1-2340, NSF grants DMS-2210833, DMS-2053832, DMS-2052926 and NIH grant 2R01-GM072611-16

1 Introduction

Over the past two decades, along with the development of technology, datasets with high-dimensionality in various fields such as biology, genomics, neuroscience and finance have been collected. One stylized feature of the high-dimensional data is the high dependence across features that give rises to near collinearity. A common structure to characterize the dependence across features is the approximate factor model [Bai, 2003, Fan et al., 2013], in which the variables are correlated with each other through several common latent factors. More specifically, we assume the observed d -dimensional covariate vector $\mathbf{x}$ follows from the model

$$\mathbf{x} = \mathbf{B}\mathbf{f} + \mathbf{u}, \quad (1.1)$$

where $\mathbf{f}$ is a K -dimensional vector of latent factors, $\mathbf{B} \in \mathbb{R}^{d \times K}$ is the corresponding factor loading matrix, and $\mathbf{u}$ is a d -dimensional vector of idiosyncratic component which is uncorrelated with $\mathbf{f}$.

To tackle the high-dimensionality of datasets, various methods have been proposed. Among these, dimensionality reduction and sparse regression are two popularly used ones to circumvent the curse of dimensionality. They also serve as the backbones for many emerging statistical methods.

In terms of dimension reduction, the factor regression model is one of the most popular methods and has been widely used [Stock and Watson, 2002, Bai and Ng, 2006, Bair et al., 2006, Bai and Ng, 2008, Fan et al., 2017b, Bing et al., 2019, Bunea et al., 2020, Bing et al., 2021]. It assumes that the latent factors drive both dependent and independent variables as follows:

$$\begin{aligned} Y &= \mathbf{f}^\top \boldsymbol{\gamma} + \varepsilon, \\ \mathbf{x} &= \mathbf{B}\mathbf{f} + \mathbf{u}. \end{aligned} \quad (1.2)$$

Here Y is the response variable and $\varepsilon \in \mathbb{R}$ is the random noise which is independent with the factor $\mathbf{f}$. When the factors are unobserved, one usually learns the latent factors based on observed $\mathbf{x}$ and substitutes the sample version into the regression model (1.2). There are several methods for estimating latent factors such as Principal Component Analysis (PCA) [Bai, 2003, Fan et al., 2013], maximum likelihood estimation [Bai and Li, 2012], and random projections [Fan and Liao, 2020]. In particular,

when the leading Principal Components are used as an estimator for $\mathbf{f}$, the sample version of (1.2) reduces to the classical Principal Component Regression (PCR) [Hotelling, 1933].

As for sparse regression, a commonly used model is the following (sparse) linear regression:

$$Y = \mathbf{x}^\top \boldsymbol{\beta} + \varepsilon. \quad (1.3)$$

In the high dimensional regime where the dimension d can be much larger than the sample size n , it is commonly assumed that the population parameter vector $\boldsymbol{\beta} \in \mathbb{R}^d$ is sparse. Over the last two decades, various regularized methods, which incorporate this notion of sparsity, have been proposed. See, for instance, LASSO [Tibshirani, 1996], SCAD [Fan and Li, 2001], Least Angle Regression [Efron et al., 2004], Dantzig selector [Candes and Tao, 2007], Adaptive LASSO [Zou, 2006], MCP [Zhang, 2010] and many others. For more details, please refer to Fan et al. [2020b] for a comprehensive account.

In this paper, we introduce the **F**actor **A**ugmented (sparse linear) **R**egression **M**odel (FARM) (1.4), which incorporates both the latent factor and the idiosyncratic component into the covariates,

$$\begin{aligned} Y &= \mathbf{f}^\top \boldsymbol{\gamma}^* + \mathbf{u}^\top \boldsymbol{\beta}^* + \varepsilon, \\ \mathbf{x} &= \mathbf{B}\mathbf{f} + \mathbf{u}, \end{aligned} \quad (1.4)$$

where $\boldsymbol{\gamma}^* \in \mathbb{R}^K$ and $\boldsymbol{\beta}^* \in \mathbb{R}^d$ are population parameter vectors quantifying the contribution of the latent factor $\mathbf{f}$ and the idiosyncratic component $\mathbf{u}$, respectively. Obviously, the factor regression model (1.2) is a special case of (2.1) in which $\boldsymbol{\beta}^* = \mathbf{0}$. To better illustrate the difference between model (1.4) and the sparse linear model (1.3), our model can be written in an equivalent form,

$$Y = \mathbf{f}^\top \boldsymbol{\varphi}^* + \mathbf{x}^\top \boldsymbol{\beta}^* + \varepsilon, \quad \mathbf{x} = \mathbf{B}\mathbf{f} + \mathbf{u}, \quad (1.5)$$

where $\boldsymbol{\varphi}^* = \boldsymbol{\gamma}^* - \mathbf{B}^\top \boldsymbol{\beta}^* \in \mathbb{R}^K$ quantifies the extra contribution of the latent factor $\mathbf{f}$ beyond the observed predictor $\mathbf{x}$. Therefore, FARM expands the space spanned by $\mathbf{x}$ into useful directions spanned by $\mathbf{f}$. It is clear that the sparse regression model (1.3) is also a special case of (1.4) with $\boldsymbol{\varphi}^* = \mathbf{0}$. Thus, our model is general enough to bridge the dimensionality reduction and the sparse regression.

The motivation of our factor augmented linear model (1.4) comes from two perspectives.

1. Firstly, it origins from [Fan et al. \[2020a\]](#). In order to get precise estimation of β^* based on highly correlated variables, they study the sparse regression estimation by substituting (1.1) into (1.3) and obtain

$$Y = (\mathbf{B}\mathbf{f} + \mathbf{u})^\top \beta^* + \varepsilon = \mathbf{f}^\top (\mathbf{B}^\top \beta^*) + \mathbf{u}^\top \beta^* + \varepsilon. \quad (1.6)$$

We observe from (1.6), when the sparse linear regression is adequate, for a given β^* , the regression coefficient on $\mathbf{f}$ is fixed at $\gamma^* = \mathbf{B}^\top \beta^*$. However, in reality, especially when the variables are highly correlated, it is very likely that the leading factors possess extra contributions to the response instead of only a fixed portion $\mathbf{B}^\top \beta^*$. This results in our proposition of model (1.4), where we augment the leading factors into sparse regression that expands the linear space spanned by $\mathbf{x}$ into useful directions.

2. Secondly, it origins from the factor regression given in (1.2). In reality, the leading common factors $\mathbf{f}$ indeed provides some important contributions to the response, but it is hard to believe that they will have fully explanation power, especially when the effect of the factors is weak. Besides, in real applications, several examples illustrate the poor performance of factor regression model or PCR, see [Jolliffe \[1982\]](#) for more details. Thus, completely ignoring the idiosyncratic component $\mathbf{u}$ will harm in model generalization. This also motivates us to propose model (1.4), in which we augment the sparse regression by incorporating the idiosyncratic component $\mathbf{u}$ into the original factor regression.

In this paper, we first study the properties of estimated parameters under the proposed model (1.4). Specifically, we assume the factors given in (1.4) are unobserved and leverage PCA to estimate them. Incorporated with penalized least-squares with the ℓ_1 -penalty, we derive the ℓ_2 -consistency results for parameter vectors γ^* and β^* . Going beyond the linear regression model and the least squares estimation, our idea can be naturally extended to more general supervised learning models through different loss functions. For instance, quantile regression [[Belloni and Chernozhukov, 2011](#), [Fan et al., 2014](#)], support vector machine [[Zhang et al., 2016](#), [Peng et al., 2016](#)], Huber regression [[Fan et al., 2017a](#), [Sun et al., 2020](#)], generalized linear model [[Van de Geer, 2008](#), [Fan et al., 2020a](#)] and many other variants.

In order to demonstrate the general applicability of our proposed methods, in our paper, we further extend our model settings to robust regression. To be more specific, we only assume the existence of $(1 + \vartheta)$ -th moment of the noise distribution for some $\vartheta > 0$. We adopt Huber loss together with adaptive tuning parameters and ℓ_1 -penalization to derive the consistency results for the parameters of our interest. Besides the aforementioned extensions, it is worth to note that our model is also applicable in the field of causal inference [Imbens and Rubin, 2015, Hernan and Robins, 2019]. To be more specific, the latent factors $\mathbf{f}$ given in our model is able to be treated as the unobserved confounding variables which affect both the covariate $\mathbf{x}$ and the response Y . From the causal perspective, we provide a methodology to conduct (robust) statistical estimation as well as inference of our model under the existence of latent confounding variables.

The aforementioned works on factor regression and sparse linear regression mainly investigate the theoretical properties based on the assumption that either of them is the true underlying model [Stock and Watson, 2002, Tibshirani, 1996, Fan and Li, 2001, Zou, 2006, Bai and Ng, 2006, Zhang, 2010, Fan et al., 2017b, 2020a, Bing et al., 2021]. However, whether a given model is adequate to explain a given dataset plays a crucial role in the model selection step. This motivates us to fill the gap by leveraging our model as the alternative one to perform hypothesis testing on the adequacy of the factor regression model as well as the sparse linear regression model when covariates admit a factor structure.

For the hypothesis test on the adequacy of the latent factor regression model, we consider testing the hypotheses

$$H_0 : Y = \mathbf{f}^\top \boldsymbol{\gamma}^* + \varepsilon \text{ versus } H_1 : Y = \mathbf{f}^\top \boldsymbol{\gamma}^* + \mathbf{u}^\top \boldsymbol{\beta}^* + \varepsilon. \quad (1.7)$$

This amounts to testing $H_0 : \boldsymbol{\beta}^* = 0$ under FARM model. To this end, we propose the **Factor-Adjusted deBiased Test** statistic (FabTest) $\tilde{\beta}_\lambda$ which serves as a de-sparsify version of the estimator $\hat{\beta}_\lambda$ obtained under ℓ_1 -regularization. The asymptotic distribution of the proposed test statistic is derived by leveraging the high-dimensional Gaussian approximation. The critical value controlling the Type-I error is estimated based on the multiplier bootstrap method. As a byproduct, we are also able to conduct entrywise and groupwise hypothesis testing on parameter $\boldsymbol{\beta}^*$ by following similar de-biasing procedure.

For validating the adequacy of the sparse linear regression model, we consider testing the hypotheses

$$H_0 : Y = \mathbf{x}^\top \boldsymbol{\beta}^* + \varepsilon \text{ versus } H_1 : Y = \mathbf{f}^\top \boldsymbol{\varphi}^* + \mathbf{x}^\top \boldsymbol{\beta}^* + \varepsilon, \quad (1.8)$$

or $\boldsymbol{\varphi}^* = 0$ under the FARM model. To tackle the testing problem, we propose a two-stage ANOVA test. In the first stage, we use marginal screening [Fan and Lv, 2008] to pre-select a group of variables which cope well the curse of high dimensionality. In the second stage, we derive the ANOVA-type test statistic. Asymptotic null distribution and the power of the test statistic are derived. In addition, we further extend the aforementioned two-stage ANOVA test to linear multi-modal models [Li and Li, 2021], whose data framework has been well applied in a wide range of scientific fields (e.g multi-omics data in genomics, multimodal neuroimaging data in neuroscience, multimodal electronic health records data in health care).

In summary, our main contributions are as follows:

1. Motivated from the factor regression and sparse regression, we propose the **Factor Augmented** (sparse linear) **Regression Model** (FARM) (1.4) [also (1.5)] and investigate in the parameter estimation properties on $\boldsymbol{\gamma}^*$ and $\boldsymbol{\beta}^*$ given in (1.4). Our work serves as an extension of Fan et al. [2020a] to a general setting with weaker assumptions. It augments the sparse linear regression in useful directions of common factors.
2. To further demonstrate the wide applicability of our methods, we extend our model to a more robust setting, where we only assume the existence of $(1 + \vartheta)$ -th moment ($\vartheta > 0$) of our noise distribution. Leveraging the ℓ_1 -penalized adaptive Huber estimation, we establish statistical estimation results for our parameters of interest. Comparing with those closely related literature [Fan et al., 2020a, 2021a], our assumption on the moment condition of the noise variable is the weakest. Our robustified factor augmented regression also serves as an extension of Sun et al. [2020] to a more general setting.
3. In terms of testing the adequacy of the factor regression, we propose the **FabTest** by incorporating the factor structure into the de-biased estimators [van de Geer et al., 2014, Zhang and

Zhang, 2014, Javanmard and Montanari, 2014]. Accompanied with Gaussian approximation, the asymptotic distribution of our test statistic is derived. As for implementation, we propose the multiplier bootstrap method to estimate the critical value in order to control the Type-I error.

4. For testing the adequacy of sparse linear regression model, we propose a two stage ANOVA-type testing procedure. Asymptotic distribution (under the null) and power (under the alternative) of our constructed test statistic are investigated. In addition, we further extend the methodology to the multi-modal sparse linear regression model [Li and Li, 2021], by testing whether the sparse linear regression for some given modals is adequate.
5. We conduct large scale simulation studies for our proposed methodology using both synthetic data and real data. Simulation results via synthetic data lend further support to our theoretical findings. As for real data, we apply our methodology to the studies of the macroeconomics dataset named FRED-MD [McCracken and Ng, 2016]. The experimental results also illustrate the high efficiency and robustness of our model (FARM) against latent factor regression as well as sparse linear regression.

1.1 Notation

For a vector $\boldsymbol{\gamma} = (\gamma_1, \dots, \gamma_m)^\top \in \mathbb{R}^m$, we denote its ℓ_q norm as $\|\boldsymbol{\gamma}\|_q = (\sum_{\ell=1}^m |\gamma_\ell|^q)^{1/q}$, $1 \leq q < \infty$, and write $\|\boldsymbol{\gamma}\|_\infty = \max_{1 \leq \ell \leq m} |\gamma_\ell|$. For any integer m , we denote $[m] = \{1, \dots, m\}$. The sub-Gaussian norm of a scalar random variable Z is defined as $\|Z\|_{\psi_2} = \inf\{t > 0 : \mathbb{E} \exp(Z^2/t^2) \leq 2\}$. For a random vector $\boldsymbol{x} \in \mathbb{R}^m$, we use $\|\boldsymbol{x}\|_{\psi_2} = \sup_{\|\boldsymbol{v}\|_2=1} \|\boldsymbol{v}^\top \boldsymbol{x}\|_{\psi_2}$ to denote its sub-Gaussian norm. Let $\mathbb{I}\{\cdot\}$ denote the indicator function and let $\mathbf{I}_K$ denotes the identity matrix in $\mathbb{R}^{K \times K}$. For a matrix $\mathbf{A} = [A_{jk}]$, we define $\|\mathbf{A}\|_{\mathbb{F}} = \sqrt{\sum_{jk} A_{jk}^2}$, $\|\mathbf{A}\|_{\max} = \max_{jk} |A_{jk}|$ and $\|\mathbf{A}\|_\infty = \max_j \sum_k |A_{jk}|$ to be its Frobenius norm, element-wise max-norm and matrix ℓ_∞ -norm, respectively. Moreover, we use $\lambda_{\min}(\mathbf{A})$ and $\lambda_{\max}(\mathbf{A})$ to denote the minimal and maximal eigenvalues of $\mathbf{A}$, respectively. We use $|\mathcal{A}|$ to denote the cardinality of set $\mathcal{A}$. For two positive sequences $\{a_n\}_{n \geq 1}$, $\{b_n\}_{n \geq 1}$, we write $a_n = O(b_n)$ if there exists a positive constant C such that $a_n \leq C \cdot b_n$ and we write $a_n = o(b_n)$ if $a_n/b_n \rightarrow 0$. In addition, $a_n = O_{\mathbb{P}}(b_n)$ and $a_n = o_{\mathbb{P}}(b_n)$ have similar meanings as above except that the relationship of a_n/b_n

holds with high probability.

1.2 RoadMap

The rest of this paper is organized as follows. We study the parameter estimation properties of our proposed model (FARM) in section 2, where theoretical results of both regular and robust estimators are analyzed. In section 3, we construct a de-biased test statistic to test the adequacy of latent factor regression model. In addition, in section 4, we construct a two-stage ANOVA test to study the adequacy of sparse linear regression under the setting with highly correlated features. Moreover, to corroborate our theoretical findings, in section 5, we conduct exhaustive simulation studies. Last but not least, we apply our methodology to study the real data FRED-MD in section 5.4.

2 Factor Augmented Regression Model

The primary objective of this section is to propose a regularized estimation method for our factor augmented sparse linear model and investigate the corresponding statistical properties. Suppose we observe n independent and identically distributed (i.i.d.) random samples $\{(\mathbf{x}_t, Y_t)\}_{t=1}^n$ from $(\mathbf{x}, Y)$, which satisfy that

$$\mathbf{x}_t = \mathbf{B}\mathbf{f}_t + \mathbf{u}_t \text{ and } Y_t = \mathbf{f}_t^\top \boldsymbol{\gamma}^* + \mathbf{u}_t^\top \boldsymbol{\beta}^* + \varepsilon_t, \quad t = 1, \dots, n, \quad (2.1)$$

where $\mathbf{f}_1, \dots, \mathbf{f}_n \in \mathbb{R}^K$, $\mathbf{u}_1, \dots, \mathbf{u}_n \in \mathbb{R}^d$ and $\varepsilon_1, \dots, \varepsilon_n \in \mathbb{R}$ are i.i.d. realizations of $\mathbf{f}$, $\mathbf{u}$ and ε , respectively. To ease the presentation, we rewrite (2.1) in a more compact matrix form as follows,

$$\begin{aligned} \mathbf{X} &= \mathbf{F}\mathbf{B}^\top + \mathbf{U}, \\ \mathbf{Y} &= \mathbf{F}\boldsymbol{\gamma}^* + \mathbf{U}\boldsymbol{\beta}^* + \boldsymbol{\varepsilon}, \end{aligned} \quad (2.2)$$

where $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_n)^\top$, $\mathbf{F} = (\mathbf{f}_1, \dots, \mathbf{f}_n)^\top$, $\mathbf{U} = (\mathbf{u}_1, \dots, \mathbf{u}_n)^\top$, $\mathbf{Y} = (Y_1, \dots, Y_n)^\top$ and $\boldsymbol{\varepsilon} = (\varepsilon_1, \dots, \varepsilon_n)^\top$. Throughout the whole paper, we assume we only get access to observations $\{(\mathbf{x}_t, Y_t)\}_{t=1}^n$. Both the latent factors $\mathbf{F}$ and the idiosyncratic components $\mathbf{U}$ are unobserved and need

to be estimated from the observed predictors $\mathbf{X}$. Thus, in the following, we shall first illustrate how to estimate $\mathbf{F}$ and $\mathbf{U}$ and then proceed with the regularized estimation for model (2.2).

2.1 Factor Estimation

Since only the predictor vector $\mathbf{x}$ is observable, the latent factor $\mathbf{f}$ and the corresponding loading matrix $\mathbf{B}$ are not identifiable under the factor model (1.1). More specifically, for any non-singular matrix $\mathbf{S} \in \mathbb{R}^{K \times K}$, we have $\mathbf{x} = \mathbf{B}\mathbf{f} + \mathbf{u} = (\mathbf{B}\mathbf{S})(\mathbf{S}^{-1}\mathbf{f}) + \mathbf{u}$. To resolve this issue, we impose the following identifiability conditions [Bai, 2003, Fan et al., 2013]:

$$\text{Cov}(\mathbf{f}) = \mathbf{I}_K \text{ and } \mathbf{B}^\top \mathbf{B} \text{ is diagonal.}$$

Consequently, the constrained least squares estimator of $(\mathbf{F}, \mathbf{B})$ based on $\mathbf{X}$ is given by

$$\begin{aligned} (\hat{\mathbf{F}}, \hat{\mathbf{B}}) &= \arg \min_{\mathbf{F} \in \mathbb{R}^{n \times K}, \mathbf{B} \in \mathbb{R}^{d \times K}} \|\mathbf{X} - \mathbf{F}\mathbf{B}^\top\|_{\mathbb{F}}^2 \\ \text{subject to } &\frac{1}{n} \mathbf{F}^\top \mathbf{F} = \mathbf{I}_K \text{ and } \mathbf{B}^\top \mathbf{B} \text{ is diagonal.} \end{aligned}$$

Elementary manipulation yields that the columns of $\hat{\mathbf{F}}/\sqrt{n}$ are the eigenvectors corresponding to the largest K eigenvalues of the matrix $\mathbf{X}\mathbf{X}^\top$ and $\hat{\mathbf{B}} = (\hat{\mathbf{F}}^\top \hat{\mathbf{F}})^{-1} \hat{\mathbf{F}}^\top \mathbf{X} = n^{-1} \hat{\mathbf{F}}^\top \mathbf{X}$. Then the least squares estimator for $\mathbf{U}$ is given by $\hat{\mathbf{U}} = \mathbf{X} - \hat{\mathbf{F}}\hat{\mathbf{B}}^\top = (\mathbf{I}_n - n^{-1} \hat{\mathbf{F}}\hat{\mathbf{F}}^\top) \mathbf{X}$.

Before presenting the asymptotic properties of the estimators $\{\hat{\mathbf{F}}, \hat{\mathbf{B}}, \hat{\mathbf{U}}\}$, we first impose some regularity conditions.

Assumption 2.1. There exists a positive constant $c_0 < \infty$ such that $\|\mathbf{f}\|_{\psi_2} \leq c_0$ and $\|\mathbf{u}\|_{\psi_2} \leq c_0$.

Assumption 2.2. There exists a constant $\tau > 1$ such that $d/\tau \leq \lambda_{\min}(\mathbf{B}^\top \mathbf{B}) \leq \lambda_{\max}(\mathbf{B}^\top \mathbf{B}) \leq d\tau$. Moreover, we assume $n \log^2 n = O(d)$.

Assumption 2.3. Let $\Sigma = \text{Cov}(\mathbf{u})$. There exists a constant $\Upsilon > 0$ such that $\|\mathbf{B}\|_{\max} \leq \Upsilon$ and

$$\mathbb{E}|\mathbf{u}^\top \mathbf{u} - \text{tr}(\Sigma)|^4 \leq \Upsilon d^2.$$

Assumption 2.4. There exist a positive constant $\kappa < 1$ such that $\kappa \leq \lambda_{\min}(\Sigma)$, $\|\Sigma\|_1 \leq 1/\kappa$ and $\min_{1 \leq k, \ell \leq d} \text{Var}(u_k u_\ell) \geq \kappa$.

Remark 1. Assumptions 2.1–2.4 are standard assumptions in the studies of large dimensional factor model. We refer to Bai [2003], Fan et al. [2013] and Li et al. [2018] for more details. $\square$

We next summarize the theoretical results related to consistent factor estimation in the following proposition which directly follows from Lemmas D.1 and D.2 in Wang and Fan [2017].

Proposition 2.1. Assume that $d = o(\exp(n))$. Let $\mathbf{H} = n^{-1}\mathbf{V}^{-1}\hat{\mathbf{F}}^\top \mathbf{F}\mathbf{B}^\top \mathbf{B}$, where $\mathbf{V} \in \mathbb{R}^{K \times K}$ is a diagonal matrix consisting of the first K largest eigenvalues of the matrix $n^{-1}\mathbf{X}\mathbf{X}^\top$. Then, under Assumptions 2.1–2.4, we have

1. $\|\hat{\mathbf{F}} - \mathbf{F}\mathbf{H}^\top\|_{\mathbb{F}}^2 = O_{\mathbb{P}}(n/d + 1/n)$.
2. For any $\mathcal{I} \subset \{1, 2, \dots, d\}$, we have $\max_{\ell \in \mathcal{I}} \sum_{t=1}^n |\hat{u}_{t\ell} - u_{t\ell}|^2 = O_{\mathbb{P}}(\log |\mathcal{I}| + n/d)$.
3. $\|\mathbf{H}^\top \mathbf{H} - \mathbf{I}_K\|_{\mathbb{F}}^2 = O_{\mathbb{P}}(1/n + 1/d)$.
4. $\max_{\ell \in [d]} \|\hat{\mathbf{b}}_\ell - \mathbf{H}\mathbf{b}_\ell\|_2^2 = O_{\mathbb{P}}\{(\log d)/n\}$.

Remark 2. In practice, the number of latent factors K is typically unknown and it is an important issue to determine K in a data-driven way. There have been various methods proposed in the literature to estimate the number K [Bai and Ng, 2002, Lam and Yao, 2012, Ahn and Horenstein, 2013, Fan et al., 2022]. Our theories always work as long as we replace K by any consistent estimator $\hat{K}$, i.e. we only require

$$\mathbb{P}(\hat{K} = K) \rightarrow 1, \text{ as } n \rightarrow \infty.$$

Thus, without loss of generality, we assume the number of factors K is known throughout all the theories developed in this paper. As for the application part, throughout this paper, we utilize the eigenvalue ratio method [Lam and Yao, 2012, Ahn and Horenstein, 2013] to select the number of factors. More specifically, we let $\lambda_k(\mathbf{X}\mathbf{X}^\top)$ denote the eigenvalues of the Gram matrix $\mathbf{X}\mathbf{X}^\top$ and the number of factors is given by

$$\hat{K} = \arg \max_{K \leq \mathcal{K}} \frac{\lambda_K(\mathbf{X}\mathbf{X}^\top)}{\lambda_{K+1}(\mathbf{X}\mathbf{X}^\top)}, \quad (2.3)$$

where $1 \leq \mathcal{K} \leq n$ is a prescribed upper bound for K . $\square$

2.2 Regularized Estimation

Under the high dimensional regime where the dimension d can be much larger than the sample size n , it is often assumed that only a small portion of the predictors contribute to the response variable, which amounts to assuming that the true parameter vector β^* is sparse. Then the regularized estimator for the unknown parameter vectors β^* and γ^* of our factor augmented linear model is defined as follows:

$$(\hat{\beta}_\lambda, \hat{\gamma}) = \arg \min_{\beta \in \mathbb{R}^d, \gamma \in \mathbb{R}^K} \left\{ \frac{1}{2n} \|\mathbf{Y} - \hat{\mathbf{U}}\beta - \hat{\mathbf{F}}\gamma\|_2^2 + \lambda \|\beta\|_1 \right\}, \quad (2.4)$$

where $\lambda > 0$ is a tuning parameter.

We let $\tilde{\mathbf{Y}} = (\mathbf{I}_n - \hat{\mathbf{P}})\mathbf{Y}$ denote the residuals of the response vector $\mathbf{Y}$ after projecting onto the column space of $\hat{\mathbf{F}}$, where $\hat{\mathbf{P}} = n^{-1}\hat{\mathbf{F}}\hat{\mathbf{F}}^\top$ is the corresponding projection matrix. Recall that $\hat{\mathbf{U}} = (\mathbf{I}_n - \hat{\mathbf{P}})\mathbf{X}$. Hence $\hat{\mathbf{F}}^\top \hat{\mathbf{U}} = 0$ and it is straightforward to verify that the solution of (2.4) is equivalent to

$$\begin{aligned} \hat{\beta}_\lambda &= \arg \min_{\beta \in \mathbb{R}^d} \left\{ \frac{1}{2n} \|\tilde{\mathbf{Y}} - \hat{\mathbf{U}}\beta\|_2^2 + \lambda \|\beta\|_1 \right\}, \\ \hat{\gamma} &= (\hat{\mathbf{F}}^\top \hat{\mathbf{F}})^{-1} \hat{\mathbf{F}}^\top \mathbf{Y} = \frac{1}{n} \hat{\mathbf{F}}^\top \mathbf{Y}. \end{aligned}$$

For any subset $\mathcal{S}$ of $\{1, \dots, d\}$, we define the convex cone $\mathcal{C}(\mathcal{S}, 3) = \{\delta \in \mathbb{R}^d : \|\delta_{\mathcal{S}^c}\|_1 \leq 3\|\delta_{\mathcal{S}}\|_1\}$. For simplicity of notation, we write

$$\mathcal{V}_{n,d} = \frac{n}{d} + \sqrt{\frac{\log d}{n}} + \sqrt{\frac{n \log d}{d}}. \quad (2.5)$$

To investigate the consistency property of $(\hat{\beta}_\lambda, \hat{\gamma})$, we impose the following moment condition on the random noise ε .

Assumption 2.5. There exists a positive constant $c_1 < \infty$ such that $\|\varepsilon\|_{\psi_2} \leq c_1$.

Theorem 2.2. Recall $\varphi^* = \gamma^* - \mathbf{B}^\top \beta^* \in \mathbb{R}^K$. Under Assumptions 2.1–2.5, we have

$$\|\hat{\gamma} - \mathbf{H}\gamma^*\|_2 = O_{\mathbb{P}} \left\{ \frac{1}{\sqrt{n}} + \left(\frac{1}{\sqrt{n}} + \frac{1}{\sqrt{d}} \right) \|\varphi^*\|_2 + \|\beta^*\|_1 \left(\sqrt{\frac{\log |\mathcal{S}_*|}{n}} + \frac{1}{\sqrt{d}} \right) \right\},$$

where $\mathcal{S}_\star = \{j \in [d] : \beta_j^\star \neq 0\}$ and $|\mathcal{S}_\star|$ is its cardinality. Furthermore, if $|\mathcal{S}_\star| \left(\frac{\log d}{n} + \frac{1}{d} \right) = o(1)$, then, by taking $\lambda = (\mathcal{I}_0/n) \|\hat{\mathbf{U}}^\top (\tilde{\mathbf{Y}} - \hat{\mathbf{U}} \beta^\star)\|_\infty$ for some constant $\mathcal{I}_0 \geq 2$, we have $\hat{\beta}_\lambda - \beta^\star \in \mathcal{C}(\mathcal{S}_\star, 3)$ and

$$\|\hat{\beta}_\lambda - \beta^\star\|_2 = O_{\mathbb{P}} \left(\sqrt{\frac{|\mathcal{S}_\star| \log d}{n}} + \frac{\mathcal{V}_{n,d} \|\varphi^\star\|_2 \sqrt{|\mathcal{S}_\star|}}{n} \right). \quad (2.6)$$

Remark 3. In most of literature investigating the regularized estimation of sparse linear regression model (1.3), it is commonly assumed that the observed covariate vector $\mathbf{x}$ is a sub-Gaussian random vector with bounded sub-Gaussian norm $\|\mathbf{x}\|_{\psi_2}$. See, for instance, [Loh and Wainwright \[2012\]](#), [Nickl and Van De Geer \[2013\]](#), [van de Geer et al. \[2014\]](#), [Zhang and Cheng \[2017\]](#) and many others. However, such assumption can be unreasonable in the presence of highly correlated covariates. To see this, suppose now both $\mathbf{f}$ and $\mathbf{u}$ are Gaussian random vectors and the underlying $\mathbf{x}$ satisfies the factor model (1.1). Then $\mathbf{x}$ is also a Gaussian random vector with $\text{Cov}(\mathbf{x}) = \mathbf{B}\mathbf{B}^\top + \Sigma$. Under the pervasiveness condition (Assumption 2.2) and Assumption 2.4, it is straightforward to verify that $\|\mathbf{x}\|_{\psi_2} = \sqrt{8/3} \lambda_{\max}(\mathbf{B}\mathbf{B}^\top + \Sigma) \asymp d$, which violates the assumption on bounded sub-Gaussian norm. In contrast, our model can circumvent such issue because we decompose the covariate $\mathbf{x}$ into $(\mathbf{f}, \mathbf{u})$, and we only need impose sub-Gaussian assumption on $(\mathbf{f}, \mathbf{u})$. As the sparse linear regression model serves as a special case to our model, our model serves as a more robust choice to conduct parameter estimation comparing with using linear regression directly, even if the sparse linear regression model is adequate. $\square$

Remark 4. Theorem 2.2 substantially generalize the results in [Fan et al. \[2020a\]](#) with weaker assumptions. First, we did not impose the irrepresentable condition on the design matrix $\mathbf{U}$, only the lower bound on $\Sigma = \text{Cov}(\mathbf{u})$ is required. In addition, although [Fan et al. \[2020a\]](#) also decompose the covariate $\mathbf{x}$ into $(\mathbf{f}, \mathbf{u})$ in order to get precise estimator for $\beta^\star$, they mainly focus on the linear model $Y = \mathbf{x}^\top \beta^\star + \varepsilon$ which corresponds to the special case with $\varphi^\star = 0$ in our results given in Theorem 2.2. $\square$

Remark 5. Our study is different from the related work by [Fan et al. \[2021a\]](#), although they also study one kind of factor augment linear regression model. First of all, they do hypothesis testing for covariance matrix of the idiosyncratic component whereas we focus on testing the adequacy of factor

regression and sparse regression and address also robustness issue. Secondly, their study focuses on panel data and concerning more on prediction rather than the inference. $\square$

2.3 Factor Augmented Robust Linear Regression

In reality, datasets, especially collected from the field of finance, are often contaminated by noises with relatively heavy tails. To resolve such issue, we leverage the adaptive Huber regression to study the parameter of interest in our FARM under the existence of heavy-tailed noise [Sun et al., 2020].

Let $\rho_\omega(\cdot)$ denote the Huber function,

$$\rho_\omega(z) = \begin{cases} z^2/2, & \text{if } |z| \leq \omega, \\ \omega z - \omega^2/2, & \text{if } |z| > \omega, \end{cases}$$

where $\omega > 0$ is the robustification parameter which balances robustness and bias. As an robust version of (2.4), our factor augmented adaptive Huber estimator for $(\beta^\star, \gamma^\star)$ is given by

$$(\hat{\beta}_h, \hat{\gamma}_h) = \arg \min_{\beta \in \mathbb{R}^d, \gamma \in \mathbb{R}^K} \left\{ \frac{1}{n} \sum_{t=1}^n \rho_\omega(y_t - \hat{\mathbf{u}}_t^\top \beta - \hat{\mathbf{f}}_t^\top \gamma) + \lambda \|\beta\|_1 \right\}, \quad (2.7)$$

where $\lambda > 0$ is a tuning parameter and ω depends on sample size, dimenality, and noise level. For simplicity of notation, we write $\hat{\phi}_h = (\hat{\beta}_h^\top, \hat{\gamma}_h^\top)^\top \in \mathbb{R}^{d+K}$ and $\tilde{\phi} = (\beta^{\star\top}, \tilde{\gamma}^\top)^\top \in \mathbb{R}^{d+K}$, where $\tilde{\gamma} = \hat{B}^\top \beta^\star + n^{-1} \hat{F}^\top F \varphi^\star$. The following theorem establishes the statistical consistency of $\hat{\phi}_h$.

Proposition 2.3. *Assume that $\mathbb{E}|\varepsilon|^{1+\vartheta} < \infty$ for some constant $\vartheta > 0$. Let*

$$\omega \asymp \left(\frac{n}{\log d} \right)^{\frac{1}{1+(\vartheta \wedge 1)}} \quad \text{and} \quad \lambda \asymp \left(\frac{\log d}{n} \right)^{\frac{\vartheta \wedge 1}{1+(\vartheta \wedge 1)}}.$$

Furthermore, we assume that $(|\mathcal{S}_\star| + K)(\log d)^{3/2} = o(n)$,

$$\frac{\log n}{n + \sqrt{d}} \|\varphi^\star\|_2 = o(\omega) \quad \text{and} \quad \mathcal{V}_{n,d} \|\varphi^\star\|_2 = O(\omega \log d). \quad (2.8)$$

Then, under Assumptions 2.1–2.4, we have

$$\|\hat{\phi}_h - \tilde{\phi}\|_1 = O_{\mathbb{P}} \left\{ (|\mathcal{S}_\star| + K) \left(\frac{\log d}{n} \right)^{\frac{\vartheta \wedge 1}{1+(\vartheta \wedge 1)}} \right\}.$$

We establish the ℓ_1 -statistical rate for the parameters in model (1.4)[also (1.5)] by only assuming the existence of $(1 + \vartheta)$ -th moment of the noise distribution. Specifically, when $\vartheta \geq 1$, the results reduce to the same rate as the sub-Gaussian assumption of ε . Our result serves as an extension of Sun et al. [2020], who study the robust estimation for high-dimensional linear regression, to a more general setting by incorporating latent factors.

Remark 6. It is worth noting that the statistical errors we obtained throughout section 2 are non-asymptotic, in the sense that the results always hold as long as n is greater than some fixed constant. As our learned covariates contain statistical errors, novel analysis is required for downstream statistical estimation and inference under both scenarios with light and heavy-tailed noises

3 Is Factor Regression Model Adequate?

The latent factor regression is widely applied in many fields as an efficient dimension reduction method. A natural question arises is whether the model is adequate and FARM (1.4) serves naturally as the alternative model. To be more specific, we consider testing the hypotheses

$$H_0 : \beta^* = 0 \text{ versus } H_1 : \beta^* \neq 0 \quad (3.1)$$

in FARM (1.4). As the penalized least-squares estimator $\hat{\beta}_\lambda$ is used for estimating β^* , it creates biases and make it difficulty for inferences. Thus, we first introduce a de-biased version of $\hat{\beta}_\lambda$ given in (2.4).

3.1 Bias Correction

We begin with the construction of bias-corrected estimator for β^* following similar idea of Zhang and Zhang [2014], van de Geer et al. [2014] and Javanmard and Montanari [2014]. Specifically, let $\hat{\Theta} \in \mathbb{R}^{d \times d}$ be an approximation for the inverse of the Gram matrix $\tilde{\Sigma} = n^{-1} \hat{U}^\top \hat{U}$, the de-biased estimator for β^* is then defined as

$$\tilde{\beta}_\lambda = \hat{\beta}_\lambda + \frac{1}{n} \hat{\Theta} \hat{U}^\top (\mathbf{Y} - \hat{U} \hat{\beta}_\lambda). \quad (3.2)$$

The rationale behind such construction is that we are able to decompose estimation error as

$$\tilde{\beta}_\lambda - \beta^\star = \frac{1}{n} \hat{\Theta} \hat{U}^\top \mathcal{E} + \frac{1}{n} \hat{\Theta} \hat{U}^\top F \varphi^\star + (I_d - \hat{\Theta} \tilde{\Sigma})(\hat{\beta}_\lambda - \beta^\star), \quad (3.3)$$

after we expand $\mathbf{Y}$ according to (2.2) and replace $\mathbf{X}$ by $\mathbf{X} = \hat{F} \hat{B} + \hat{U}$. The first term on the right hand side of (3.3) quantifies the uncertainty of our estimator $\tilde{\beta}_\lambda$ and the last two terms are biases which will be shown to be of smaller order.

One observes that constructing the de-biased estimator $\tilde{\beta}_\lambda$ given above requires an estimator $\hat{\Theta}$. There are many methods for estimating such precision matrix, for example, the node-wise regression proposed in Zhang and Zhang [2014] and van de Geer et al. [2014], and the CLIME-type estimator given in Cai et al. [2011], Javanmard and Montanari [2014] and Avella-Medina et al. [2018]. In our work, we do not restrict $\hat{\Theta}$ to be any specific one, but require to satisfy the following general conditions.

Assumption 3.1. Let $\Theta = \Sigma^{-1}$ with Σ defined in Assumption 2.3. There exist positive $\Lambda_{\max}$ and Δ_∞ such that

$$\|I_d - \hat{\Theta} \tilde{\Sigma}\|_{\max} = O_{\mathbb{P}}(\Lambda_{\max}) \quad \text{and} \quad \|\hat{\Theta} - \Theta\|_\infty = O_{\mathbb{P}}(\Delta_\infty).$$

Without loss of generality, here we assume that $\Delta_\infty \leq \|\Theta\|_\infty$.

Remark 7. To give a concrete example, under the mild conditions therein, Assumption 3.1 is satisfied with

$$\Lambda_{\max} = O\left(\sqrt{\frac{\log d}{n}} + \frac{1}{d}\right) \quad \text{and} \quad \Delta_\infty = O\left(\max_{j \in [d]} |\mathcal{S}_j| \sqrt{\frac{\log d}{n}} + \frac{1}{d}\right),$$

by using node-wise regression [Zhang and Zhang, 2014, van de Geer et al., 2014], where $|\mathcal{S}_j| = \sum_{k=1}^d \mathbb{I}\{\Theta_{jk} \neq 0\}$ quantifies the sparsity of j -th column of the precision matrix Θ for each $1 \leq j \leq d$. In Appendix ??, we will provide a detailed analysis on estimating $\tilde{\Sigma}^{-1}$ via node-wise regression and establish precise theoretical upper bounds for the statistical rates given in Assumption 3.1. $\square$

3.2 Gaussian Approximation

The goal of this section is to derive the asymptotic distribution of $\|\tilde{\beta}_\lambda - \beta^\star\|_\infty$ in the high dimensional setting. To this end, we apply the Gaussian approximation result given in Chernozhukov et al. [2013, 2017, 2020] for high dimensional random vectors. More specifically, we let $\mathbf{Z} = (Z_1, \dots, Z_d)^\top \in \mathbb{R}^d$ be a zero-mean Gaussian random vector with the same covariance matrix as that of $n^{-1/2}\Theta\mathbf{U}^\top\mathcal{E}$, that is,

$$\text{Cov}(\mathbf{Z}) = \text{Cov}\left(\frac{1}{\sqrt{n}}\Theta\mathbf{U}^\top\mathcal{E}\right) = \sigma^2\Theta. \quad (3.4)$$

We next present the theoretical results on Gaussian approximation of our test statistics under some mild conditions.

Theorem 3.1. *Recall $\varphi^\star = \gamma^\star - \mathbf{B}^\top\beta^\star \in \mathbb{R}^K$. We assume that $(\log d)^5/n \rightarrow 0$,*

$$(\Lambda_{\max}|\mathcal{S}_\star| + \Delta_\infty) \log d \rightarrow 0 \text{ and } \left(\mathcal{V}_{n,d}\|\varphi^\star\|_2 + \sqrt{\frac{n}{d}} + \sqrt{\log d}\right) \|\Theta\|_\infty \sqrt{\frac{\log d}{n}} \rightarrow 0, \quad (3.5)$$

with $\mathcal{V}_{n,d}$. Then under Assumption 3.1, we have

$$\sup_{x>0} \left| \mathbb{P}\left(\sqrt{n}\|\tilde{\beta}_\lambda - \beta^\star\|_\infty \leq x\right) - \mathbb{P}(\|\mathbf{Z}\|_\infty \leq x) \right| \rightarrow 0.$$

For any $\alpha \in (0, 1)$, let $c_{1-\alpha}$ denote the $(1-\alpha)$ -th quantile of the distribution of $\|\mathbf{Z}\|_\infty$. Theorem 3.1 leads to an approximately level α test for (3.1) as follows:

$$\psi_{\infty,\alpha} = \mathbb{I}\left\{\sqrt{n}\|\tilde{\beta}_\lambda\|_\infty > c_{1-\alpha}\right\}. \quad (3.6)$$

3.3 Gaussian multiplier bootstrap

The critical value $c_{1-\alpha}$ depends on the unknown σ^2 and Θ , which can be estimated by the following Gaussian multiplier bootstrap.

1. Generate i.i.d. random variables $\xi_1, \dots, \xi_n \sim N(0, 1)$ and compute

$$\hat{L} = \frac{1}{\sqrt{n}} \|\hat{\Theta}\hat{\mathbf{U}}^\top\boldsymbol{\xi}\|_\infty, \text{ where } \boldsymbol{\xi} = (\xi_1, \xi_2, \dots, \xi_n)^\top.$$

2. Repeat the first step independently for B times and obtain $\hat{L}_1, \dots, \hat{L}_B$. Estimate the critical value $c_{1-\alpha}$ via $1 - \alpha$ quantile of the empirical distribution of the bootstrap statistics:

$$\hat{c}_{1-\alpha} = \inf\{t \geq 0 : H_B(t) \geq 1 - \alpha\}, \text{ where } H_B(t) = \frac{1}{B} \sum_{b=1}^B \mathbb{I}\{\hat{L}_b \leq t\}.$$

Reject the null hypothesis H_0 when $\sqrt{n}\|\tilde{\beta}_\lambda\|_\infty/\hat{\sigma} > \hat{c}_{1-\alpha}$, for a given consistent estimator $\hat{\sigma}$ of σ . To validate the procedure, we need some additional conditions on $\hat{\Theta}$ and $\hat{\sigma}$.

Assumption 3.2. There exists a $\Delta_{\max} > 0$ such that $\|\hat{\Theta} - \Theta\|_{\max} = O_{\mathbb{P}}(\Delta_{\max})$.

Assumption 3.3. There exists a $0 < \Delta_\sigma \leq 1$ such that $|\hat{\sigma}/\sigma - 1| = O_{\mathbb{P}}(\Delta_\sigma)$.

Remark 8. The estimation of σ^2 for high dimensional linear regression has been studied in the literature. For example, [Fan et al. \[2012\]](#) proposed refitted cross-validation to construct a consistent estimator with clearly quantified uncertainty of $\hat{\sigma}$ in ultra-high dimension. In addition, [Sun and Zhang \[2012\]](#) and [Yu and Bien \[2019\]](#) derived scaled-Lasso and organic Lasso respectively for estimating σ . Like our case of estimating Θ , we also do not restrict estimating σ by any fixed method mentioned above, our theory works as long as the general condition of Assumption 3.3 holds. $\square$

Let $\mathbb{P}^*(\cdot) = \mathbb{P}(\cdot | \mathbf{X}, \mathbf{Y})$ denote the conditional probability. In the following theorem, we establish the validity of the proposed bootstrap procedure.

Theorem 3.2. *Let Assumptions 3.1–3.3 hold. Assume that*

$$\Lambda_{\max}\|\Theta\|_\infty + \Delta_{\max} + \Delta_\sigma = o\left(\frac{1}{\log d}\right). \quad (3.7)$$

Then, under conditions of Theorem 3.1, we have

$$\sup_{x>0} \left| \mathbb{P}\left(\sqrt{n}\|\tilde{\beta}_\lambda - \beta^*\|_\infty \leq x\right) - \mathbb{P}^*\left(\hat{L} \leq x\right) \right| \xrightarrow{\mathbb{P}} 0.$$

Remark 9. Following the same de-biasing procedure as given in (3.2), we are also able to construct entrywise [\[Javanmard and Montanari, 2014\]](#) and groupwise [\[Zhang and Cheng, 2017, Dezeure et al.,](#)

2017] simultaneous confidence intervals for β^* . For each $1 \leq j \leq d$, a $(1 - \alpha)$ -confidence interval for β_j^* is given by

$$\mathcal{CI}_\alpha(\beta_j^*) = \left\{ \tilde{\beta}_{j,\lambda} - \hat{\sigma} z_{1-\alpha/2} \sqrt{\frac{\hat{\Theta}_{jj}}{n}}, \tilde{\beta}_{j,\lambda} + \hat{\sigma} z_{1-\alpha/2} \sqrt{\frac{\hat{\Theta}_{jj}}{n}} \right\},$$

where $z_{1-\alpha/2}$ is the $(1 - \alpha/2)$ -th quantile of standard normal distribution. By choosing this cut-off value, we obtain a tighter confidence interval comparing with using $c_{1-\alpha}$. For simultaneous groupwise inference of β^* , let G be a subset of $\{1, \dots, d\}$ of interest and consider testing the hypotheses

$$H_{0,G} : \beta_j^* = \beta_j^\circ \text{ for all } j \in G \text{ versus } H_{1,G} : \beta_j^* \neq \beta_j^\circ \text{ for some } j \in G.$$

In particular, when $\beta_j^\circ = 0$ for all $j \in G$, this reduces to testing the significance of a group of parameters. We obtain that the asymptotic distribution of $\max_{j \in G} \sqrt{n} |\tilde{\beta}_{j,\lambda} - \beta_j^*|$ converges to the distribution of $\max_{j \in G} |Z_j|$ by leveraging the Gaussian approximation. The remaining steps follow directly by conducting the Gaussian multiplier bootstrap. $\square$

4 Is Sparse Linear Model Adequate?

Sparse linear regression, which serves as the backbone of high dimensional statistics, has been widely applied in many areas of science, engineering, and social sciences. However, its adequacy has never been validated. This section focuses on testing the adequacy of the sparse linear model.

4.1 Main Results

As mentioned in introduction, the proposed model (1.5) contains the sparse linear regression model as a special case. Thus, we consider testing the hypotheses

$$H_0 : Y = \mathbf{x}^\top \beta^* + \varepsilon \text{ versus } H_1 : Y = \mathbf{f}^\top \varphi^* + \mathbf{x}^\top \beta^* + \varepsilon, \quad (4.1)$$

which is equivalent to test whether $\varphi^* = \gamma^* - \mathbf{B}^\top \beta^* = 0$. Since $\mathbf{B}$ is an unknown dense matrix, simultaneously testing this linear equation will suffer from the curse of dimensionality.

On the other hand, for any set $\mathcal{S} \subset [d]$ with $\mathcal{S}_\star \subset \mathcal{S}$, we have $\mathbf{B}_\mathcal{S}^\top \boldsymbol{\beta}_\mathcal{S}^\star = \mathbf{B}^\top \boldsymbol{\beta}^\star$. Hence, it suffices to compare the following two linear models in reduced dimension:

$$H_0 : Y = \mathbf{x}_\mathcal{S}^\top \boldsymbol{\beta}_\mathcal{S}^\star + \varepsilon \text{ versus } H_1 : Y = \mathbf{f}^\top \boldsymbol{\varphi}^\star + \mathbf{x}_\mathcal{S}^\top \boldsymbol{\beta}_\mathcal{S}^\star + \varepsilon. \quad (4.2)$$

This hinges applying a sure screening method to reduce the dimensionality. There exist several methods which lead to the sure screening property. Among those, the commonly used one is the marginal screening method [Fan and Lv, 2008, Fan and Song, 2010, Zhu et al., 2011, Li et al., 2012, Liu et al., 2014, Barut et al., 2016, Chu et al., 2016, Wang and Leng, 2016].

We propose an ANOVA-type test for (4.1) with two stages. In the first stage, the data set is split into two data sets $(\mathbf{Y}^{(1)}, \mathbf{X}^{(1)})$ and $(\mathbf{Y}^{(2)}, \mathbf{X}^{(2)})$, with sample sizes m and $n - m$, respectively. We use $(\mathbf{Y}^{(1)}, \mathbf{X}^{(1)})$ to screen variables. Let $\hat{\mathcal{S}}_1$ denote the set of variables selected. In the second stage, we leverage the selected $\hat{\mathcal{S}}_1$ and remaining data $(\mathbf{Y}^{(2)}, \mathbf{X}^{(2)})$ to perform hypothesis testing based on the ANOVA-type test statistic for low-dimensional model (4.2) with $\mathcal{S}$ replaced by $\hat{\mathcal{S}}_1$. As the first step is based on marginal screening and is relatively crude, the sample size m is relatively small in comparing with the second step. We impose a general assumption on the set $\hat{\mathcal{S}}_1$.

Assumption 4.1 (Sure screening property). There exists an $s_n > 0$ such that

$$\mathbb{P} \left(|\hat{\mathcal{S}}_1| \leq s_n \text{ and } \mathcal{S}_\star \subset \hat{\mathcal{S}}_1 \right) \rightarrow 1, \text{ as } n \rightarrow \infty.$$

A simple procedure that satisfies the above assumption is the follow factor-adjusted marginal screening based on the data $(\mathbf{Y}^{(1)}, \mathbf{X}^{(1)})$.

1. Estimation. Compute the latent factor estimator $\hat{\mathbf{F}}^{(1)}$, idiosyncratic component $\hat{\mathbf{U}}^{(1)}$ based on $\mathbf{X}^{(1)}$, and $\tilde{\mathbf{Y}}^{(1)} = (\mathbf{I}_m - \hat{\mathbf{F}}^{(1)}(\hat{\mathbf{F}}^{(1)\top} \hat{\mathbf{F}}^{(1)})^{-1} \hat{\mathbf{F}}^{(1)\top}) \mathbf{Y}^{(1)}$, the residual after factor regression.
2. Marginal regression. Compute the least square estimate $\hat{\beta}_{\ell, M} = \hat{\mathbf{U}}_\ell^{(1)\top} \tilde{\mathbf{Y}}^{(1)} / (\hat{\mathbf{U}}_\ell^{(1)\top} \hat{\mathbf{U}}_\ell^{(1)})$ for each $1 \leq \ell \leq d$.
3. Screening. Let $\hat{\mathcal{S}}_1 := \hat{\mathcal{S}}_\phi = \{\ell \in [d] : |\hat{\beta}_{\ell, M}| > \phi\}$ for some prescribed $\phi > 0$.

Here $\hat{\mathbf{U}}_\ell^{(1)} \in \mathbb{R}^d$ stands for the ℓ -th column of the matrix $\hat{\mathbf{U}}^{(1)}$. We next provide a sufficient condition for the Assumption 4.1 to hold.

Proposition 4.1. Assume that $m = o(d \log d)$ and

$$\mathcal{O}\left(\frac{\|\psi^*\|_2}{d} + \|\beta^*\|_2 \sqrt{\frac{\log d}{m}}\right) \leq \phi \leq \mathcal{O}\left(\min_{\ell \in [d]} \beta_{\ell, M}^*\right), \quad (4.3)$$

where $\beta_{\ell, M}^* := \Sigma_\ell^\top \beta^* / \Sigma_{\ell\ell}$. Here Σ_ℓ denotes the ℓ -th column of Σ . Then, under the Assumptions 2.1–2.5, we have

$$\mathbb{P}\left(\mathcal{S}_\star \subset \hat{\mathcal{S}}_\phi\right) \rightarrow 1, \text{ as } m \rightarrow \infty.$$

Furthermore, we assume that $\min_{\ell \in \mathcal{S}_\star} |\beta_{\ell, M}^*| \geq c_\star m^{-\kappa}$ for some positive constant $\kappa < 1/2$. Then for any $\phi = c_\diamond m^{-\kappa}$ with $c_\diamond \leq c_\star / (1 + \bar{c})$, we have

$$\mathbb{P}\left\{|\hat{\mathcal{S}}_\phi| \leq \frac{c_\diamond^2 m^{2\kappa} \|\Sigma \beta^*\|_2^2}{\lambda_{\min}^2(\Sigma)(1 - \bar{c})^2}\right\} \rightarrow 1 \text{ as } m \rightarrow \infty.$$

Remark 10. From the conclusion of Proposition 4.1, we obtain sure screening property by using our first data set with sample size $m = n^\alpha$ for some $\alpha < 1$ as long as the signal satisfies $\min_{\ell \in \mathcal{S}_\star} |\beta_{\ell, M}^*| \geq c_\star m^{-\kappa}$. Thus, the size of the remaining data set for constructing the test statistic in our second step is $n - n^\alpha \approx n$. It is worth to note that this does not lose any efficiency in terms of the asymptotic power in our hypothesis test when n goes to infinity. $\square$

Remark 11. Fan et al. [2020a] proposed a similar sure screening estimator which is a special case of our Proposition 4.1 with $\varphi^* = \gamma^* - B^\top \beta^* = 0$. Moreover, we also provide an upper bound for the number of selected variables whereas Fan et al. [2020a] only provided a sufficient condition for the sure screening property. $\square$

Next, we proceed to the second stage of our hypothesis testing. In this step, we construct an ANOVA test statistic for (4.2) with $\mathcal{S}$ replaced by $\hat{\mathcal{S}}_1$, which is given by

$$Q_n^{(2)} = \left\| \left(I_{n-m} - P_{X_{\hat{\mathcal{S}}_1}^{(2)}} \right) Y^{(2)} \right\|_2^2 - \left\| \left(I_{n-m} - P_{\hat{F}^{(2)}} - P_{\hat{U}_{\hat{\mathcal{S}}_1}^{(2)}} \right) Y^{(2)} \right\|_2^2. \quad (4.4)$$

We then summarize our results on the asymptotic behaviors of $Q_n^{(2)}$ in the following Theorem 4.2.

Theorem 4.2. *Let Assumptions 2.1–2.5 and Assumption 4.1 hold with*

$$s_n \left(\frac{\log d}{n} + \frac{1}{d} \right) \rightarrow 0 \text{ and } \Delta_\sigma \rightarrow 0.$$

We obtain

$$\sup_{x>0} |\mathbb{P}(Q_n^{(2)} \leq x\hat{\sigma}^2 | H_0) - \mathbb{P}(\chi_K^2 \leq x)| \rightarrow 0, \text{ as } n \rightarrow \infty.$$

Theorem 4.2 yields a level α test for (4.1), with critical region $\{Q_n^{(2)} > \hat{\sigma}^2 \chi_{K,1-\alpha}^2\}$, where $\chi_{K,1-\alpha}^2$ is the $(1 - \alpha)$ -th quantile of χ_K^2 -distribution.

Remark 12. Under stronger conditions such as irrepresentable condition [Zhao and Yu, 2006] or RIP condition [Candes and Tao, 2007], the $\hat{S}$ achieved by certain explicit regularization [Zhao and Yu, 2006, Fan and Lv, 2011, Shi et al., 2019, Fan et al., 2020a] or implicit regularization accompanied with early stopping and signal truncation [Zhao et al., 2019, Fan et al., 2021b] enjoys variable selection consistency $\mathbb{P}(\hat{S} = S_\star) \rightarrow 1$. In this scenario, we take the test statistic as

$$Q_n = \left\| \left(P_{\hat{F}} + P_{\hat{U}_{\hat{S}}} - P_{X_{\hat{S}}} \right) Y \right\|_2^2$$

without using sample splitting. Under Assumptions 2.1–2.5, we obtain

$$\sup_{x>0} |\mathbb{P}(Q_n \leq x\hat{\sigma}^2 | H_0) - \mathbb{P}(\chi_K^2 \leq x)| \rightarrow 0, \quad (4.5)$$

by following similar proof idea with Theorem 4.2. □

We now present the power of the test statistic (4.4).

Theorem 4.3. *Define*

$$\mathcal{D}(\alpha, \theta) = \left\{ \varphi \in \mathbb{R}^K : \frac{n \|\varphi\|^2}{1 + K s_n \|\mathbf{B}\|_{\max}^2 / \lambda_{\min}(\Sigma)} \geq \sigma^2 (2 + \delta) (\chi_{K,1-\alpha}^2 + \chi_{K,1-\theta}^2) \right\},$$

where $\delta > 0$ is some constant, s_n is the size of selected set from the first stage and K is the number of factors. In addition, parameter θ is a threshold such that for any $\varphi^* \in \mathcal{D}(\alpha, \theta)$, the power of the test is larger than $1 - \theta$. To be more specific, we assume that

$$\|\varphi^*\|_2 \left(\sqrt{n/d} + 1/\sqrt{n} \right) \rightarrow 0. \quad (4.6)$$

Then, under the conditions of Theorem 4.2, we have

$$\inf_{\varphi^* \in \mathcal{D}(\alpha, \theta)} \mathbb{P}(\psi_\alpha = 1 | H_1) \geq 1 - \theta,$$

where $\psi_\alpha = \mathbb{I}_{\{Q_n^{(2)} > \hat{\sigma}^2 \chi_{K, 1-\alpha}^2\}}$.

Remark 13. Dataset with multiple types are now frequently collected for a common set of experimental subjects. This new data structure is also called multimodal data. It is worth to mention, the above hypothesis test can be further extended to test the adequacy of multi-modal sparse linear regression model [Li and Li, 2021]. Interested readers are referred to Appendix F.4 for more details. $\square$

5 Numerical Studies

5.1 Accuracy of Estimation

For data generation, we let number of factors $K = 2$, dimension of covariate $d = 1000$, $\gamma^* = (0.5, 0.5)$, the first $s = 3$ entries of β^* be 0.5 and remaining $d - s$ entries be 0. Throughout this subsection, we generate every entry of F, U from the standard Gaussian distribution and let every entry of B be generated from the uniform distribution $\text{Unif}(-1, 1)$. We choose the noise distribution of ε given in model (2.1) from (i) standard Gaussian, (ii) uniform, and (iii) t_3 distribution respectively.

Distributions (i) and (ii) have sub-Gaussian tails. For these two cases, we select sample size n so that $s\sqrt{\log d/n}$ takes uniform grids in $[0.15, 0.5]$. Then we generate n response variables from model (2.1) and estimate our parameters via (2.4). The results are shown as the red lines in Figure 1. They lend further support to our theoretical findings given in section 2 as the statistical rates there are upper bounded by $O(s\sqrt{\log d/n})$. Moreover, we also show the estimation results by using Lasso directly on measurements (X, Y) . Results are shown as the blue lines given in the first two figures in Figure 1. Using Lasso directly on (X, Y) leads to much worse results due in part to the inadequacy of the model. In addition, as shown in Fan et al. [2020a], even when the sparse regression model is correct, we still have better estimation accuracy using factor adjusted regression.

Distribution (iii) has only the bounded second moment, but no third moment. Likewise, we select corresponding number of observations n so that $(s+K)\sqrt{\log d/n}$ takes uniform grids in $[0.4, 0.7]$. The

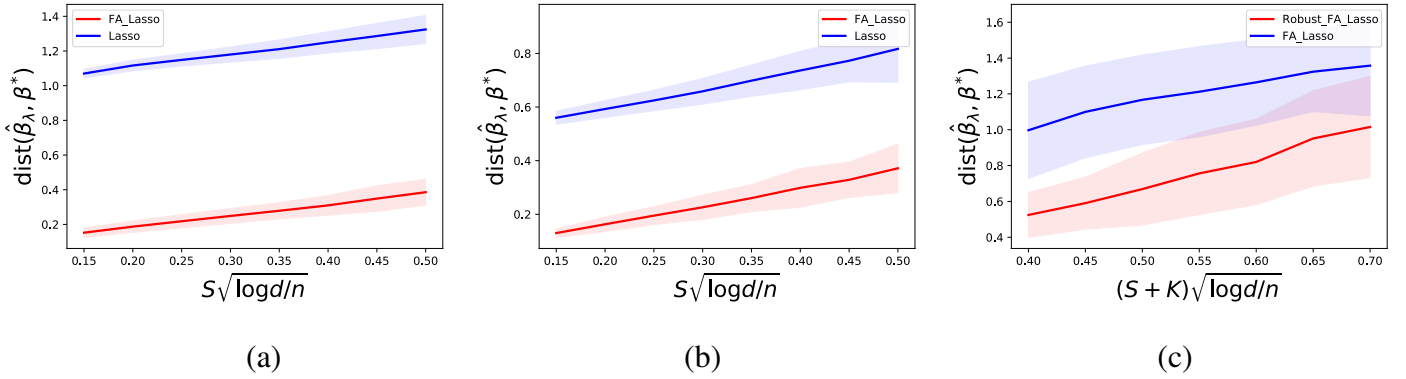


Figure 1: Accuracy for $\hat{\beta}_\lambda$ with $\text{dist}(\hat{\beta}_\lambda, \beta^*) := \|\hat{\beta}_\lambda - \beta^*\|_1$ based on 500 replications. The light color regions indicate the standard errors across the simulation. Figure (a) and (b) depict the estimation results of model (1.4) with noise ϵ following the standard Gaussian and uniform distributions respectively. In (a) and (b), the red lines denote the estimation results using the method (2.4) (labeled as FA_Lasso in the figure) and the blue lines represent the results using Lasso with data $(\mathbf{X}, \mathbf{Y})$. In Figure (c), the noise ϵ follows t_3 -distribution. The red line in (c) represents the result of robust factor adjusted regression (Robust_FA_Lasso) via adaptive Huber estimation together with ℓ_1 -penalty given in (2.7) and the blue line represents the result achieved by using FA_Lasso.

reduced sample sizes help reduce the computation cost on the regularized adaptive Huber estimation using cross-validation to choose the parameter ω . We compare the results for the robust estimator (2.7) with that of the factor adjusted regression (2.4). The results are shown as the red and blue lines in part (c) of Figure 1 respectively. They provide stark evidence that it is necessary to conduct the robust version of factor adjusted regression (2.7) when noises have heavy tails.

5.2 Adequacy of Factor Regression

Data Generation Processes. We choose $n = 200$, $K = 2$ and d either 200 or 500 and the matrix $\mathbf{X} = \mathbf{F}\mathbf{B}^\top + \mathbf{U}$ using the following two models with entries of $\mathbf{B}$ generated from $\text{Unif}(-1, 1)$.

1. We generate every row of $\mathbf{F} \in \mathbb{R}^{n \times K}$, $\mathbf{U} \in \mathbb{R}^{n \times d}$ from $\mathcal{N}(\mathbf{0}, \mathbf{I}_K)$ and $\mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ respectively.
2. We let the t -th row $\mathbf{f}_t \in \mathbb{R}^K$ of $\mathbf{F} \in \mathbb{R}^{n \times K}$ follow $\mathbf{f}_t = \Phi \mathbf{f}_{t-1} + \boldsymbol{\xi}_t$ where $\Phi \in \mathbb{R}^{K \times K}$ with

$\Phi_{i,j} = 0.5^{|i-j|+1}$, $i, j \in [K]$. In addition, $\{\xi_t\}_{t \geq 1}$ are drawn independently from $N(\mathbf{0}, \mathbf{I}_K)$. We generate every row of $\mathbf{U}$ from $N(\mathbf{0}, \Sigma)$ where $\Sigma_{i,j} = 0.6^{|i-j|}$, $i, j \in [d]$.

The response vector follows $\mathbf{Y} = \mathbf{F}\boldsymbol{\gamma}^* + \mathbf{U}\boldsymbol{\beta}^* + \mathcal{E}$ in (4.1) with every entry of $\mathcal{E} \in \mathbb{R}^n$ being generated independently from either from $N(0, 0.5^2)$ or uniform distribution $\text{Unif}(-\sqrt{3}/2, \sqrt{3}/2)$. We set $\boldsymbol{\gamma}^* = (0.5, 0.5)$ and $\boldsymbol{\beta}^* = (w, w, w, 0, \dots, 0)$, where $w \geq 0$. When $w = 0$, the null hypothesis $\mathbf{Y} = \mathbf{F}\boldsymbol{\gamma}^* + \mathcal{E}$ holds and the simulation results correspond to the size of the test. Otherwise, they correspond to the power of the test.

For $n = 200, K = 2, d \in \{200, 500\}$ and all $w \in \{0, 0.05, 0.10, 0.15, 0.20\}$, we generate the data from each model and compute the testing results based on procedures in §3 and 2000 simulations with $\alpha = 0.05$. For every replication, we conduct bootstrap 2000 times to compute the critical value $\hat{c}_{1-\alpha}$ given in §3. The results are depicted in the Table 1. The column named Gaussian(i), $i \in \{1, 2\}$ represents the simulation results under model i with Gaussian noise. Similar labels applied to the uniform noise distribution.

		Gaussian (1)	Gaussian (2)	Uniform (1)	Uniform (2)
$p = 200$	$w = 0$	0.044	0.047	0.046	0.048
	$w = 0.05$	0.067	0.119	0.065	0.108
	$w = 0.10$	0.326	0.714	0.311	0.653
	$w = 0.15$	0.859	0.989	0.854	0.984
	$w = 0.20$	0.998	1.000	0.996	1.000
$p = 500$	$w = 0$	0.043	0.040	0.048	0.436
	$w = 0.05$	0.067	0.080	0.059	0.071
	$w = 0.10$	0.253	0.632	0.237	0.563
	$w = 0.15$	0.787	0.974	0.780	0.962
	$w = 0.20$	0.993	1.000	0.987	1.000

Table 1: Simulation results of section 3 under different regimes.

Table 1 reveals that our test gives approximately the right size (subject to simulation error; see the rows with $w = 0$). This is consistent with our theoretical findings given in section 3. In addition, when $0 < w < 0.2$, the power of our test increases rapidly to 1 which reveals the efficiency of our test

statistic.

5.3 Adequacy of Sparse Regression

This subsection provides finite-sample validations for the results in section 4. We take the number of data used for screening $m = \lceil n^{0.8} \rceil$, use Iterative Sure Independence Screening method [Fan and Lv, 2008, Saldana and Feng, 2018, Zhang et al., 2019] to select $\hat{\mathcal{S}}_1$ and apply the refitted cross-validation [Fan et al., 2012] to estimate σ^2 . The size and the power of the test are computed based on 2000 simulations.

Data Generation Processes. We let $n = 250$, $K = 3$ and d be either 250 or 600. The noises ε are i.i.d from $N(0, 0.5^2)$ or $\text{Unif}(-\sqrt{3}/2, \sqrt{3}/2)$. The covariate $\mathbf{X} \in \mathbb{R}^{n \times d}$ follows the factor model $\mathbf{X} = \mathbf{F}\mathbf{B}^\top + \mathbf{U}$. We generate $\mathbf{F}$, $\mathbf{U}$ and $\mathbf{B}$ in the same way as those in section 5.2. In addition, the response variable follows $\mathbf{Y} = \mathbf{F}\boldsymbol{\varphi}^* + \mathbf{X}\boldsymbol{\beta}^* + \boldsymbol{\varepsilon}$ in (4.1) with $\boldsymbol{\beta}^* = (0.8, 0.8, 0.8, 0.8, 0, \dots, 0)$ and $\boldsymbol{\varphi}^* = v \cdot \mathbf{1}_{K \times 1}$ for several different values of $v \geq 0$. The case $v = 0$ corresponds to the null hypothesis and it is designed to test the validity of the size.

Results. For $n = 250$, $K = 3$, $d \in \{250, 600\}$ and $v \in \{0, 0.04, 0.08, 0.12, 0.16\}$, we implement the proposed method for every model in section 5.2. The simulation results are depicted in Table 2. The column named Gaussian (or uniform) (i), $i \in \{1, 2\}$ represents the results under model i with Gaussian (or uniform) noise mentioned in section 5.2. When $v = 0$, the null hypothesis holds, our Type-I error is approximately 0.05 which matches with the theoretical value. In addition, when we increase the size of v from $v = 0.04$ to $v = 0.16$, the power of our test statistic increases sharply to 1, which reveals its efficiency.

We next discuss the necessity of using sample splitting. Suppose we do not split samples and use the whole dataset to do sure screening and construct the test statistic. This will result in the high correlation between the selected set $\hat{\mathcal{S}}$ and covariates when $\hat{\mathcal{S}}$ is not a consistent estimator of $\mathcal{S}_*$. In this case, the asymptotic behavior of our test statistic is hard to capture. To demonstrate this point, we simulate the null distribution of the test statistic constructed without using sample splitting and compare it with the asymptotic distribution (χ_K^2) via the quantile-quantile plot in Figure 7 in appendix

		Gaussian (1)	Gaussian (2)	Uniform (1)	Uniform (2)
$p = 250$	$v = 0$	0.051	0.054	0.056	0.053
	$v = 0.04$	0.215	0.278	0.233	0.286
	$v = 0.08$	0.659	0.740	0.655	0.750
	$v = 0.12$	0.965	0.993	0.965	0.996
	$v = 0.16$	1.000	1.000	1.000	1.000
$p = 600$	$v = 0$	0.051	0.052	0.050	0.052
	$v = 0.04$	0.208	0.362	0.197	0.353
	$v = 0.08$	0.624	0.802	0.604	0.785
	$v = 0.12$	0.941	0.994	0.934	0.999
	$v = 0.16$	1.000	1.000	0.999	1.000

Table 2: Simulation results of section 4 under different regimes.

§B.2. Figure 7 reveals that the test statistic constructed without using sample splitting has heavier right tail than that of the χ_K^2 distribution. The sizes of the test are much larger than the results in Table 2 when $v = 0$.

We summarize the numerical results as follows. In terms of statistical estimation, our estimated parameters of FARM behave much better than those estimated parameters via sparse linear regression model due to mis-specification. As for prediction, we also conduct additional simulations on comparing FARM with the latent factor regression and sparse linear regression model. Interested readers are referred to §B.1 for more details. For high-dimensional inference, as illustrated in §5.2 and §5.3, when the null hypotheses hold, the size of the test is well-controlled. On the other hand, when the null hypotheses do not hold, the power of our test statistics grow rapidly to 1 even for weak signals.

5.4 Empirical Applications

In this section, we use a macroeconomic dataset named FRED-MD [McCracken and Ng, 2016] to illustrate the performance of our factor augmented regression model (FARM) and investigate whether the latent factor regression model and sparse linear model are adequate.

There are 134 monthly U.S. macroeconomic variables in this dataset. As they measure certain

aspects of economic health, these variables are driven by latent factors and hence correlated. They can be well explained by a few principal components. In our study, we pick out two variables named 'HOUSTNE' and 'GS5' as our responses respectively and let the remaining variables be the covariates. Here 'HOUSTNE' represents the housing starts in the northeast region. Studying the number of housing starts helps one to understand the residents' life condition and economic environment. 'GS5' is correlated with many important variables such as interests rates, inflation and economic growth. It is an important indicator on the financial condition and economics environment of a country.

There exist significant structural breaks for many variables around the year of financial crisis in 2008 which makes our data non-stationary even after performing the suggested transformations. Thus, we analyze the dataset in two separate time periods independently. Specifically, we study the monthly data collected from February 1992 to October 2007 and from August 2010 to February 2020 respectively after examining the missingness and stationarity of the data.

We next compare the performance of our proposed FARM against several benchmarks presented in a few related references which study the same or similar datasets. In specific, we compare the forecasting results of FARM with Lasso (sparse linear regression), PCR (latent factor regression), Ridge (Ridge regression), El-Net (Elastic Net) used in [Coulombe et al. \[2021b\]](#), [Smeekes and Wijler \[2018\]](#), [Hall et al. \[2018\]](#), RF (Random Forest) used in [Goulet Coulombe \[2020\]](#), [Coulombe et al. \[2021a\]](#), [Bianchi et al. \[2021\]](#) and FarmSelect (Factor adjusted Lasso) used in [Fan et al. \[2020a\]](#). For every given time period and model, we perform the prediction by using the moving window approach with window size 90 months. Indexing the panel data from 1 for each of the two time periods, for all $t > 90$, we use the 90 previous measurements $\{(\mathbf{x}_{t-90}, Y_{t-90}), \dots, (\mathbf{x}_{t-1}, Y_{t-1})\}$ to train a model (FARM, sparse linear regression model, latent factor regression model, ridge regression, elastic net, factor adjusted Lasso, random forest), and output a prediction $\hat{Y}_t$ as well as the in-sample average $\bar{Y}_t := \frac{1}{90} \sum_{i=t-90}^{t-1} Y_i$ (the detailed implementations for different methods are presented in the appendix §B.3). We measure the prediction accuracy by using out-of-sample R^2 :

$$R^2 = 1 - \frac{\sum_{t=91}^T (Y_t - \hat{Y}_t)^2}{\sum_{t=91}^T (Y_t - \bar{Y}_t)^2},$$

where T denotes the number of total data points in a given time period. Table 3 presents the out-

of-sample R^2 obtained by the aforementioned several models in the two time periods for predicting 'HOUSTNE' and 'GS5'.

From the results, we observe that FARM outperforms all other benchmarks. In specific, in comparison with PCR, our performance is better. This is due to the possibility that the latent factor regression did not adequately explain the data. Additionally, applying traditional penalized regression methods like Lasso or Elastic Net (El-Net) directly will result in an erroneous estimator $\hat{\beta}$ and worse prediction outcomes when the covariates have a factor structure (highly-correlated). [Fan et al. \[2020a\]](#) propose a factor-adjusted lasso estimator (FarmSelect) to mitigate the impact of latent factors, but they still assume the sparse regression as a sufficient method. From this point of view, their model could be misspecified and performs worse than ours. In terms of the Ridge regression, it assumes the underlying signal β^* is dense instead of sparse. From the outcomes, we conclude that the dense model may not explain this dataset. Finally, our FARM also outperforms the random forest, a well-used model for making predictions via machine learning.

Time period	Data	FARM	Lasso	PCR	Ridge	El-Net	RF	FarmSelect
02.1992-10.2007	HOUSTNE	0.769	0.684	0.372	0.221	0.699	0.497	0.741
	GS5	0.720	0.702	0.056	0.249	0.699	0.557	0.709
08.2010-02.2020	HOUSTNE	0.743	0.374	0.079	0.125	0.348	0.421	0.569
	GS5	0.681	0.650	0.032	0.342	0.653	0.557	0.626

Table 3: Out-of-sample R^2 for predicting 'HOUSTNE' and 'GS5' data using different models in different time periods. In this table, the values in the FARM column denote the prediction results through the factor-augmented linear regression model. In addition, we also compare the prediction results with several benchmarks, Lasso (sparse linear regression), PCR (latent factor regression), Ridge (Ridge regression), El-Net (Elastic Net), RF (Random Forest), and FarmSelect (Factor adjusted Lasso).

We next conduct the hypothesis testing on the adequacy of latent factor regression and sparse linear regression respectively by using FARM as the alternative model. As computing the bootstrap estimate of the null distribution is expensive for testing the adequacy of the factor model, we only conduct the hypothesis testing using the data in the entire two subperiods: 02.1992-10.2007 and 08.2010-02.2020.

The P-values for the tests are given in Table 4. Taking the significant level 0.05, the hypothesis testing results indicate that in most of the cases, the latent factor regression and sparse linear regression, are not sufficient to explain the dataset. These results match well with our prediction results.

Time period	Data	LA_factor	SP_Linear
02.1992-10.2007	HOUSTNE	$< 10^{-3}$	$< 10^{-3}$
	GS5	$1.5 \cdot 10^{-3}$	$4.73 \cdot 10^{-3}$
08.2010-02.2020	HOUSTNE	$< 10^{-3}$	$1.64 \cdot 10^{-1}$
	GS5	$1.98 \cdot 10^{-1}$	$2.94 \cdot 10^{-2}$

Table 4: p -values for testing the adequacy of the latent factor regression and sparse linear regression models to explain 'HOUSTNE' and 'GS5' data in two different time periods. The LA_Factor and SP_Linear have the same meaning as those in Table 3.

6 Conclusion and Discussion

In this paper, we propose a model named Factor Augmented (sparse linear) Regression Model (FARM), which contains the latent factor regression and the sparse linear regression as our special cases. The model expands the space spanned by covariates into useful principal component directions and hence use additional information beyond the linear space spanned by the predictors. We provide theoretical guarantees for our model estimation under the existence of light-tailed and heavy-tailed noises respectively. In addition, we leverage the FARM model as the alternative one to test the adequacy of the latent factor regression model and sparse regression model. We believe that the study is among the first of this kind in high-dimensional inference. The practical performance of our model estimation and our constructed test statistics are proven by extensive simulation studies including both synthetic data and real data. Moreover, it is worth to mention that our model and methodology can be extended to more general supervised learning problems such as nonparametric regression, quantile regression, regression and classification trees, support vector machines, among others where the factor augmentation idea is always useful.

Next, we provide some discussion with several related works. First, we make comparison with [Luciani \[2014\]](#). It is worth noting although our work shares similar idea with [Luciani \[2014\]](#) in terms of incorporating factors into sparse regression, our framework is more general, systematic and contains theirs as special case. Moreover, our intuition is different. To be more specific, they provide a conceptual idea without specifying any statistical model and do not provide any theoretical guarantees. In contrast, we provide a thorough study of the factor augmented sparse regression model (FARM), from the perspective of (robust) estimation to uncertainty quantification with well established methodology and theoretical results. We further utilize our model to test the goodness-of-fit of two important models, namely, factor regression and sparse regression. More importantly, the starting points of deriving our FARM is different from that in [Luciani \[2014\]](#). Specifically, our model is intuitive from the inadequacy of factor regression and sparse regression in many situations, whereas they mainly focus on using past factors and idiosyncratic components to forecast time series data.

Next, we discuss our connection with the sparse and dense model proposed by [Giannone et al. \[2021\]](#). They consider the model

$$Y = \mathbf{x}_t^\top \boldsymbol{\beta}^* + \mathbf{z}_t^\top \boldsymbol{\varphi}^* + \epsilon,$$

where $\mathbf{z}_t$ is a low-dimensional vector whose regression parameter $\boldsymbol{\varphi}^*$ is considered to be dense and $\mathbf{x}_t$ is a high-dimensional covariate whose regression coefficient $\boldsymbol{\beta}^*$ is considered to be sparse. It is worth noting that their model acts as a special case to our FARM when we assume the factor is observable with $\mathbf{z}_t = \mathbf{f}_t$ and $\mathbf{x}_t$ possess factor structure. Moreover, they aim at identifying explainable regression variables and degree of sparseness using tools from Bayesian statistics and empirical illustration, whereas we provide consistency estimation results and also conduct hypothesis testing on $\boldsymbol{\beta}^*$ elementwisely or groupwisely via both theories and numerical studies. Last but not least, it is mentioned in their paper, when they analyze the macroeconomic data, although their posterior sparse level is low, the heat map shows high uncertainty on whether certain predictors should be included in the model due to high collinearity of predictors. In such a scenario, conducting regression directly using their method will fail to recover the true support of the covariates [[Zhao and Yu, 2006](#)] and thus results in unstable estimation results. To remedy this issue, one needs to decompose covariates into the factor and idiosyncratic component and run factor adjusted regression via FARM. This showcases the neces-

sity of estimating this sparse and dense model using FARM instead of ordinary linear regression under strongly correlated covariate.

It is worth mentioning that, we also discuss the connections with several other related literature, such as [Lin and Michailidis \[2020\]](#), [Kneip and Sarda \[2011\]](#), discuss the model selection consistency, and test the contribution of a particular x_i to Y . Due to the space limit, we put them in the appendix §A.

References

- S. C. Ahn and A. R. Horenstein. Eigenvalue ratio test for the number of factors. *Econometrica*, 81(3):1203–1227, 2013.
- M. Avella-Medina, H. S. Battey, J. Fan, and Q. Li. Robust estimation of high-dimensional covariance and precision matrices. *Biometrika*, 105(2):271–284, 2018.
- J. Bai. Inferential theory for factor models of large dimensions. *Econometrica*, 71(1):135–171, 2003.
- J. Bai and K. Li. Statistical analysis of factor models of high dimension. *Ann. Statist.*, 40(1):436–465, 2012.
- J. Bai and S. Ng. Determining the number of factors in approximate factor models. *Econometrica*, 70(1):191–221, 2002.
- J. Bai and S. Ng. Confidence intervals for diffusion index forecasts and inference for factor-augmented regressions. *Econometrica*, 74(4):1133–1150, 2006.
- J. Bai and S. Ng. Forecasting economic time series using targeted predictors. *Journal of Econometrics*, 146(2):304–317, 2008. ISSN 0304-4076.
- E. Bair, T. Hastie, D. Paul, and R. Tibshirani. Prediction by supervised principal components. *Journal of the American Statistical Association*, 101(473):119–137, 2006.
- E. Barut, J. Fan, and A. Verhasselt. Conditional sure independence screening. *J. Amer. Statist. Assoc.*, 111(515):1266–1277, 2016.
- A. Belloni and V. Chernozhukov. ℓ_1 -penalized quantile regression in high-dimensional sparse models. *Ann. Statist.*, 39(1):82–130, 2011.
- D. Bianchi, M. Büchner, and A. Tamoni. Bond risk premiums with machine learning. *The Review of Financial Studies*, 34(2):1046–1089, 2021.
- X. Bing, F. Bunea, and M. Wegkamp. Inference in latent factor regression with clusterable features. *arXiv:1905.12696*, 2019.
- X. Bing, F. Bunea, S. Strimas-Mackey, and M. Wegkamp. Prediction under latent factor regression: Adaptive pcr, interpolating predictors and beyond. *Journal of Machine Learning Research*, 22(177):1–50, 2021.

- F. Bunea, S. Strimas-Mackey, and M. Wegkamp. Interpolating predictors in high-dimensional factor regression. *arXiv:2002.02525*, 2020.
- T. Cai, W. Liu, and X. Luo. A constrained ℓ_1 minimization approach to sparse precision matrix estimation. *J. Amer. Statist. Assoc.*, 106(494):594–607, 2011.
- E. Candes and T. Tao. The Dantzig selector: statistical estimation when p is much larger than n . *Ann. Statist.*, 35(6):2313–2351, 2007.
- V. Chernozhukov, D. Chetverikov, and K. Kato. Gaussian approximations and multiplier bootstrap for maxima of sums of high-dimensional random vectors. *Ann. Statist.*, 41(6):2786–2819, 2013.
- V. Chernozhukov, D. Chetverikov, and K. Kato. Central limit theorems and bootstrap in high dimensions. *Ann. Probab.*, 45(4):2309–2352, 2017.
- V. Chernozhukov, D. Chetverikov, and Y. Koike. Nearly optimal central limit theorem and bootstrap approximations in high dimensions. *arXiv preprint arXiv:2012.09513*, 2020.
- W. Chu, R. Li, and M. Reimherr. Feature screening for time-varying coefficient models with ultrahigh dimensional longitudinal data. *Ann. Appl. Stat.*, 10(2):596, 2016.
- P. G. Coulombe, M. Leroux, D. Stevanovic, and S. Surprenant. Macroeconomic data transformations matter. *International Journal of Forecasting*, 37(4):1338–1354, 2021a.
- P. G. Coulombe, M. Marcellino, and D. Stevanović. Can machine learning catch the covid-19 recession? *National Institute Economic Review*, 256:71–109, 2021b.
- R. Dezeure, P. Bühlmann, and C.-H. Zhang. High-dimensional simultaneous inference with the bootstrap. *Test*, 26(4):685–719, 2017.
- B. Efron, T. Hastie, I. Johnstone, and R. Tibshirani. Least angle regression. *Ann. Statist.*, 32(2):407–499, 2004.
- J. Fan and R. Li. Variable selection via nonconcave penalized likelihood and its oracle properties. *J. Amer. Statist. Assoc.*, 96(456):1348–1360, 2001.
- J. Fan and Y. Liao. Learning latent factors from diversified projections and its applications to over-estimated and weak factors. *Journal of the American Statistical Association*, 0(0):1–16, 2020.
- J. Fan and J. Lv. Sure independence screening for ultrahigh dimensional feature space. *J. R. Stat. Soc. Ser. B Stat. Methodol.*, 70(5):849–911, 2008.
- J. Fan and J. Lv. Nonconcave penalized likelihood with NP-dimensionality. *IEEE Trans. Inform. Theory*, 57(8):5467–5484, 2011.
- J. Fan and R. Song. Sure independence screening in generalized linear models with NP-dimensionality. *Ann. Statist.*, 38(6):3567–3604, 2010.

- J. Fan, S. Guo, and N. Hao. Variance estimation using refitted cross-validation in ultrahigh dimensional regression. *J. R. Stat. Soc. Ser. B. Stat. Methodol.*, 74(1):37–65, 2012.
- J. Fan, Y. Liao, and M. Mincheva. Large covariance estimation by thresholding principal orthogonal complements. *J. R. Stat. Soc. Ser. B. Stat. Methodol.*, 75(4):603–680, 2013. With 33 discussions by 57 authors and a reply by Fan, Liao and Mincheva.
- J. Fan, Y. Fan, and E. Barut. Adaptive robust variable selection. *Ann. Statist.*, 42(1):324, 2014.
- J. Fan, Q. Li, and Y. Wang. Estimation of high dimensional mean regression in the absence of symmetry and light tail assumptions. *J. R. Stat. Soc. Ser. B Stat. Methodol.*, 79(1):247, 2017a.
- J. Fan, L. Xue, and j. Yao. Sufficient forecasting using factor models. *Journal of Econometrics*, 201(2):292–306, 2017b. ISSN 0304-4076.
- J. Fan, Y. Ke, and K. Wang. Factor-adjusted regularized model selection. *J. Econometrics*, 216(1):71–85, 2020a.
- J. Fan, R. Li, C.-H. Zhang, and H. Zou. Statistical foundations of data science. 2020b.
- J. Fan, R. Masini, and M. C. Medeiros. Bridging factor and sparse models. *arXiv:2102.11341*, 2021a.
- J. Fan, Z. Yang, and M. Yu. Understanding implicit regularization in over-parameterized single index model. *arXiv:2007.08322v3*, 2021b.
- J. Fan, J. Guo, and S. Zheng. Estimating number of factors by adjusted eigenvalues thresholding. *Journal of the American Statistical Association*, 117:852–861, 2022.
- D. Giannone, M. Lenza, and G. E. Primiceri. Economic predictions with big data: The illusion of sparsity. *Econometrica*, 89(5):2409–2437, 2021.
- P. Goulet Coulombe. The macroeconomy as a random forest. *Available at SSRN 3633110*, 2020.
- A. S. Hall et al. Machine learning approaches to macroeconomic forecasting. *The Federal Reserve Bank of Kansas City Economic Review*, 103(63):2, 2018.
- M. Hernan and J. Robins. *Causal Inference*. Chapman & Hall/CRC Monographs on Statistics & Applied Probab. CRC Press LLC, 2019. ISBN 9781420076165.
- H. Hotelling. Analysis of a complex of statistical variables into principal components. *Journal of educational psychology*, 24(6):417, 1933.
- G. Imbens and D. Rubin. Causal inference for statistics, social, and biomedical sciences: An introduction. 2015.
- A. Javanmard and A. Montanari. Confidence intervals and hypothesis testing for high-dimensional regression. *J. Mach. Learn. Res.*, 15:2869–2909, 2014.
- I. T. Jolliffe. A note on the use of principal components in regression. *Journal of the Royal Statistical Society. Series C (Applied Statistics)*, 31(3):300–303, 1982.

- A. Kneip and P. Sarda. Factor models and variable selection in high-dimensional regression analysis. *The Annals of Statistics*, 39(5):2410–2447, 2011.
- C. Lam and Q. Yao. Factor modeling for high-dimensional time series: Inference for the number of factors. *The Annals of Statistics*, 40(2):694 – 726, 2012.
- G. Li, H. Peng, J. Zhang, and L. Zhu. Robust rank correlation based screening. *Ann. Statist.*, 40(3):1846–1877, 2012.
- Q. Li and L. Li. Integrative factor regression and its inference for multimodal data analysis. *J. Amer. Statist. Assoc.*, 113(521):1–15, 2021.
- Q. Li, G. Cheng, J. Fan, and Y. Wang. Embracing the blessing of dimensionality in factor models. *J. Amer. Statist. Assoc.*, 113(521):380–389, 2018.
- J. Lin and G. Michailidis. System identification of high-dimensional linear dynamical systems with serially correlated output noise components. *IEEE Transactions on Signal Processing*, 68:5573–5587, 2020.
- J. Liu, R. Li, and R. Wu. Feature selection for varying coefficient models with ultrahigh-dimensional covariates. *J. Amer. Statist. Assoc.*, 109(505):266–274, 2014.
- P.-L. Loh and M. J. Wainwright. High-dimensional regression with noisy and missing data: Provable guarantees with nonconvexity. *Ann. Statist.*, 40(3):1637–1664, 2012.
- M. Luciani. Forecasting with approximate dynamic factor models: the role of non-pervasive shocks. *International Journal of Forecasting*, 30(1):20–29, 2014.
- M. W. McCracken and S. Ng. Fred-md: A monthly database for macroeconomic research. *Journal of Business & Economic Statistics*, 34(4):574–589, 2016.
- R. Nickl and S. Van De Geer. Confidence sets in sparse regression. *Ann. Statist.*, 41(6):2852–2876, 2013.
- B. Peng, L. Wang, and Y. Wu. An error bound for L_1 -norm support vector machine coefficients in ultra-high dimension. *J. Mach. Learn. Res.*, 17(1):8279–8304, 2016.
- D. F. Saldana and Y. Feng. Sis: An r package for sure independence screening in ultrahigh-dimensional statistical models. *Journal of Statistical Software*, 83(2):1–25, 2018.
- C. Shi, R. Song, Z. Chen, and R. Li. Linear hypothesis testing for high dimensional generalized linear models. *Ann. Statist.*, 47(5):2671–2703, 2019.
- S. Smeeke and E. Wijler. Macroeconomic forecasting using penalized regression methods. *International journal of forecasting*, 34(3):408–430, 2018.
- J. H. Stock and M. W. Watson. Forecasting using principal components from a large number of predictors. *Journal of the American Statistical Association*, 97(460):1167–1179, 2002.

- Q. Sun, W.-X. Zhou, and J. Fan. Adaptive Huber regression. *Journal of the American Statistical Association*, 115(529):254–265, 2020.
- T. Sun and C.-H. Zhang. Scaled sparse linear regression. *Biometrika*, 99(4):879–898, 2012.
- R. Tibshirani. Regression shrinkage and selection via the lasso. *J. Roy. Statist. Soc. Ser. B*, 58(1):267–288, 1996.
- S. Van de Geer. High-dimensional generalized linear models and the lasso. *Ann. Statist.*, 36(2):614–645, 2008.
- S. van de Geer, P. Bühlmann, Y. Ritov, and R. Dezeure. On asymptotically optimal confidence regions and tests for high-dimensional models. *Ann. Statist.*, 42(3):1166–1202, 2014.
- W. Wang and J. Fan. Asymptotics of empirical eigenstructure for high dimensional spiked covariance. *Ann. Statist.*, 45(3):1342–1374, 2017.
- X. Wang and C. Leng. High dimensional ordinary least squares projection for screening variables. *J. R. Stat. Soc. Ser. B Stat. Methodol.*, pages 589–611, 2016.
- G. Yu and J. Bien. Estimating the error variance in a high-dimensional linear model. *Biometrika*, 106(3):533–546, 2019.
- C.-H. Zhang. Nearly unbiased variable selection under minimax concave penalty. *Ann. Statist.*, 38(2):894–942, 2010.
- C.-H. Zhang and S. S. Zhang. Confidence intervals for low dimensional parameters in high dimensional linear models. *J. R. Stat. Soc. Ser. B. Stat. Methodol.*, 76(1):217–242, 2014.
- N. Zhang, W. Jiang, and Y. Lan. On the sure screening properties of iteratively sure independence screening algorithms. *arXiv:1812.01367*, 2019.
- X. Zhang and G. Cheng. Simultaneous inference for high-dimensional linear models. *J. Amer. Statist. Assoc.*, 112(518):757–768, 2017.
- X. Zhang, Y. Wu, L. Wang, and R. Li. Variable selection for support vector machines in moderately high dimensions. *J. R. Stat. Soc. Ser. B Stat. Methodol.*, 78(1):53, 2016.
- P. Zhao and B. Yu. On model selection consistency of Lasso. *J. Mach. Learn. Res.*, 7:2541–2563, 2006.
- P. Zhao, Y. Yang, and Q.-C. He. Implicit regularization via hadamard product over-parametrization in high-dimensional linear regression. *arXiv:1903.09367*, 2019.
- L.-P. Zhu, L. Li, R. Li, and L.-X. Zhu. Model-free feature screening for ultrahigh-dimensional data. *J. Amer. Statist. Assoc.*, 106(496):1464–1475, 2011.
- H. Zou. The adaptive lasso and its oracle properties. *J. Amer. Statist. Assoc.*, 101(476):1418–1429, 2006.

Supplement material for “Are Latent Factor Regression and Sparse Regression Adequate?”

A More Discussion

1. Comparison with related work

- Comparison with [Kneip and Sarda \[2011\]](#). [Kneip and Sarda \[2011\]](#) study a similar problem by incorporating factors into sparse regression under stronger assumptions. Their model requires the idiosyncratic component $\mathbf{u}_t$ to be uncorrelated, which is quite a strong assumption in the high dimension, whereas we only require the conditional number of the covariance of $\mathbf{u}_t$ to be bounded. In addition, they require the noise distribution to be Gaussian, whereas our framework not only allow for general sub-Gaussian noises but heavy-tailed noises as well. Most importantly, in addition to estimating the parameters of interest, we also focus on high-dimensional inference: testing the goodness of fit of two important models, factor regression and sparse regression which are not considered in [Kneip and Sarda \[2011\]](#).
- Connection with [Lin and Michailidis \[2020\]](#). [Lin and Michailidis \[2020\]](#) study factor model $\mathbf{x}_t = \mathbf{B}\mathbf{f}_t + \mathbf{u}_t$ where factors, idiosyncratic component following certain vector autoregression (VAR) model. In this case, they predict the current covariate $\mathbf{x}_t$ by augmenting factors to past covariate $\mathbf{x}_{t-1}$. Compared with them, we study different objectives in a sense that we focus on high-dimensional inference for regression problems with general responses. It is worth to note that when our $Y_t = x_{t,i}, i \in [d]$, their model becomes our special case.

2. [Model Selection Consistency] Model selection consistency (recovering the true support of β^*) plays an essential role in model prediction, especially if the dataset is in high dimension but the underlying model is sparse. However, when the features are strongly co-linear, the irrepresentable condition fails, and we can not achieve model selection consistency [[Zhao and Yu, 2006](#)]. Just as the case mentioned in [Giannone et al. \[2021\]](#), one may delete the actual important variables but involve useless ones while making predictions. However, it is worth noting that even when sparse regression is adequate, if one decomposes $\mathbf{X}$ into $(\mathbf{F}, \mathbf{U})$, which are two less correlated ones, one will eliminate the high-colinearity across the features and thus obtain model selection consistency [[Fan et al., 2020](#)]. Thus, even if the sparse model is adequate, FARM still has many merits in the sense that we can achieve model selection consistency while preserving mean square errors.

3. Contribution of a particular $\mathbf{x}_i$ to Y . There are two ways for $\mathbf{x}_i$ to influence the response Y .

The first is through idiosyncratic component $\mathbf{u}_i$ and the other is indirect effect through $\mathbf{f}_t$. To test the first effect, we leverage the debiasing idea given in section 3 to test whether $\beta_i^* = 0$.

For the indirect influence of $\mathbf{x}_i$ on Y through $\mathbf{f}_t$, this procedure is more involved. This is equivalent to testing whether the i -th row of the $\mathbf{B}$ is zero or not. Recall that in our model formulation, we have the eigenspace with respect to spiked eigenvalues of $\text{Cov}(\mathbf{x})$ is spanned by columns of $\mathbf{B}$. Thus, testing whether the i -th row of $\mathbf{B}$ is zero is equivalent to testing whether the i -th row of spiked

eigenvectors $\mathbf{V}_{i,\cdot}$ of $\text{Cov}(\mathbf{x})$ is zero. To accommodate this issue, we leverage the recent development on doing statistical inference on eigenspace [Yan et al., 2021]. They derive asymptotic distribution for $\hat{\mathbf{V}}_{i,\cdot}\mathbf{R} - \mathbf{V}_{i,\cdot} \rightarrow N(0, \Sigma_{V,i})$. Here $\hat{\mathbf{V}}_{i,\cdot} \in \mathbb{R}^K$ is the i -th row of the estimated spiked eigenvectors of $\mathbf{X}$, $\mathbf{R} \in \mathbb{R}^{K \times K}$ is a rotation matrix and $\Sigma_{V,i}$ is a covariance matrix that can be estimated via data. Under the null hypothesis, we have $\mathbf{V}_{i,\cdot} = 0$. In this scenario, we have $\hat{\mathbf{V}}_{i,\cdot}\mathbf{R}\mathbf{R}^\top\hat{\mathbf{V}}_{i,\cdot}^\top \rightarrow \|\mathbf{Z}\|_2^2$ with $\mathbf{Z} \sim N(0, \Sigma_{U,i})$. Thus, the hypothesis testing can be done by comparing $\hat{\mathbf{V}}_{i,\cdot}\hat{\mathbf{V}}_{i,\cdot}^\top$ with $1 - \alpha$ quantile of the distribution of $\|\mathbf{Z}\|_2^2$, which can be estimated via bootstrap.

B Additional Numerical Results

B.1 Model Prediction

In this subsection, we provide a detailed simulation study on the prediction performance of latent factor regression and sparse regression with respect to FARM when our FARM is the true underlying model. First, we demonstrate how latent factor regression and sparse regression behave when they have different levels of misspecification.

As for data generation, we let $n = 200$, $K = 5$, dimension of covariate $d = 1000$, $\boldsymbol{\varphi}^\star = 0.8 \cdot \mathbf{1}_K$, $\boldsymbol{\beta}^\star = (v \cdot \mathbf{1}_{20}^\top, \mathbf{0}_{p-20}^\top)^\top$ with the magnitude of v varies uniformly from 0 to 1.20. Throughout this subsection, we generate every entry of $\mathbf{F}$ and $\mathbf{U}$ from the standard Gaussian distribution and let every entry of $\mathbf{B}$ be generated from uniform distribution with $\text{Unif}(-2, 2)$. The noise distribution ϵ is given by standard Gaussian distribution. The experiment is repeated 500 times and we record the out-of-sample MSE. We then illustrate our comparison with factor regression in Table 1. As for our comparison with sparse

$\ \boldsymbol{\beta}^\star\ _\infty$	0.00	0.20	0.40	0.60	0.80	1.00	1.20
FA Reg	0.25	1.02	3.69	8.33	13.42	18.74	28.14
FARM	0.27	0.65	0.69	0.67	0.65	0.74	0.96

Table 1: Out-of-sample MSE of our model (FARM) with factor regression (FA Reg).

regression, we adopt the aforementioned setting, except that we fix $\boldsymbol{\beta}^\star = (0.8 \cdot \mathbf{1}_{20}^\top, \mathbf{0}_{p-20}^\top)^\top$ and let $\boldsymbol{\varphi}^\star = v \cdot \mathbf{1}_K$, with v ranges uniformly from 0 to 1.20. We summarize the comparison results in Table 2.

$\ \boldsymbol{\varphi}^\star\ _\infty$	0.00	0.20	0.40	0.60	0.80	1.00	1.20
SP Reg	1.94	2.14	2.15	2.37	2.24	2.40	2.54
FARM	0.55	0.59	0.61	0.62	0.59	0.60	0.61

Table 2: Out-of-sample MSE of our model (FARM) with sparse regression (SP Reg).

Finally, we illustrate the model prediction performances of factor regression and sparse regression

with respect to our FARM when the same size n increases. All settings are the same with those mentioned above, except that we fix $\beta^* = (0.8 \cdot \mathbf{1}_{20}^\top, \mathbf{0}_{p-20}^\top)^\top$ and let $\varphi^* = 0.8 \cdot \mathbf{1}_K$. The out-of-sample MSE are recorded below in Table 3. In addition, we also compare FARM with the other two models (Sparse regression and Latent factor regression) using huber loss when there exist heavy-tailed noise. We keep all the aforementioned settings except for changing the noise distribution from standard Gaussian distribution to t_3 distribution. The simulation results are summarized in the following Table 4.

n	200	229	257	286	314	343	371	400
FA Reg	13.59	12.03	14.37	11.81	12.27	13.20	11.79	12.46
SP Reg	2.25	1.86	1.87	1.63	1.59	1.49	1.36	1.38
FARM	0.58	0.48	0.45	0.39	0.35	0.39	0.34	0.31

Table 3: Out-of-sample MSE of our model (FARM) with sparse regression (SP Reg) and factor regression (FA Reg) when n varies.

n	200	229	257	286	314	343	371	400
FA Reg	16.39	18.00	16.60	14.37	15.12	15.97	18.19	14.85
SP Reg	13.73	15.06	14.07	12.63	13.22	13.90	16.67	13.14
FARM	6.67	5.99	4.71	4.04	4.02	3.53	4.77	2.93

Table 4: Out-of-sample MSE of our model (FARM) with sparse regression (SP Reg) and factor regression (FA Reg) when n varies when heavy-tailed error exists.

B.2 Additional Plots

B.3 Implementation Details of Empirical Applications

In this section, we describe the implementation details of different methods used in Empirical Applications. For the implementation of Lasso [Tibshirani, 1996], Ridge, Elastic Net [Zou and Hastie, 2005] we use *glmnet* package in R directly to conduct the numerical studies. The tuning parameters λ and α in the models are chosen via leave-one-out cross-validation (by definition, the parameter α only needs to be tuned in Elastic Net). For the stableness of the algorithms, for FARM, PCR, FarmSelect [Fan et al., 2020], throughout this empirical application, we use PCA to estimate the factors, as suggested in §2, and let the number of factors be the maximum number of (2.3) and 2. For the implementation of FARM, FarmSelect, their first steps of estimating $\hat{\beta}$ are the same. In specific, we decompose the covariate $\mathbf{X}$ into factor and idiosyncratic component $[\hat{\mathbf{F}}, \hat{\mathbf{U}}]$ and then use Lasso to estimate the loadings $\hat{\beta}$ as we described in §2. The prediction procedures are different, namely, we use covariate $\mathbf{x}_{\text{new}}$ and $[\hat{\mathbf{f}}_{\text{new}}, \hat{\mathbf{u}}_{\text{new}}]$ to make prediction, for FarmSelect and FARM, respectively. In terms of using Random Forest, we use the package *randomForest* in R to implement the model estimation and prediction

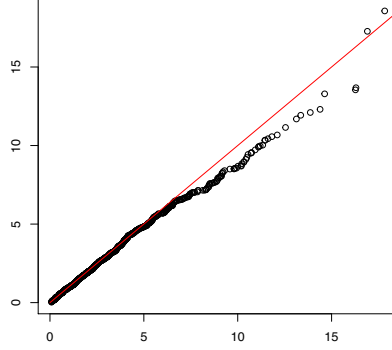


Figure 7: Quantiles of the χ_K^2 distribution against those of the test statistic without sample splitting. The x-axis represents the quantiles of the test statistic whereas y-axis is the quantiles of χ_K^2 distribution.

directly. For more details of our implementation and reproduction, we put the relevant codes in github path given in the ACC form.

C Technical Lemmas

Recall that $\tilde{\Sigma} = n^{-1}\hat{\mathbf{U}}^\top \hat{\mathbf{U}}$ and $\mathcal{C}(\mathcal{S}, 3) = \{\mathbf{v} \in \mathbb{R}^d : \|\mathbf{v}_{\mathcal{S}}\|_1 \leq 3\|\mathbf{v}_{\mathcal{S}^c}\|_1\}$ for any subset $\mathcal{S} \subset [d]$. First, we introduce the following Lemma C.1. In this Lemma, when the RSC condition is satisfied, we are able to achieve some certain ℓ_2 -statistical rates for $\hat{\beta}_\lambda$ and $\hat{\mathbf{U}}\hat{\beta}_\lambda$.

Lemma C.1. Assume that $\lambda \geq (2/n)\|\hat{\mathbf{U}}^\top(\tilde{\mathbf{Y}} - \hat{\mathbf{U}}\beta^*)\|_\infty$ and for some positive constant $\kappa(\mathcal{S}_*, 3)$,

$$\min_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_*, 3)} \frac{\mathbf{v}^\top \tilde{\Sigma} \mathbf{v}}{\|\mathbf{v}\|_2^2} \geq \kappa(\mathcal{S}_*, 3).$$

Then we have $\hat{\beta}_\lambda - \beta^* \in \mathcal{C}(\mathcal{S}_*, 3)$,

$$\|\hat{\beta}_\lambda - \beta^*\|_2 \leq \frac{3\lambda\sqrt{|\mathcal{S}_*|}}{\kappa(\mathcal{S}_*, 3)} \quad \text{and} \quad \|\hat{\mathbf{U}}(\hat{\beta}_\lambda - \beta^*)\|_2^2 \leq \frac{9n\lambda^2|\mathcal{S}_*|}{\kappa(\mathcal{S}_*, 3)}.$$

Proof of Lemma C.1. Write $\delta = \hat{\beta}_\lambda - \beta^*$. Following the proof of Theorem 7.2 in Bickel et al. [2009], we obtain $\delta \in \mathcal{C}(\mathcal{S}_*, 3)$ and

$$\kappa(\mathcal{S}_*, 3)\|\delta\|_2^2 \leq \delta^\top \tilde{\Sigma} \delta = \frac{1}{n}\|\hat{\mathbf{U}}\delta\|_2^2 \leq 3\lambda\|\delta_{\mathcal{S}_*}\|_1 \leq 3\lambda\sqrt{|\mathcal{S}_*|}\|\delta\|_2.$$

Then we have

$$\|\boldsymbol{\delta}\|_2 \leq \frac{3\lambda\sqrt{|\mathcal{S}_\star|}}{\kappa(\mathcal{S}_\star, 3)} \text{ and } \|\widehat{\mathbf{U}}\boldsymbol{\delta}\|_2^2 \leq 3n\lambda\sqrt{|\mathcal{S}_\star|}\|\boldsymbol{\delta}\|_2 \leq \frac{9n\lambda^2|\mathcal{S}_\star|}{\kappa(\mathcal{S}_\star, 3)}.$$

□

In the next Lemma C.2, we prove that under quite mild conditions, the RSC condition given in Lemma C.1 holds with high probability.

Lemma C.2. *Under Assumptions 2.1–2.4, for any set $\mathcal{S} \subset \{1, \dots, d\}$ with*

$$|\mathcal{S}| \left(\frac{\log d}{n} + \frac{1}{d} \right) \rightarrow 0, \quad (\text{C.1})$$

there exists a constant $\kappa(\mathcal{S}, 3) > 0$ such that

$$\mathbb{P} \left(\min_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}, 3)} \frac{\mathbf{v}^\top \tilde{\Sigma} \mathbf{v}}{\|\mathbf{v}\|_2^2} \geq \kappa(\mathcal{S}, 3) \right) \rightarrow 1, \text{ as } n \rightarrow \infty.$$

Proof of Lemma C.2. For any vector $\mathbf{v} \in \mathbb{R}^d$ such that $\|\mathbf{v}_{\mathcal{S}^c}\|_1 \leq 3\|\mathbf{v}_{\mathcal{S}}\|_1$, we have

$$\|\mathbf{v}\|_1^2 \leq 16\|\mathbf{v}_{\mathcal{S}}\|_1^2 \leq 16|\mathcal{S}|\|\mathbf{v}_{\mathcal{S}}\|_2^2 \leq 16|\mathcal{S}|\|\mathbf{v}\|_2^2,$$

which implies that

$$\frac{\mathbf{v}^\top \tilde{\Sigma} \mathbf{v}}{\|\mathbf{v}\|_2^2} \geq \frac{\mathbf{v}^\top \widehat{\Sigma} \mathbf{v}}{\|\mathbf{v}\|_2^2} - \frac{\|\tilde{\Sigma} - \widehat{\Sigma}\|_{\max} \|\mathbf{v}\|_1^2}{\|\mathbf{v}\|_2^2} \geq \frac{\mathbf{v}^\top \widehat{\Sigma} \mathbf{v}}{\|\mathbf{v}\|_2^2} - 16|\mathcal{S}|\|\tilde{\Sigma} - \widehat{\Sigma}\|_{\max}, \quad (\text{C.2})$$

where $\widehat{\Sigma} = n^{-1}\mathbf{U}^\top \mathbf{U}$. Combining this with Lemma C.6 and (C.1), we obtain

$$\mathbb{P}(\mathcal{E}_1) := \mathbb{P} \left(16|\mathcal{S}|\|\tilde{\Sigma} - \widehat{\Sigma}\|_{\max} \leq \frac{\lambda_{\min}(\Sigma)}{8} \right) \rightarrow 1, \text{ as } n \rightarrow \infty.$$

For $\phi > 0$, define the sparse set $\mathbb{K}(\phi) = \{\mathbf{v} \in \mathbb{R}^d : \|\mathbf{v}\|_0 \leq \phi, \|\mathbf{v}\|_2 \leq 1\}$. Taking $\phi_\diamond = 128|\mathcal{S}|$. It follows from (C.1) that $\phi_\diamond = o(n/\log d)$. Then, by Lemma 15 in Loh and Wainwright [2012] and Assumption 2.1, it follows that

$$\mathbb{P}(\mathcal{E}_2) := \mathbb{P} \left(\sup_{\mathbf{v} \in \mathbb{K}(2\phi_\diamond)} |\mathbf{v}^\top (\widehat{\Sigma} - \Sigma) \mathbf{v}| \leq \frac{\lambda_{\max}(\Sigma)}{54} \right) \rightarrow 1, \text{ as } n \rightarrow \infty.$$

Under the event $\mathcal{E}_2$, by Lemma 13 in Loh and Wainwright [2012],

$$\mathbf{v}^\top \widehat{\Sigma} \mathbf{v} \geq \frac{\lambda_{\min}(\Sigma)}{2} \|\mathbf{v}\|_2^2 - \frac{\lambda_{\min}(\Sigma)}{\phi_\diamond} \|\mathbf{v}\|_1^2 \geq \frac{\lambda_{\min}(\Sigma)}{2} \|\mathbf{v}\|_2^2 - \frac{16|\mathcal{S}|\lambda_{\min}(\Sigma)}{\phi_\diamond} \|\mathbf{v}\|_2^2.$$

Combining this with (C.2), under $\mathcal{E}_1 \cap \mathcal{E}_2$, we obtain

$$\frac{\mathbf{v}^\top \tilde{\Sigma} \mathbf{v}}{\|\mathbf{v}\|_2^2} \geq \frac{\lambda_{\min}(\Sigma)}{2} - 16|\mathcal{S}|\|\tilde{\Sigma} - \hat{\Sigma}\|_{\max} - \frac{16|\mathcal{S}|\lambda_{\min}(\Sigma)}{\phi_\diamond} \geq \frac{\lambda_{\min}(\Sigma)}{4}.$$

□

Lemma C.3. Under Assumptions 2.1–2.4, for any vector $\phi \in \mathbb{R}^K$ with $\|\phi\|_2 = 1$, we have

$$\|\hat{\mathbf{U}}^\top \mathbf{F} \phi\|_\infty = O_{\mathbb{P}}(\mathcal{V}_{n,d}). \quad (\text{C.3})$$

Proof of Lemma C.3. Define $\mathcal{E}_\lambda = \{\lambda_{\min}(\mathbf{H}^\top \mathbf{H}) \geq 1/2\}$. It follows from Lemma 2.1 that $\mathbb{P}(\mathcal{E}_\lambda) \rightarrow 1$. Since $\hat{\mathbf{F}}^\top \hat{\mathbf{U}} = \mathbf{O}$, under $\mathcal{E}_\lambda$, we have

$$\begin{aligned} \|\hat{\mathbf{U}}^\top \mathbf{F} \phi\|_\infty &= \|\hat{\mathbf{U}}^\top (\mathbf{F} - \hat{\mathbf{F}} \mathbf{H}^{-\top}) \phi\|_\infty \leq \|(\hat{\mathbf{U}} - \mathbf{U})^\top (\mathbf{F} \mathbf{H}^\top - \hat{\mathbf{F}}) \mathbf{H}^{-\top} \phi\|_\infty \\ &\quad + \|\mathbf{U}^\top (\mathbf{F} \mathbf{H}^\top - \hat{\mathbf{F}}) \mathbf{H}^{-\top} \phi\|_\infty =: \Delta^\diamond + \Delta^\circ. \end{aligned}$$

We first bound $\Delta^\diamond$. By Lemma 2.1 and the Cauchy-Schwarz inequality, it follows that

$$\begin{aligned} \Delta^\diamond &\leq \left(\max_{j \in [d]} \sum_{t=1}^n |\hat{u}_{tj} - u_{tj}|^2 \right)^{1/2} \|\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top\|_{\mathbb{F}} \|\mathbf{H}^{-\top} \phi\|_2 \\ &= O_{\mathbb{P}} \left\{ \sqrt{\left(\log d + \frac{n}{d} \right) \left(\frac{n}{d} + \frac{1}{n} \right)} \right\} = O_{\mathbb{P}}(\mathcal{V}_{n,d}). \end{aligned}$$

We now bound Δ° . Recall that $\mathbf{H} = n^{-1} \mathbf{V}^{-1} \hat{\mathbf{F}}^\top \mathbf{F} \mathbf{B}^\top \mathbf{B}$ and $\hat{\mathbf{F}} \mathbf{V} = n^{-1} \mathbf{X} \mathbf{X}^\top \hat{\mathbf{F}}$. Hence

$$\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top = \frac{1}{n} \mathbf{F} \mathbf{B}^\top \mathbf{U}^\top \hat{\mathbf{F}} \mathbf{V}^{-1} + \frac{1}{n} \mathbf{U} \mathbf{B} \mathbf{F}^\top \hat{\mathbf{F}} \mathbf{V}^{-1} + \frac{1}{n} \mathbf{U} \mathbf{U}^\top \hat{\mathbf{F}} \mathbf{V}^{-1}.$$

By the triangle inequality, it follows that

$$\begin{aligned} \Delta^\circ &\leq \frac{1}{n} \left(\|\mathbf{U}^\top \mathbf{F} \mathbf{B}^\top \mathbf{U}^\top \hat{\mathbf{F}} \mathbf{V}^{-1} \mathbf{H}^{-\top} \phi\|_\infty + \|\mathbf{U}^\top \mathbf{U} \mathbf{B} \mathbf{F}^\top \hat{\mathbf{F}} \mathbf{V}^{-1} \mathbf{H}^{-\top} \phi\|_\infty \right. \\ &\quad \left. + \|\mathbf{U}^\top \mathbf{U} \mathbf{U}^\top \hat{\mathbf{F}} \mathbf{V}^{-1} \mathbf{H}^{-\top} \phi\|_\infty \right) =: \Delta_1^\circ + \Delta_2^\circ + \Delta_3^\circ. \end{aligned}$$

By Assumptions 2.1 and 2.3, we have $\mathbb{E} \|\mathbf{B}^\top \mathbf{U}^\top\|_{\mathbb{F}}^2 = n \text{tr}(\mathbf{B}^\top \Sigma \mathbf{B})$ and

$$\mathbb{E} \|\mathbf{B}^\top \mathbf{U}^\top \mathbf{F}\|_{\mathbb{F}}^2 = n \mathbb{E} \|\mathbf{B}^\top \mathbf{u}\|_2^2 \|\mathbf{f}\|_2^2 \leq n \sqrt{\mathbb{E} \|\mathbf{B}^\top \mathbf{u}\|_2^4} \sqrt{\mathbb{E} \|\mathbf{f}\|_2^4} = O(nd).$$

Consequently, by Lemma 2.1, it follows that

$$\|\mathbf{B}^\top \mathbf{U}^\top \hat{\mathbf{F}}\|_{\mathbb{F}} \leq \|\mathbf{B}^\top \mathbf{U}^\top (\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top)\|_{\mathbb{F}} + \|\mathbf{B}^\top \mathbf{U}^\top \mathbf{F} \mathbf{H}^\top\|_{\mathbb{F}}$$

$$\begin{aligned}
&\leq \| \mathbf{B}^\top \mathbf{U}^\top \|_{\mathbb{F}} \| \hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top \|_{\mathbb{F}} + \| \mathbf{H} \|_2 \| \mathbf{B}^\top \mathbf{U}^\top \mathbf{F} \|_{\mathbb{F}} \\
&= O_{\mathbb{P}} \left\{ \sqrt{nd \left(\frac{n}{d} + \frac{1}{n} \right)} + \sqrt{nd} \right\} = O_{\mathbb{P}} \left(n + \sqrt{nd} \right).
\end{aligned}$$

Combining this with $\max_{j \in [d]} \| \mathbf{e}_j^\top \mathbf{U}^\top \mathbf{F} \|_2^2 = \max_{j \in [d]} \| \sum_{t=1}^n u_{tj} \mathbf{f}_t \|_2^2 = O_{\mathbb{P}}(n \log d)$, we obtain

$$\begin{aligned}
\Delta_1^\circ &= \max_{j \in [d]} \frac{1}{n} | \mathbf{e}_j^\top \mathbf{U}^\top \mathbf{F} \mathbf{B}^\top \mathbf{U}^\top \hat{\mathbf{F}} \mathbf{V}^{-1} \mathbf{H}^{-\top} \phi | \\
&\leq \max_{j \in [d]} \frac{1}{n} \| \mathbf{e}_j^\top \mathbf{U}^\top \mathbf{F} \|_2 \| \mathbf{B}^\top \mathbf{U}^\top \hat{\mathbf{F}} \|_{\mathbb{F}} \| \mathbf{V}^{-1} \|_2 \| \mathbf{H}^{-\top} \|_2 \\
&= O_{\mathbb{P}} \left\{ \sqrt{(\log d)/d} + \sqrt{n \log d/d} \right\}.
\end{aligned}$$

Write $\mathbf{B} = (\tilde{\mathbf{b}}_1, \dots, \tilde{\mathbf{b}}_K) \in \mathbb{R}^{d \times K}$. By Assumptions 2.3 and 2.4,

$$\max_{k \in [K]} \max_{j \in [d]} \mathbb{E} \left(u_j \mathbf{u}^\top \tilde{\mathbf{b}}_k \right) = \max_{k \in [K]} \max_{j \in [d]} \Sigma_j^\top \tilde{\mathbf{b}}_k \leq \max_{k \in [K]} \max_{j \in [d]} \| \Sigma_j \|_1 \| \tilde{\mathbf{b}}_k \|_\infty \leq \frac{\Upsilon}{\kappa},$$

where $\Sigma_j \in \mathbb{R}^d$ denotes the j -th column of the covariance matrix Σ . Moreover, by Assumption 2.1, for each $k \in [K]$, $\{u_{tj} \mathbf{u}_t^\top \tilde{\mathbf{b}}_k\}_{t=1}^n$ is a sequence of i.i.d. sub-exponential random variables with

$$\max_{k \in [K]} \max_{j \in [d]} \| u_{tj} \mathbf{u}_t^\top \tilde{\mathbf{b}}_k \|_{\psi_1} \leq \max_{j \in [d]} \| u_{tj} \|_{\psi_2} \max_{k \in [K]} \| \mathbf{u}_t^\top \tilde{\mathbf{b}}_k \|_{\psi_2} \leq c_0^2 \max_{k \in [K]} \| \tilde{\mathbf{b}}_k \|_2 \leq c_0^2 \Upsilon \sqrt{d}.$$

Hence, it follows that

$$\max_{j \in [d]} \left\| \sum_{t=1}^n u_{tj} \mathbf{u}_t^\top \mathbf{B} \right\|_2 \leq n \max_{j \in [d]} \| \mathbb{E}(u_j \mathbf{u}^\top \mathbf{B}) \|_2 + \max_{j \in [d]} \left\| \sum_{t=1}^n \mathbb{E}_0(u_{tj} \mathbf{u}_t^\top \mathbf{B}) \right\|_2 = O_{\mathbb{P}} \left(n + \sqrt{nd \log d} \right),$$

where $\mathbb{E}_0(\cdot) = \cdot - \mathbb{E}(\cdot)$. Consequently, we obtain

$$\begin{aligned}
\Delta_2^\circ &= \max_{j \in [d]} \frac{1}{n} \left| \sum_{t=1}^n u_{tj} \mathbf{u}_t^\top \mathbf{B} \sum_{s=1}^n \mathbf{f}_s \hat{\mathbf{f}}_s^\top \mathbf{V}^{-1} \mathbf{H}^{-\top} \phi \right| \\
&\leq \frac{1}{n} \| \mathbf{V}^{-1} \|_2 \| \mathbf{H}^{-\top} \|_2 \left\| \sum_{s=1}^n \hat{\mathbf{f}}_s \mathbf{f}_s^\top \right\|_2 \max_{j \in [d]} \left\| \sum_{t=1}^n u_{tj} \mathbf{u}_t^\top \mathbf{B} \right\|_2 \\
&= O_{\mathbb{P}} \left\{ n/d + \sqrt{n(\log d)/d} \right\}.
\end{aligned}$$

By the triangle inequality,

$$\Delta_3^\circ \leq \frac{1}{n} \left(\max_{j \in [d]} \| \mathbf{e}_j \mathbf{U}^\top \mathbf{U} \mathbf{U}^\top \mathbf{F} \mathbf{H}^\top \|_2 + \max_{j \in [d]} \| \mathbf{e}_j \mathbf{U}^\top \mathbf{U} \mathbf{U}^\top (\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top) \|_2 \right) \| \mathbf{V}^{-1} \|_2 \| \mathbf{H}^{-\top} \|_2$$

$$=: \Delta_{31}^\circ + \Delta_{32}^\circ$$

By Assumptions 2.3 and Lemma 2.1, it follows that

$$\sum_{t=1}^n \mathbb{E} \left\| \sum_{s=1}^n \mathbb{E}_0(\mathbf{u}_t^\top \mathbf{u}_s) \mathbf{f}_s \right\|_2^2 \lesssim n \mathbb{E} \|\mathbb{E}_0(\mathbf{u}^\top \mathbf{u}) \mathbf{f}\|_2^2 + n^2 \mathbb{E} \|\mathbf{u}_1^\top \mathbf{u}_2 \mathbf{f}_2\|_2^2 \lesssim n^2 d.$$

Combining this with the fact that $\max_{j \in [d]} \|\mathbf{e}_j^\top \mathbf{U}^\top \mathbf{F}\|_2^2 = O_{\mathbb{P}}(n \log d)$, we obtain

$$\begin{aligned} \Delta_{31}^\circ &\leq \max_{j \in [d]} \frac{\text{tr}(\Sigma)}{n} \left\| \sum_{t=1}^n u_{tj} \mathbf{f}_t \right\|_2 \|\mathbf{H}\|_2 \|\mathbf{V}^{-1}\|_2 \|\mathbf{H}^{-\top}\|_2 \\ &\quad + \max_{j \in [d]} \frac{1}{n} \left\| \sum_{t=1}^n u_{tj} \sum_{s=1}^n \mathbb{E}_0(\mathbf{u}_t^\top \mathbf{u}_s) \mathbf{f}_s \right\|_2 \|\mathbf{V}^{-1}\|_2 \|\mathbf{H}^{-\top}\|_2 \\ &= O_{\mathbb{P}} \left\{ \sqrt{(\log d)/n} + \sqrt{n/d} \right\}. \end{aligned}$$

Similarly, by Lemma 2.1, we have $\Delta_{32}^\circ = O_{\mathbb{P}}(n/d + 1/n + \sqrt{n/d})$. Putting all these pieces together, we obtain (C.3). $\square$

Lemma C.4. *Under Assumptions 2.1–2.5, we have*

$$\|(\hat{\mathbf{U}} - \mathbf{U})^\top \mathcal{E}\|_\infty = O_{\mathbb{P}} \left(\sqrt{\log d + \frac{n}{d}} \right).$$

Proof of Lemma C.4. Recall that $\hat{\mathbf{U}} = \mathbf{X} - \hat{\mathbf{F}} \hat{\mathbf{B}}^\top$. By the triangle inequality,

$$\begin{aligned} \|(\hat{\mathbf{U}} - \mathbf{U})^\top \mathcal{E}\|_\infty &\leq \|(\hat{\mathbf{B}} - \mathbf{B} \mathbf{H}^\top) \hat{\mathbf{F}}^\top \mathcal{E}\|_\infty + \|\mathbf{B} \mathbf{H}^\top (\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top)^\top \mathcal{E}\|_\infty \\ &\quad + \|\mathbf{B} (\mathbf{H}^\top \mathbf{H} - \mathbf{I}_K) \mathbf{F}^\top \mathcal{E}\|_\infty =: \Delta_1 + \Delta_2 + \Delta_3. \end{aligned}$$

By Lemma 2.1 and the Cauchy-Schwarz inequality, we have

$$\begin{aligned} \Delta_1 &\leq \max_{j \in [d]} \|\hat{\mathbf{b}}_j - \mathbf{H} \mathbf{b}_j\|_2 \|\hat{\mathbf{F}}^\top \mathcal{E}\|_2 = O_{\mathbb{P}} \left(\sqrt{\log d} \right) \\ \Delta_2 &\leq \max_{j \in [d]} \|\mathbf{b}_j\|_2 \|\mathbf{H}\|_2 \|(\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top)^\top \mathcal{E}\|_2 = O_{\mathbb{P}} \left(1/\sqrt{n} + \sqrt{n/d} \right) \\ \Delta_3 &\leq \max_{j \in [d]} \|\mathbf{b}_j\|_2 \|\mathbf{H}^\top \mathbf{H} - \mathbf{I}_K\|_{\mathbb{F}} \|\mathbf{F}^\top \mathcal{E}\|_2 = O_{\mathbb{P}} \left(1 + \sqrt{n/d} \right). \end{aligned}$$

$\square$

Lemma C.5. *Under Assumptions 2.1–2.5, we have*

$$\|\hat{\mathbf{U}}^\top (\tilde{\mathbf{Y}} - \hat{\mathbf{U}} \boldsymbol{\beta}^*)\|_\infty = O_{\mathbb{P}} \left(\sqrt{n \log d} + \mathcal{V}_{n,d} \|\boldsymbol{\varphi}^*\|_2 \right).$$

Proof of Lemma C.5. By Lemmas C.3 and C.4, we have $\|\hat{\mathbf{U}}^\top \mathbf{F} \boldsymbol{\varphi}^\star\|_\infty = O_{\mathbb{P}}(\mathcal{V}_{n,d} \|\boldsymbol{\varphi}^\star\|_2)$ and

$$\|\hat{\mathbf{U}}^\top \boldsymbol{\varepsilon}\|_\infty \leq \|(\hat{\mathbf{U}} - \mathbf{U})^\top \boldsymbol{\varepsilon}\|_\infty + \|\mathbf{U}^\top \boldsymbol{\varepsilon}\|_\infty = O_{\mathbb{P}}\left(\sqrt{n \log d}\right).$$

Consequently, as $\hat{\mathbf{U}}^\top (\tilde{\mathbf{Y}} - \hat{\mathbf{U}} \boldsymbol{\beta}^\star) = \hat{\mathbf{U}}^\top \boldsymbol{\varepsilon} + \hat{\mathbf{U}}^\top \mathbf{F} \boldsymbol{\varphi}^\star$, it follows that

$$\|\hat{\mathbf{U}}^\top (\tilde{\mathbf{Y}} - \hat{\mathbf{U}} \boldsymbol{\beta}^\star)\|_\infty \leq \|\hat{\mathbf{U}}^\top \boldsymbol{\varepsilon}\|_\infty + \|\hat{\mathbf{U}}^\top \mathbf{F} \boldsymbol{\varphi}^\star\|_\infty = O_{\mathbb{P}}\left(\sqrt{n \log d} + \mathcal{V}_{n,d} \|\boldsymbol{\varphi}^\star\|_2\right).$$

□

Lemma C.6. Under Assumptions 2.1–2.4, we have

$$\|\hat{\mathbf{U}}^\top \hat{\mathbf{U}} - \mathbf{U}^\top \mathbf{U}\|_{\max} = O_{\mathbb{P}}\left(\frac{n}{d} + \log d\right).$$

Proof of Lemma C.6. By the triangle inequality, we have

$$\|\hat{\mathbf{U}}^\top \hat{\mathbf{U}} - \mathbf{U}^\top \mathbf{U}\|_{\max} \leq 2\|\hat{\mathbf{U}}^\top (\hat{\mathbf{U}} - \mathbf{U})\|_{\max} + \|(\hat{\mathbf{U}} - \mathbf{U})^\top (\hat{\mathbf{U}} - \mathbf{U})\|_{\max}.$$

Recall that $\mathbf{X} = \mathbf{F} \mathbf{B}^\top + \mathbf{U} = \hat{\mathbf{F}} \hat{\mathbf{B}}^\top + \hat{\mathbf{U}}$ and $\hat{\mathbf{F}}^\top \hat{\mathbf{U}} = \mathbf{O}$. Hence, it follows from Lemma C.3 that

$$\|\hat{\mathbf{U}}^\top (\hat{\mathbf{U}} - \mathbf{U})\|_{\max} = \|\hat{\mathbf{U}}^\top \mathbf{F} \mathbf{B}^\top\|_{\max} = \max_{j \in [d]} \|\hat{\mathbf{U}} \mathbf{F} \mathbf{b}_j\|_\infty = O_{\mathbb{P}}(\mathcal{V}_{n,d}) = O_{\mathbb{P}}\left(\frac{n}{d} + \log d\right).$$

Similarly, by Lemma 2.1,

$$\|(\hat{\mathbf{U}} - \mathbf{U})^\top (\hat{\mathbf{U}} - \mathbf{U})\|_{\max} \leq \max_{j \in [d]} \sum_{t=1}^n |\hat{u}_{tj} - u_{tj}|^2 = O_{\mathbb{P}}\left(\frac{n}{d} + \log d\right).$$

□

D Proof of Results in Section 2

D.1 Proof of Theorem 2.2

Proof of Theorem 2.2. Recall that $\hat{\mathbf{F}}^\top \hat{\mathbf{U}} = \mathbf{O}$ and $\hat{\boldsymbol{\gamma}} = n^{-1} \hat{\mathbf{F}}^\top \mathbf{Y}$. Hence

$$\hat{\boldsymbol{\gamma}} - \mathbf{H} \boldsymbol{\gamma}^\star = \frac{1}{n} \hat{\mathbf{F}}^\top \boldsymbol{\varepsilon} + \frac{1}{n} \hat{\mathbf{F}}^\top (\mathbf{F} - \hat{\mathbf{F}} \mathbf{H}) \boldsymbol{\varphi}^\star + (\hat{\mathbf{B}}^\top - \mathbf{H} \mathbf{B}^\top) \boldsymbol{\beta}^\star.$$

By Proposition 2.1, we have

$$\|\hat{\mathbf{F}}^\top (\mathbf{F} - \hat{\mathbf{F}} \mathbf{H}) \boldsymbol{\varphi}^\star\|_2 \leq \|\hat{\mathbf{F}}\|_{\mathbb{F}} \|\mathbf{F} - \hat{\mathbf{F}} \mathbf{H}\|_{\mathbb{F}} \|\boldsymbol{\varphi}^\star\|_2 = O_{\mathbb{P}}\left\{\left(\sqrt{n} + \frac{n}{\sqrt{d}}\right) \|\boldsymbol{\varphi}^\star\|_2\right\},$$

$$\|(\hat{\mathbf{B}} - \mathbf{B}\mathbf{H}^\top)^\top \boldsymbol{\beta}^\star\|_2 \leq \max_{j \in \mathcal{S}_\star} \|\hat{\mathbf{b}}_j - \mathbf{H}\mathbf{b}_j\|_2 \|\boldsymbol{\beta}^\star\|_1 = O_{\mathbb{P}} \left(\|\boldsymbol{\beta}^\star\|_1 \sqrt{\frac{\log |\mathcal{S}_\star|}{n} + \frac{1}{d}} \right).$$

Combined with the fact that $\|\hat{\mathbf{F}}^\top \mathcal{E}\|_2 = O_{\mathbb{P}}(\sqrt{n})$, we obtain the bound for $\|\hat{\boldsymbol{\gamma}} - \mathbf{H}\boldsymbol{\gamma}^\star\|_2$ by the triangle inequality. We now bound $\|\hat{\boldsymbol{\beta}}_\lambda - \boldsymbol{\beta}^\star\|_2$. Applying Lemma C.1 with Lemmas C.5 and C.2, we obtain (2.6). $\square$

D.2 Proof of Proposition 2.3

Proof. We introduce several key ingredients for proving Proposition 2.3 in the following several Lemmas. First, the following Lemma D.1 provides upper bounds for several key terms. Here we let $\hat{\boldsymbol{\nu}}_t = (\hat{\mathbf{u}}_t^\top, \hat{\mathbf{f}}_t^\top)^\top \in \mathbb{R}^{d+K}$.

Lemma D.1. *Under Assumptions 2.1–2.4, we have*

$$\max_{t \in [n]} \|\hat{\boldsymbol{\nu}}_t\|_\infty = O_{\mathbb{P}} \left(\sqrt{\log d} \right) \quad \text{and} \quad \max_{t \in [n]} |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^\star| = O_{\mathbb{P}} \left(\frac{\log n}{n + \sqrt{d}} \|\boldsymbol{\varphi}^\star\|_2 \right).$$

Furthermore, under the assumption, for any positive constant $\tau < \infty$, we have

$$\max_{j \in [d]} \sum_{t=1}^n \mathbb{I}\{|\varepsilon_t| > \tau \omega\} |\hat{\nu}_{tj}|^2 = O_{\mathbb{P}}(\log d).$$

Second, we provide an upper bound for the first order derivative of our loss function $\rho_\omega(\cdot)$ evaluated at the global optimizer. In the following, we let $\psi_\omega(\cdot)$ be the first-order derivative function of $\rho_\omega(\cdot)$ and let $\mathcal{S}_\diamond = \mathcal{S}_\star \cup \{d+1, \dots, d+K\}$. Recall that $\hat{\boldsymbol{\phi}}_h = (\hat{\boldsymbol{\beta}}_h^\top, \hat{\boldsymbol{\gamma}}_h^\top)^\top \in \mathbb{R}^{d+K}$ and $\tilde{\boldsymbol{\phi}} = (\boldsymbol{\beta}^{\star\top}, \tilde{\boldsymbol{\gamma}}^\top)^\top \in \mathbb{R}^{d+K}$, where $\tilde{\boldsymbol{\gamma}} = \hat{\mathbf{B}}^\top \boldsymbol{\beta}^\star + n^{-1} \hat{\mathbf{F}}^\top \mathbf{F} \boldsymbol{\varphi}^\star$. Hence we have $\mathbf{Y} - \hat{\mathbf{F}} \tilde{\boldsymbol{\gamma}} - \hat{\mathbf{U}} \boldsymbol{\beta}^\star = (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^\star + \mathcal{E}$.

Lemma D.2. *Assume that $n = O(d \log d)$ and $\mathbb{E}|\varepsilon|^{1+\vartheta} < \infty$ for some constant $\vartheta > 0$. Then, under Assumptions 2.1–2.4, we have*

$$\left\| \sum_{t=1}^n \psi_\omega(y_t - \hat{\boldsymbol{\nu}}_t^\top \tilde{\boldsymbol{\phi}}) \hat{\boldsymbol{\nu}}_t \right\|_\infty = O_{\mathbb{P}} \left(\mathcal{V}_{n,d} \|\boldsymbol{\varphi}^\star\|_2 + \sqrt{n \omega^{1-(\vartheta \wedge 1)} \log d} + \omega \log d \right).$$

Third, we prove the local strong convexity of our loss function $\rho_\omega(\cdot)$ in the following Lemma D.3.

Lemma D.3. *Let $\hat{\boldsymbol{\Gamma}}_n = n^{-1} \sum_{t=1}^n \mathbb{I}\{|\varepsilon_t| \leq \omega/3\} \hat{\boldsymbol{\nu}}_t \hat{\boldsymbol{\nu}}_t^\top$. Assume that $|\mathcal{S}_\diamond| \log d = o(n)$ and $\mathbb{E}|\varepsilon|^{1+\vartheta} < \infty$ for some constant $\vartheta > 0$. Then, under Assumptions 2.1–2.4, there exists a constant $\varrho_0 > 0$ such that*

$$\min_{\mathbf{0} \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{\mathbf{v}^\top \hat{\boldsymbol{\Gamma}}_n \mathbf{v}}{\|\mathbf{v}\|_2^2} \geq \varrho_0. \quad (\text{D.1})$$

For some constant $\kappa_0 > 0$, we define

$$\mathcal{E}_\Delta = \left\{ \max_{t \in [n]} |e_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^*| \leq \frac{\omega}{3} \right\} \quad \text{and} \quad \mathcal{E}_\nu = \left\{ \max_{t \in [n]} \|\hat{\boldsymbol{\nu}}_t\|_\infty \leq \frac{\omega \varrho_0}{36(1 + \kappa_0)\lambda |\mathcal{S}_\diamond|} =: M_\nu \right\}.$$

Next, we prove that $\mathbb{P}(\mathcal{E}_\Delta) \rightarrow 1$ and $\mathbb{P}(\mathcal{E}_\nu) \rightarrow 1$. By Lemma D.1 and (2.8), we have

$$\max_{t \in [n]} |e_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^*| = O_{\mathbb{P}} \left(\frac{\log n}{n + \sqrt{d}} \|\boldsymbol{\varphi}^*\|_2 \right) = o_{\mathbb{P}}(\omega).$$

Therefore $\mathbb{P}(\mathcal{E}_\Delta) \rightarrow 1$. As $\max_{t \in [n]} \|\hat{\boldsymbol{\nu}}_t\|_\infty = O_{\mathbb{P}}(\sqrt{\log d})$ and $|\mathcal{S}_\diamond|(\log d)^{3/2} = o(n)$, then

$$\frac{\max_{t \in [n]} \|\hat{\boldsymbol{\nu}}_t\|_\infty}{M_\nu} = O_{\mathbb{P}} \left(\frac{\lambda |\mathcal{S}_\diamond| \sqrt{\log d}}{\omega} \right) = o_{\mathbb{P}}(1).$$

Hence $\mathbb{P}(\mathcal{E}_\nu) \rightarrow 1$.

By the choice of ω , we have

$$\frac{2}{n} \left\| \sum_{t=1}^n \psi_\omega \left(y_t - \hat{\boldsymbol{\nu}}_t^\top \tilde{\boldsymbol{\phi}} \right) \hat{\boldsymbol{\nu}}_t \right\|_\infty = O_{\mathbb{P}} \left(\mathcal{V}_{n,d} \|\boldsymbol{\varphi}^*\|_2 + \sqrt{n\omega^{1-(\vartheta \wedge 1)} \log d} + \omega \log d \right) = O_{\mathbb{P}}(\omega \log d).$$

Combining all conclusions given above, we conclude the proof of Proposition 2.3 by the following Lemma D.4.

Lemma D.4. Assume that $\mathbb{E}|\varepsilon|^{1+\vartheta} < \infty$ for some constant $\vartheta > 0$, and $\max_{t \in [n]} \|\hat{\boldsymbol{\nu}}_t\|_\infty \leq M_\nu$ for some $M_\nu \geq 0$. Let

$$\lambda \geq \frac{2}{n} \left\| \sum_{t=1}^n \psi_\omega \left(y_t - \hat{\boldsymbol{\nu}}_t^\top \tilde{\boldsymbol{\phi}} \right) \hat{\boldsymbol{\nu}}_t \right\|_\infty \quad \text{and} \quad \omega \geq 3 \max_{t \in [n]} |e_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^*|. \quad (\text{D.2})$$

Furthermore, we assume that there exists a constant $\varrho_0 > 36\lambda M_\nu |\mathcal{S}_\diamond|/\omega$ such that

$$\min_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{\mathbf{v}^\top \hat{\mathbf{\Gamma}}_n \mathbf{v}}{\|\mathbf{v}\|_2^2} \geq \varrho_0. \quad (\text{D.3})$$

Then we have $\|\hat{\boldsymbol{\phi}}_h - \tilde{\boldsymbol{\phi}}\|_1 \leq 12\lambda |\mathcal{S}_\diamond|/\varrho_0$.

□

D.3 Proof of Lemma D.1

Proof. By Lemma 2.1,

$$\max_{t \in [n]} \|\hat{\mathbf{f}}_t\|_\infty \leq \max_{t \in [n]} \|\mathbf{H}\mathbf{f}_t\|_\infty + \max_{t \in [n]} \|\hat{\mathbf{f}}_t - \mathbf{H}\mathbf{f}_t\|_\infty = O_{\mathbb{P}} \left(\sqrt{\log n} + \frac{\log n}{n + \sqrt{d}} \right) = O_{\mathbb{P}} \left(\sqrt{\log n} \right).$$

Decompose $\mathbf{u}_t - \hat{\mathbf{u}}_t = (\hat{\mathbf{B}} - \mathbf{B}\mathbf{H}^{-1})(\hat{\mathbf{f}}_t - \mathbf{H}\mathbf{f}_t) + (\hat{\mathbf{B}} - \mathbf{B}\mathbf{H}^{-1})\mathbf{H}\mathbf{f}_t + \mathbf{B}\mathbf{H}^{-1}(\hat{\mathbf{f}}_t - \mathbf{H}\mathbf{f}_t)$. Then, by Lemma 2.1,

$$\begin{aligned} \max_{t \in [n]} \|\mathbf{u}_t - \hat{\mathbf{u}}_t\|_\infty &\leq K \|\hat{\mathbf{B}} - \mathbf{B}\mathbf{H}^{-1}\|_{\max} \|\hat{\mathbf{F}} - \mathbf{F}\mathbf{H}^\top\|_{\max} \\ &\quad + \sqrt{K} \|\hat{\mathbf{B}} - \mathbf{B}\mathbf{H}^{-1}\|_{\max} \|\mathbf{H}\|_2 \max_{t \in [n]} \|\mathbf{f}_t\|_2 \\ &\quad + K \|\mathbf{H}^{-1}\|_2 \|\mathbf{B}\|_{\max} \|\hat{\mathbf{F}} - \mathbf{F}\mathbf{H}^\top\|_{\max} \\ &= O_{\mathbb{P}} \left(\sqrt{(\log d)(\log n)/n} \right) = O_{\mathbb{P}} \left(\sqrt{\log d} \right). \end{aligned}$$

Consequently, by Assumption 2.1, we have

$$\max_{t \in [n]} \|\hat{\mathbf{u}}_t\|_\infty \leq \max_{t \in [n]} \|\mathbf{u}_t\|_\infty + \max_{t \in [n]} \|\hat{\mathbf{u}}_t - \mathbf{u}_t\|_\infty = O_{\mathbb{P}} \left(\sqrt{\log d} \right).$$

We now bound $\max_{t \in [n]} |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^*|$. Observe that $(\mathbf{I}_n - \hat{\mathbf{P}}) \hat{\mathbf{F}} = \mathbf{O}$. Hence

$$\begin{aligned} \max_{t \in [n]} |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^*| &\leq \max_{t \in [n]} |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) (\mathbf{F}\mathbf{H}^\top - \hat{\mathbf{F}}) \mathbf{H}^{-\top} \boldsymbol{\varphi}^*| \\ &\leq \max_{t \in [n]} |(\hat{\mathbf{f}}_t - \mathbf{H}\mathbf{f}_t)^\top \mathbf{H}^{-\top} \boldsymbol{\varphi}^*| + \max_{t \in [n]} \frac{1}{n} |\hat{\mathbf{f}}_t^\top \hat{\mathbf{F}}^\top (\mathbf{F}\mathbf{H}^\top - \hat{\mathbf{F}}) \mathbf{H}^{-\top} \boldsymbol{\varphi}^*| \\ &\leq \max_{t \in [n]} \|\hat{\mathbf{f}}_t - \mathbf{H}\mathbf{f}_t\|_2 \|\mathbf{H}^{-\top}\|_2 \|\boldsymbol{\varphi}^*\|_2 + \max_{t \in [n]} \frac{1}{n} \|\hat{\mathbf{f}}_t\|_2 \|\hat{\mathbf{F}}\|_{\mathbb{F}} \|\mathbf{F}\mathbf{H}^\top - \hat{\mathbf{F}}\|_{\mathbb{F}} \|\mathbf{H}^{-\top}\|_2 \|\boldsymbol{\varphi}^*\|_2 \\ &= O_{\mathbb{P}} \left(\frac{\log n}{n + \sqrt{d}} \|\boldsymbol{\varphi}^*\|_2 + \frac{1}{n} \sqrt{\log n} \sqrt{n} \sqrt{\frac{1}{n} + \frac{n}{d}} \|\boldsymbol{\varphi}^*\|_2 \right) = O_{\mathbb{P}} \left(\frac{\log n}{n + \sqrt{d}} \|\boldsymbol{\varphi}^*\|_2 \right). \end{aligned}$$

Recall that $n^{-1} \hat{\mathbf{F}}^\top \hat{\mathbf{F}} = \mathbf{I}_K$. Hence

$$\mathbb{E} \left(\max_{k \in [K]} \sum_{t=1}^n \mathbb{I}\{|\varepsilon_t| > \tau\omega\} |\hat{f}_{tk}|^2 \right) \leq \mathbb{E} \left(\sum_{t=1}^n \mathbb{I}\{|\varepsilon_t| > \tau\omega\} \|\hat{\mathbf{f}}_t\|_2^2 \right) \leq \frac{nK\mathbb{E}|\varepsilon|^{1+(\vartheta \wedge 1)}}{(\tau\omega)^{1+(\vartheta \wedge 1)}}.$$

Then it suffices to bound $\max_{j \in [d]} \sum_{t=1}^n \sum_{i=1}^n \mathbb{I}\{|\varepsilon_t| > \tau\omega\} u_{ij}^2$ as $\max_{j \in [d]} \sum_{t=1}^n (\hat{u}_{tj} - u_{tj})^2 = O_{\mathbb{P}}(\log d + n/d)$. By Assumption 2.1, we have $\mathbb{E} \exp(u_{ij}^2/c_0^2) \leq 2$ uniformly for $j \in [d]$. By Jensen's inequality and

$$\mathbb{E} \left(\max_{j \in [d]} \frac{1}{c_0^2} \sum_{t=1}^n \mathbb{I}\{|\varepsilon_t| > \tau\omega\} u_{tj}^2 \right) \leq \log \left\{ \sum_{j=1}^d \mathbb{E} \exp \left(\frac{1}{c_0^2} \sum_{t=1}^n \mathbb{I}\{|\varepsilon_t| > \tau\omega\} u_{tj}^2 \right) \right\}$$

$$\begin{aligned}
&\leq \log \left[d \max_{j \in [d]} \left\{ \mathbb{E} \exp \left(\frac{u_{tj}^2}{c_0^2} \right) \mathbb{P}(|\varepsilon| > \tau\omega) + \mathbb{P}(|\varepsilon| \leq \tau\omega) \right\}^n \right] \\
&\leq \log d + \frac{n \mathbb{E}|\varepsilon|^{1+(\vartheta \wedge 1)}}{(\tau\omega)^{1+(\vartheta \wedge 1)}}.
\end{aligned}$$

□

D.4 Proof of Lemma D.2

Proof. Recall that $\mathbf{Y} - \hat{\mathbf{U}}\beta^\star - \hat{\mathbf{F}}\tilde{\gamma} = \mathcal{E} + (\mathbf{I}_n - \hat{\mathbf{P}})\mathbf{F}\varphi^\star$. Hence, by Taylor's formula,

$$\begin{aligned}
\sum_{t=1}^n \psi_\omega(y_t - \hat{\nu}_t^\top \tilde{\phi}) \hat{\nu}_t &= \sum_{t=1}^n \psi_\omega(\varepsilon_t)(\hat{\nu}_t - \tilde{\nu}_t) + \sum_{t=1}^n \psi_\omega(\varepsilon_t) \tilde{\nu}_t \\
&\quad + \sum_{t=1}^n \mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star \hat{\nu}_t \int_0^1 \mathbb{I} \left\{ |\varepsilon_t + s \mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star| \leq \omega \right\} ds \\
&=: \Delta_1 + \Delta_2 + \Delta_3,
\end{aligned}$$

where $\tilde{\nu}_t = (\mathbf{u}_t^\top, \mathbf{f}_t^\top \mathbf{H}^\top)^\top \in \mathbb{R}^{d+K}$. Note that $\mathbb{E}|\psi_\omega(\varepsilon)|^2 \leq \mathbb{E}|\varepsilon|^{1+(\vartheta \wedge 1)} \omega^{1-(\vartheta \wedge 1)}$. Hence, by Lemma 2.1 and the Cauchy-Schwarz inequality,

$$\|\Delta_1\|_\infty^2 \leq \sum_{t=1}^n \{\psi_\omega(\varepsilon_t)\}^2 \max_{1 \leq j \leq d+K} \sum_{t=1}^n (\hat{\nu}_{tj} - \nu_{tj})^2 = O_{\mathbb{P}} \left\{ n \omega^{1-(\vartheta \wedge 1)} \left(\log d + \frac{n}{d} \right) \right\}.$$

By Lemma C.6 in Sun et al. [2020] and Lemma 2.1, we have

$$\|\Delta_2\|_\infty = O_{\mathbb{P}} \left(\sqrt{n \omega^{1-(\vartheta \wedge 1)} \log d} + \omega \log d \right).$$

Decompose

$$\begin{aligned}
\Delta_3 &= \sum_{t=1}^n \mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star \hat{\nu}_t - \sum_{t=1}^n \mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star \hat{\nu}_t \int_0^1 \mathbb{I} \left\{ |\varepsilon_t + s \mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star| > \omega \right\} ds \\
&=: \Delta_{31} + \Delta_{32}.
\end{aligned}$$

Observe that $\sum_{t=1}^n \mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star \hat{\mathbf{f}}_t = \hat{\mathbf{F}}^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star = 0$. Hence, by Lemma C.3,

$$\|\Delta_{31}\|_\infty = \|\hat{\mathbf{U}}^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star\|_\infty = O_{\mathbb{P}}(\mathcal{V}_{n,d} \|\varphi^\star\|_2).$$

Since

$$\sum_{t=1}^n |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \varphi^\star|^2 \leq \|\mathbf{F} \mathbf{H}^\top - \hat{\mathbf{F}}\|_{\mathbb{F}}^2 \|\mathbf{H}^{-\top}\|_2^2 \|\varphi^\star\|_2^2 = O_{\mathbb{P}} \left\{ \left(\frac{1}{n} + \frac{n}{d} \right) \|\varphi^\star\|_2^2 \right\}.$$

Combined with the fact that $\mathbb{P}(|\varepsilon_t| > 2\omega/3) \lesssim (\log d)/n$, we obtain

$$\Delta_{32}^\diamond := \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| > \frac{2\omega}{3} \right\} |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^*|^2 = o_{\mathbb{P}} \left(\frac{\mathcal{V}_{n,d}^2}{\log d} \right).$$

Consequently, by the Cauchy-Schwarz inequality and Lemma D.1,

$$\|\Delta_{32}\|_\infty \leq \left(\Delta_{32}^\diamond \max_{k \in [K]} \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| > \frac{2\omega}{3} \right\} \hat{\nu}_{tj}^2 \right)^{1/2} = o_{\mathbb{P}}(\mathcal{V}_{n,d} \|\boldsymbol{\varphi}^*\|_2).$$

□

D.5 Proof of Lemma D.3

Proof. For simplicity of notation, write

$$\tilde{\Gamma}_n = \frac{1}{n} \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| \leq \frac{\omega}{3} \right\} \tilde{\boldsymbol{\nu}}_t \tilde{\boldsymbol{\nu}}_t^\top \text{ and } \Gamma_n = \frac{1}{n} \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| \leq \frac{\omega}{3} \right\} \boldsymbol{\nu}_t \boldsymbol{\nu}_t^\top,$$

where $\boldsymbol{\nu}_t = (\mathbf{u}_t^\top, \mathbf{f}_t^\top)^\top \in \mathbb{R}^{d+K}$. For any $\mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)$, we write $\mathbf{v} = (\mathbf{v}_{[d]}^\top, \mathbf{v}_{[-d]}^\top)^\top \in \mathbb{R}^{d+K}$, where $\mathbf{v}_{[d]} \in \mathbb{R}^d$ and $\mathbf{v}_{[-d]} \in \mathbb{R}^K$. We first bound $\sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} |\mathbf{v}^\top (\hat{\Gamma}_n - \tilde{\Gamma}_n) \mathbf{v}|$. Decompose

$$\begin{aligned} \mathbf{v}^\top (\hat{\Gamma}_n - \tilde{\Gamma}_n) \mathbf{v} &= \frac{1}{n} \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| \leq \frac{\omega}{3} \right\} \{ (\hat{\mathbf{u}}_t^\top \mathbf{v}_{[d]})^2 - (\mathbf{u}_t^\top \mathbf{v}_{[d]})^2 \} \\ &\quad + \frac{2}{n} \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| \leq \frac{\omega}{3} \right\} \left(\mathbf{v}_{[d]}^\top \hat{\mathbf{u}}_t \hat{\mathbf{f}}_t^\top \mathbf{v}_{[-d]} - \mathbf{v}_{[d]}^\top \mathbf{u}_t \mathbf{f}_t^\top \mathbf{H}^\top \mathbf{v}_{[-d]} \right) \\ &\quad + \frac{1}{n} \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| \leq \frac{\omega}{3} \right\} \{ (\hat{\mathbf{f}}_t^\top \mathbf{v}_{[-d]})^2 - (\mathbf{f}_t^\top \mathbf{H}^\top \mathbf{v}_{[-d]})^2 \} \\ &=: \frac{1}{n} \{ \Delta_1(\mathbf{v}) + 2\Delta_2(\mathbf{v}) + \Delta_3(\mathbf{v}) \}. \end{aligned}$$

Let $\mathbf{D} = \text{diag}\{D_{11}, \dots, D_{nn}\} \in \mathbb{R}^{n \times n}$ denote an identity matrix, where $D_{tt} = \mathbb{I}\{|\varepsilon_t| \leq \omega/3\}$ for each $t \in [n]$. Then, for any $\mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)$,

$$\begin{aligned} \Delta_1(\mathbf{v}) &= \mathbf{v}_{[d]}^\top \left(\hat{\mathbf{U}}^\top \mathbf{D} \hat{\mathbf{U}} - \mathbf{U}^\top \mathbf{D} \mathbf{U} \right) \mathbf{v}_{[d]} \\ &= \mathbf{v}_{[d]}^\top \left(\hat{\mathbf{U}}^\top \hat{\mathbf{U}} - \mathbf{U}^\top \mathbf{U} \right) \mathbf{v}_{[d]} + \mathbf{v}_{[d]}^\top \left(\hat{\mathbf{U}} - \mathbf{U} \right)^\top \bar{\mathbf{D}} \mathbf{U} \mathbf{v}_{[d]} + \mathbf{v}_{[d]}^\top \hat{\mathbf{U}}^\top \bar{\mathbf{D}} \left(\hat{\mathbf{U}} - \mathbf{U} \right) \mathbf{v}_{[d]} \\ &=: \Delta_{11}(\mathbf{v}) + \Delta_{12}(\mathbf{v}) + \Delta_{13}(\mathbf{v}). \end{aligned}$$

where $\bar{D} = \text{diag}\{1 - D_{11}, \dots, 1 - D_{nn}\} \in \mathbb{R}^{n \times n}$. By Lemma C.6, we have

$$\begin{aligned} \sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{|\Delta_{11}(\mathbf{v})|}{\|\mathbf{v}\|_2^2} &\leq \sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{\|\hat{\mathbf{U}}^\top \hat{\mathbf{U}} - \mathbf{U}^\top \mathbf{U}\|_{\max} \|\mathbf{v}_{[d]}\|_1^2}{\|\mathbf{v}\|_2^2} \\ &\leq 32|\mathcal{S}_\diamond| \|\hat{\mathbf{U}}^\top \hat{\mathbf{U}} - \mathbf{U}^\top \mathbf{U}\|_{\max} = O_{\mathbb{P}}(|\mathcal{S}_\diamond| \log d). \end{aligned}$$

By the Cauchy-Schwarz inequality and Lemma D.1,

$$\begin{aligned} \sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{|\Delta_{12}(\mathbf{v})|}{\|\mathbf{v}\|_2^2} &\leq 32|\mathcal{S}_\diamond| \max_{1 \leq j, \ell \leq d} \left| \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| > \frac{\omega}{3} \right\} \hat{u}_{tj} (\hat{u}_{t\ell} - u_{t\ell}) \right| \\ &\leq 32|\mathcal{S}_\diamond| \max_{1 \leq j, \ell \leq d} \left[\sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| > \frac{\omega}{3} \right\} u_{tj}^2 \sum_{s=1}^n \mathbb{I} \left\{ |\varepsilon_t| > \frac{\omega}{3} \right\} (\hat{u}_{t\ell} - u_{t\ell})^2 \right]^{1/2} \\ &= O_{\mathbb{P}} \left\{ |\mathcal{S}_\diamond| (\log d)^{1/2} \left(\log d + \frac{n}{d} \right)^{1/2} \right\} = O_{\mathbb{P}}(|\mathcal{S}_\diamond| \log d). \end{aligned}$$

Similarly, we have $\sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \{|\Delta_{13}(\mathbf{v})|/\|\mathbf{v}\|_2^2\} = O_{\mathbb{P}}(|\mathcal{S}_\diamond| \log d)$. We now upper bound $|\Delta_2(\mathbf{v})|$. Decompose

$$\begin{aligned} \Delta_2(\mathbf{v}) &= \mathbf{v}_{[d]}^\top \left(\hat{\mathbf{U}}^\top \mathbf{D} \hat{\mathbf{F}} - \mathbf{U}^\top \mathbf{D} \mathbf{F} \mathbf{H}^\top \right) \mathbf{v}_{[-d]} \\ &= \mathbf{v}_{[d]}^\top \mathbf{U}^\top \mathbf{F} \mathbf{H}^\top \mathbf{v}_{[-d]} + \mathbf{v}_{[d]}^\top \left(\hat{\mathbf{U}} - \mathbf{U} \right) \bar{\mathbf{D}} \mathbf{F} \mathbf{H}^\top \mathbf{v}_{[-d]} + \mathbf{v}_{[d]}^\top \hat{\mathbf{U}}^\top \bar{\mathbf{D}} \left(\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top \right) \mathbf{v}_{[-d]} \\ &=: \Delta_{21}(\mathbf{v}) + \Delta_{22}(\mathbf{v}) + \Delta_{23}(\mathbf{v}). \end{aligned}$$

By Assumption 2.1 and the Cauchy-Schwarz inequality, it follows that for any $0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)$,

$$\begin{aligned} \frac{|\Delta_{21}(\mathbf{v})|}{\|\mathbf{v}\|_2^2} &\leq \frac{\|\mathbf{v}_{[d]}\|_1 \|\mathbf{U}^\top \mathbf{F} \mathbf{H}^\top \mathbf{v}_{[-d]}\|_\infty}{\|\mathbf{v}\|_2^2} \\ &\leq \max_{j \in [d]} \left\| \sum_{t=1}^n u_{tj} \mathbf{f}_t \right\|_2 \frac{\|\mathbf{H}^\top\|_2 (4\sqrt{|\mathcal{S}_\star|} \|\mathbf{v}_{\mathcal{S}_\star}\|_2 + 3\sqrt{K} \|\mathbf{v}_{[-d]}\|_2) \|\mathbf{v}_{[-d]}\|_2}{\|\mathbf{v}\|_2^2} \\ &\leq 7 \max_{j \in [d]} \left\| \sum_{t=1}^n u_{tj} \mathbf{f}_t \right\|_2 \|\mathbf{H}^\top\|_2 \sqrt{|\mathcal{S}_\diamond|} = O_{\mathbb{P}} \left(\sqrt{n|\mathcal{S}_\diamond| \log d} \right). \end{aligned}$$

Similarly, we have

$$\begin{aligned} \sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{|\Delta_{22}(\mathbf{v})|}{\|\mathbf{v}\|_2^2} &\leq 7 \max_{j \in [d]} \left\| \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| > \frac{\omega}{3} \right\} (\hat{u}_{tj} - u_{tj}) \mathbf{f}_t \right\|_2 \|\mathbf{H}^\top\|_2 \sqrt{|\mathcal{S}_\diamond|} \\ &= O_{\mathbb{P}} \left\{ \sqrt{|\mathcal{S}_\diamond| (\log d) \left(\log d + \frac{n}{d} \right)} \right\} = o_{\mathbb{P}} \left(\sqrt{n|\mathcal{S}_\diamond| \log d} \right), \end{aligned}$$

and $\sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \{|\Delta_{23}(\mathbf{v})|/\|\mathbf{v}\|_2^2\} = o_{\mathbb{P}}(\sqrt{n|\mathcal{S}_\diamond| \log d})$. Therefore $\sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \{|\Delta_2(\mathbf{v})|/\|\mathbf{v}\|_2^2\} = O_{\mathbb{P}}(\sqrt{n|\mathcal{S}_\diamond| \log d})$. Decompose

$$\begin{aligned}\Delta_3(\mathbf{v}) &= \mathbf{v}_{[-d]}^\top \left(\hat{\mathbf{F}}^\top \mathbf{D} \hat{\mathbf{F}} - \mathbf{H} \mathbf{F}^\top \mathbf{D} \mathbf{F} \mathbf{H}^\top \right) \mathbf{v}_{[-d]} \\ &= \mathbf{v}_{[-d]}^\top (n\mathbf{I}_K - \mathbf{H} \mathbf{F}^\top \mathbf{F} \mathbf{H}^\top) \mathbf{v}_{[-d]} + \mathbf{v}_{[-d]}^\top \left(\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top \right)^\top \bar{\mathbf{D}} \mathbf{F} \mathbf{H}^\top \mathbf{v}_{[-d]} \\ &\quad + \mathbf{v}_{[-d]}^\top \hat{\mathbf{F}}^\top \bar{\mathbf{D}} \left(\hat{\mathbf{F}} - \mathbf{F} \mathbf{H}^\top \right) \mathbf{v}_{[-d]} \\ &=: \Delta_{31}(\mathbf{v}) + \Delta_{32}(\mathbf{v}) + \Delta_{33}(\mathbf{v}).\end{aligned}$$

By Lemma 2.1, we have

$$\begin{aligned}\sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{|\Delta_{31}(\mathbf{v})|}{\|\mathbf{v}\|_2^2} &\leq \sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{\|n\mathbf{I}_K - \mathbf{H} \mathbf{F}^\top \mathbf{F} \mathbf{H}^\top\|_2 \|\mathbf{v}_{[-d]}\|_2^2}{\|\mathbf{v}\|_2^2} \\ &\leq n\|\mathbf{I}_K - \mathbf{H} \mathbf{H}^\top\|_2 + \|\mathbf{H}\|_2 \|n\mathbf{I}_K - \mathbf{F}^\top \mathbf{F}\|_2 \|\mathbf{H}^\top\|_2 \\ &= O_{\mathbb{P}}\left(\sqrt{n} + n/\sqrt{d}\right), \\ \sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{|\Delta_{32}(\mathbf{v})|}{\|\mathbf{v}\|_2^2} &\leq \left\| \sum_{t=1}^n \mathbb{I}\left\{|\varepsilon_t| > \frac{\omega}{3}\right\} (\hat{\mathbf{f}}_t - \mathbf{H} \mathbf{f}_t) \mathbf{f}_t^\top \right\|_2 \|\mathbf{H}^\top\|_2 \\ &\leq \left(\sum_{t=1}^n \mathbb{I}\left\{|\varepsilon_t| > \frac{\omega}{3}\right\} \|\hat{\mathbf{f}}_t - \mathbf{H} \mathbf{f}_t\|_2^2 \sum_{t=1}^n \mathbb{I}\left\{|\varepsilon_t| > \frac{\omega}{3}\right\} \|\mathbf{f}_t\|_2^2 \right)^{1/2} \|\mathbf{H}^\top\|_2 \\ &= o_{\mathbb{P}}\left(\sqrt{n} + n/\sqrt{d}\right),\end{aligned}$$

and $\sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \{|\Delta_{33}(\mathbf{v})|/\|\mathbf{v}\|_2^2\} = o_{\mathbb{P}}(\sqrt{n} + n/\sqrt{d})$. Putting all these pieces together, we obtain

$$\sup_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{|\mathbf{v}^\top (\hat{\Gamma}_n - \tilde{\Gamma}_n) \mathbf{v}|}{\|\mathbf{v}\|_2^2} = O_{\mathbb{P}}\left(\frac{|\mathcal{S}_\diamond| \log d}{n} + \sqrt{\frac{|\mathcal{S}_\diamond| \log d}{n}} + \frac{1}{\sqrt{n} + \sqrt{d}}\right) = o_{\mathbb{P}}(1).$$

For any $\mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)$, write $\tilde{\mathbf{v}} = (\mathbf{v}_{[d]}^\top, \mathbf{v}_{[-d]}^\top \mathbf{H})^\top \in \mathbb{R}^{d+K}$. Define $\mathcal{E}_H = \{\lambda_{\min}(\mathbf{H}^\top \mathbf{H}) \geq 1/4\}$. It follows from Lemma 2.1 that $\mathbb{P}(\mathcal{E}_H) \rightarrow 1$. Under $\mathcal{E}_H$, we have

$$\|\tilde{\mathbf{v}}_{\mathcal{S}_\diamond^c}\|_1 = \|\mathbf{v}_{\mathcal{S}_\diamond^c}\|_1 \leq 3\|\mathbf{v}_{\mathcal{S}_\star}\|_1 + 3\|\mathbf{v}_{[-d]}\|_1 \leq 3\|\mathbf{v}_{\mathcal{S}_\star}\|_1 + 6\sqrt{K}\|\mathbf{H}^\top \mathbf{v}_{[-d]}\|_1 \leq 6\sqrt{K}\|\tilde{\mathbf{v}}_{\mathcal{S}_\diamond}\|_1.$$

Combined with the fact that

$$\|\tilde{\mathbf{v}}\|_2^2 = \|\mathbf{v}_{[d]}\|_2^2 + \|\mathbf{H}^\top \mathbf{v}_{[-d]}\|_2^2 \geq \|\mathbf{v}_{[d]}\|_2^2 + \frac{1}{4}\|\mathbf{v}_{[-d]}\|_2^2 \geq \frac{1}{4}\|\mathbf{v}\|_2^2,$$

we obtain

$$\inf_{0 \neq \mathbf{v} \in \mathcal{C}(\mathcal{S}_\diamond, 3)} \frac{\mathbf{v}^\top \tilde{\Gamma}_n \mathbf{v}}{\|\mathbf{v}\|_2^2} \geq \inf_{0 \neq \tilde{\mathbf{v}} \in \mathcal{C}(\mathcal{S}_\diamond, 6\sqrt{K})} \frac{\tilde{\mathbf{v}}^\top \Gamma_n \tilde{\mathbf{v}}}{\|\tilde{\mathbf{v}}\|_2^2} \geq \inf_{0 \neq \tilde{\mathbf{v}} \in \mathcal{C}(\mathcal{S}_\diamond, 6\sqrt{K})} \frac{\tilde{\mathbf{v}}^\top \Gamma_n \tilde{\mathbf{v}}}{4\|\tilde{\mathbf{v}}\|_2^2}.$$

Observe that $\mathbb{I}\{|\varepsilon_t| \leq \omega/3\} \boldsymbol{\nu}_t \in \mathbb{R}^{d+K}$, $t = 1, \dots, n$, are i.i.d. sub-Gaussian random vectors. Hence, similar to Lemma C.2, $\mathbb{P}(|\varepsilon| \leq \omega/3) \geq 1/2$ and $\lambda_{\min}\{\text{Cov}(\boldsymbol{\nu}_t)\} \geq (\lambda_{\min}(\boldsymbol{\Sigma}) \wedge 1)$, we obtain

$$\mathbb{P}\left\{\inf_{0 \neq \tilde{\mathbf{v}} \in \mathcal{C}(\mathcal{S}_\circ, 6\sqrt{K})} \frac{\tilde{\mathbf{v}}^\top \mathbf{\Gamma}_n \tilde{\mathbf{v}}}{\|\tilde{\mathbf{v}}\|_2^2} \geq \frac{1}{8}(\lambda_{\min}(\boldsymbol{\Sigma}) \wedge 1)\right\} \rightarrow 1.$$

Putting all these pieces together, we obtain (D.1). $\square$

D.6 Proof of Lemma D.4

Proof. By the definition of $(\hat{\boldsymbol{\beta}}_h, \hat{\boldsymbol{\gamma}}_h)$, we have

$$\frac{1}{n} \sum_{t=1}^n \rho_\omega \left(y_t - \hat{\boldsymbol{\nu}}_t^\top \hat{\boldsymbol{\phi}}_h \right) + \lambda \|\hat{\boldsymbol{\beta}}_h\|_1 \leq \frac{1}{n} \sum_{t=1}^n \rho_\omega \left(y_t - \hat{\boldsymbol{\nu}}_t^\top \tilde{\boldsymbol{\phi}} \right) + \lambda \|\boldsymbol{\beta}^\star\|_1. \quad (\text{D.4})$$

Let $\boldsymbol{\delta}_\phi = \hat{\boldsymbol{\phi}}_h - \tilde{\boldsymbol{\phi}}$. Since $\rho_\omega(\cdot)$ is convex, it follows from (D.2) that

$$\lambda \|\boldsymbol{\beta}^\star\|_1 - \lambda \|\hat{\boldsymbol{\beta}}_h\|_1 \geq \frac{1}{n} \sum_{t=1}^n \psi_\omega \left(y_t - \hat{\boldsymbol{\nu}}_t^\top \tilde{\boldsymbol{\phi}} \right) \hat{\boldsymbol{\nu}}_t^\top \boldsymbol{\delta}_\phi \geq -\frac{\lambda}{2} \|\boldsymbol{\delta}_\phi\|_1.$$

Observe that $\|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ^c}\|_1 = \|\hat{\boldsymbol{\beta}}_{h, \mathcal{S}_\star^c}\|_1$ and $\|\boldsymbol{\beta}^\star\|_1 = \|\tilde{\boldsymbol{\phi}}_{\mathcal{S}_\star}\|_1$. Hence

$$\begin{aligned} 0 &\leq \lambda \|\tilde{\boldsymbol{\phi}}_{\mathcal{S}_\star}\|_1 - \lambda \|\hat{\boldsymbol{\phi}}_{h, \mathcal{S}_\star}\|_1 - \lambda \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ^c}\|_1 + \frac{\lambda}{2} \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ}\|_1 + \frac{\lambda}{2} \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ^c}\|_1 \\ &\leq \lambda \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ}\|_1 - \frac{\lambda}{2} \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ^c}\|_1 + \frac{\lambda}{2} \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ}\|_1 \\ &= \frac{3\lambda}{2} \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ}\|_1 - \frac{\lambda}{2} \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\circ^c}\|_1. \end{aligned}$$

Therefore $\boldsymbol{\delta}_\phi \in \mathcal{C}(\mathcal{S}_\circ, 3)$ and it follows from (D.3) that

$$\boldsymbol{\delta}_\phi^\top \hat{\mathbf{\Gamma}}_n \boldsymbol{\delta}_\phi \geq \varrho_0 \|\boldsymbol{\delta}_\phi\|_2^2. \quad (\text{D.5})$$

Since $\max_{t \in [n]} \|\hat{\boldsymbol{\nu}}_t\|_\infty \leq M_\nu$ and $\max_{t \in [n]} |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^\star| \leq \omega/3$, for each $t \in [n]$, we have

$$\begin{aligned} &\mathbb{I} \left\{ |\varepsilon_t + \mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^\star + s \hat{\boldsymbol{\nu}}_t^\top \boldsymbol{\delta}_\phi| \leq \omega \right\} \\ &\geq \mathbb{I} \left\{ |\varepsilon_t| + \max_{t \in [n]} |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^\star| + |s| \max_{t \in [n]} \|\hat{\boldsymbol{\nu}}_t\|_\infty \|\boldsymbol{\delta}_\phi\|_1 \leq \omega \right\} \\ &\geq \mathbb{I} \left\{ |\varepsilon_t| \leq \frac{\omega}{3} \right\} \mathbb{I} \left\{ |s| \leq \frac{\omega}{3M_\nu \|\boldsymbol{\delta}_\phi\|_1} \right\}. \end{aligned}$$

Combined with (D.4) and (D.5), we obtain

$$\begin{aligned}
\frac{3\lambda}{2} \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\diamond}\|_1 &\geq \frac{1}{n} \sum_{t=1}^n (\hat{\boldsymbol{\nu}}_t^\top \boldsymbol{\delta}_\phi)^2 \int_0^1 (1-s) \mathbb{I} \left\{ |\varepsilon_t + \mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}}) \mathbf{F} \boldsymbol{\varphi}^\star + s \hat{\boldsymbol{\nu}}_t^\top \boldsymbol{\delta}_\phi| \leq \omega \right\} ds \\
&\geq \frac{1}{n} \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| \leq \frac{\omega}{3} \right\} (\hat{\boldsymbol{\nu}}_t^\top \boldsymbol{\delta}_\phi)^2 \int_0^{1 \wedge \frac{\omega}{3M_\nu \|\boldsymbol{\delta}_\phi\|_1}} (1-s) ds \\
&\geq \frac{1}{2} \left(1 \wedge \frac{\omega}{3M_\nu \|\boldsymbol{\delta}_\phi\|_1} \right) \frac{1}{n} \sum_{t=1}^n \mathbb{I} \left\{ |\varepsilon_t| \leq \frac{\omega}{3} \right\} (\hat{\boldsymbol{\nu}}_t^\top \boldsymbol{\delta}_\phi)^2 \\
&\geq \frac{\varrho_0}{2} \left(1 \wedge \frac{\omega}{3M_\nu \|\boldsymbol{\delta}_\phi\|_1} \right) \|\boldsymbol{\delta}_\phi\|_2^2 \\
&\geq \frac{\varrho_0}{2} \left(1 \wedge \frac{\omega}{12M_\nu \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\diamond}\|_1} \right) \frac{\|\boldsymbol{\delta}_{\phi, \mathcal{S}_\diamond}\|_1^2}{|\mathcal{S}_\diamond|}.
\end{aligned}$$

If $12M_\nu \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\diamond}\|_1 \leq \omega$, then $\|\boldsymbol{\delta}_{\phi, \mathcal{S}_\diamond}\|_1 \leq 3\lambda |\mathcal{S}_\diamond| / \varrho_0$. Consequently, since $\omega \varrho_0 > 36\lambda M_\nu |\mathcal{S}_\diamond|$, it follows that

$$\|\boldsymbol{\delta}_{\phi, \mathcal{S}_\diamond}\|_1 \leq \min \left\{ \frac{\omega}{12M_\nu}, \frac{3\lambda |\mathcal{S}_\diamond|}{\varrho_0} \right\} = \frac{3\lambda |\mathcal{S}_\diamond|}{\varrho_0}.$$

Otherwise, if $12M_\nu \|\boldsymbol{\delta}_{\phi, \mathcal{S}_\diamond}\|_1 > \omega$, then $36\lambda M_\nu |\mathcal{S}_\diamond| \geq \omega \varrho_0$, which contradicts to the assumption that $\varrho_0 > 36\lambda M_\nu |\mathcal{S}_\diamond| / \omega$. $\square$

E Proof of Results in Section 3

E.1 Example with node-wise regression

Example E.1. In this example, we use node-wise regression in [van de Geer et al. \[2014\]](#) to estimate $\hat{\boldsymbol{\Theta}}$. More specifically, for each $j \in [d]$, we first get

$$\hat{\boldsymbol{\omega}}_j = \arg \min_{\boldsymbol{\omega} \in \mathbb{R}^{d-1}} \left\{ \frac{1}{2n} \sum_{t=1}^n |\hat{u}_{tj} - \boldsymbol{\omega}^\top \hat{\mathbf{u}}_{t,-j}|^2 + \lambda_j \|\boldsymbol{\omega}\|_1 \right\}$$

with components of $\hat{\boldsymbol{\omega}}_j = \{\hat{\omega}_{j\ell}; \ell = 1, 2, \dots, d, \ell \neq j\}$, where $\lambda_j > 0$ is a tuning parameter. We next obtain

$$\hat{\nu}_j^2 = \frac{1}{n} \sum_{t=1}^n |\hat{u}_{tj} - \hat{\boldsymbol{\omega}}_j^\top \hat{\mathbf{u}}_{t,-j}|^2 + \lambda_j \|\hat{\boldsymbol{\omega}}_j\|_1.$$

Then the estimator for $\boldsymbol{\Theta}$ is given by $\hat{\boldsymbol{\Theta}} = (\hat{\Theta}_{j\ell}) \in \mathbb{R}^{d \times d}$, where $\hat{\Theta}_{jj} = 1/\hat{\nu}_j^2$ and $\hat{\Theta}_{j\ell} = -\hat{\omega}_{j\ell}/\hat{\nu}_j^2$ for $j \neq \ell$.

We shall first impose a sparsity assumption on the columns of the precision matrix Θ . To be more specific, for each $k \in [K]$, we let $\mathcal{S}_j = \{\ell \neq j : \Theta_{j\ell} \neq 0\}$ denote the support of the j -th column of Θ . We then define the population parameter

$$\boldsymbol{\omega}_j^* = \arg \min_{\boldsymbol{\omega} \in \mathbb{R}^{d-1}} \mathbb{E}|u_j - \boldsymbol{\omega}^\top \mathbf{u}_{-j}|^2 \text{ and } \nu_j^2 = \mathbb{E}|u_j - \boldsymbol{\omega}_j^{*\top} \mathbf{u}_{-j}|^2.$$

In addition, for each $j \in [d]$, let $\tilde{\boldsymbol{\omega}}_j^*$ denote the d -dimensional vector with components $\tilde{\omega}_{jj}^* = 1$ and $\tilde{\omega}_{j\ell}^* = -\omega_{j\ell}^*$ for $\ell \neq j$.

The following proposition provides theoretical guarantees for the estimator $\hat{\Theta}$ constructed above.

Proposition E.1. *Let Assumptions 2.1–2.4 hold. Assume that*

$$\max_{j \in [d]} |\mathcal{S}_j| \left(\frac{\log d}{n} + \frac{1}{d} \right) \rightarrow 0 \text{ and } \max_{j \in [d]} \frac{\mathcal{V}_{n,d}}{n} \sqrt{|\mathcal{S}_j|} \|\mathbf{B}^\top \tilde{\boldsymbol{\omega}}_j^*\|_2 \rightarrow 0. \quad (\text{E.1})$$

Then, with suitably chosen $\lambda_j \asymp n^{-1} \|\hat{\mathbf{U}}_{-j}^\top \hat{\mathbf{U}} \tilde{\boldsymbol{\omega}}_j^\|_\infty$ uniformly for $j \in [d]$, we have*

$$\begin{aligned} \|\mathbf{I}_d - \hat{\Theta} \tilde{\Sigma}\|_{\max} &= O_{\mathbb{P}} \left(\sqrt{\frac{\log d}{n}} + \frac{1}{\sqrt{d}} + \max_{j \in [d]} \frac{\mathcal{V}_{n,d}}{n} \|\mathbf{B}^\top \tilde{\boldsymbol{\omega}}_j^*\|_2 \right), \\ \|\hat{\Theta} - \Theta\|_{\max} &= O_{\mathbb{P}} \left(\max_{j \in [d]} \sqrt{|\mathcal{S}_j| \left(\frac{\log d}{n} + \frac{1}{d} \right)} + \max_{j \in [d]} \frac{\mathcal{V}_{n,d}}{n} \sqrt{|\mathcal{S}_j|} \|\mathbf{B}^\top \tilde{\boldsymbol{\omega}}_j^*\|_2 \right), \\ \|\hat{\Theta} - \Theta\|_\infty &= O_{\mathbb{P}} \left(\max_{j \in [d]} |\mathcal{S}_j| \sqrt{\frac{\log d}{n} + \frac{1}{d}} + \max_{j \in [d]} \frac{\mathcal{V}_{n,d}}{n} |\mathcal{S}_j| \|\mathbf{B}^\top \tilde{\boldsymbol{\omega}}_j^*\|_2 \right). \end{aligned} \quad (\text{E.2})$$

E.2 Proof of Proposition E.1

Lemma E.2. *Under Assumptions 2.1–2.4, uniformly for $j \in [d]$, we have*

$$\|\hat{\mathbf{U}}_{-j}^\top \hat{\mathbf{U}} \tilde{\boldsymbol{\omega}}_j^*\|_\infty = O_{\mathbb{P}} \left(\sqrt{n \log d} + \frac{n}{\sqrt{d}} + \mathcal{V}_{n,d} \|\mathbf{B}^\top \tilde{\boldsymbol{\omega}}_j^*\|_2 \right).$$

Proof of Lemma E.2. By the triangle inequality, for each $j \in [d]$,

$$\|\hat{\mathbf{U}}_{-j}^\top \hat{\mathbf{U}} \tilde{\boldsymbol{\omega}}_j^*\|_\infty \leq \|\mathbf{U}_{-j}^\top \boldsymbol{\chi}_j\|_\infty + \|\hat{\mathbf{U}}_{-j}^\top (\hat{\mathbf{U}} - \mathbf{U}) \tilde{\boldsymbol{\omega}}_j^*\|_\infty + \|(\hat{\mathbf{U}} - \mathbf{U})^\top \boldsymbol{\chi}_j\|_\infty$$

where $\boldsymbol{\chi}_j = \mathbf{U} \tilde{\boldsymbol{\omega}}_j^*$. By the definition of $\tilde{\boldsymbol{\omega}}_j^*$ and Assumption 2.4, it follows that

$$\|\tilde{\boldsymbol{\omega}}_j^*\|_2^2 \leq \frac{\tilde{\boldsymbol{\omega}}_j^{*\top} \Sigma \tilde{\boldsymbol{\omega}}_j^*}{\lambda_{\min}(\Sigma)} = \frac{1}{\Theta_{jj} \lambda_{\min}(\Sigma)} \leq \frac{\lambda_{\max}(\Sigma)}{\lambda_{\min}(\Sigma)} \leq \frac{1}{\kappa^2}.$$

Hence $\{\mathbf{u}_t^\top \tilde{\boldsymbol{\omega}}_j^*\}_{t=1}^n$ are i.i.d. zero-mean sub-Gaussian random variables with $\|\mathbf{u}_t^\top \tilde{\boldsymbol{\omega}}_j^*\|_{\psi_2} \leq c_0/\kappa$ and for any $\ell \neq j$, $\{u_{t\ell} \mathbf{u}_t^\top \tilde{\boldsymbol{\omega}}_j^*\}_{t=1}^n$ are i.i.d. zero-mean sub-exponential random variables with $\max_{\ell \neq j} \|u_{t\ell} \mathbf{u}_t^\top \tilde{\boldsymbol{\omega}}_j^*\|_{\psi_1} \leq c_0^2/\kappa$. Consequently, by Lemmas 2.1 and C.3, we obtain $\|\mathbf{U}_{-j}^\top \boldsymbol{\chi}_j\|_\infty^2 = O_{\mathbb{P}}(n \log d)$, $\|\hat{\mathbf{U}}_{-j}^\top (\hat{\mathbf{U}} - \mathbf{U}) \tilde{\boldsymbol{\omega}}_j^*\|_\infty \leq \|\hat{\mathbf{U}}^\top \mathbf{F} \mathbf{B}^\top \tilde{\boldsymbol{\omega}}_j^*\|_\infty = O_{\mathbb{P}}(\mathcal{V}_{n,d} \|\mathbf{B}^\top \tilde{\boldsymbol{\omega}}_j^*\|_2)$ and

$$\|(\hat{\mathbf{U}} - \mathbf{U})^\top \boldsymbol{\chi}_j\|_\infty^2 \leq \|\boldsymbol{\chi}_j\|_2^2 \max_{j \in [d]} \sum_{t=1}^n |\hat{u}_{tj} - u_{tj}|^2 = O_{\mathbb{P}} \left(n \log d + \frac{n^2}{d} \right).$$

□

Proof of Proposition E.1. Following the proof of Lemma C.2, under (E.1) and Assumptions 2.1–2.4, there exists a positive constant $\tilde{\kappa}$ such that

$$\mathbb{P} \left(\min_{j \in [d]} \inf_{\mathbf{h} \in \mathbb{R}^{d-1}: \|\mathbf{h}_{\mathcal{S}_j^c}\|_1 \leq 3\|\mathbf{h}_{\mathcal{S}_j}\|_1} \frac{\mathbf{h}^\top \hat{\mathbf{U}}_{-j}^\top \hat{\mathbf{U}}_{-j} \mathbf{h}}{n \|\mathbf{h}\|_2^2} \geq \tilde{\kappa} \right) \rightarrow 1, \text{ as } n \rightarrow \infty.$$

Then, by Lemma C.1, we have $\|\hat{\boldsymbol{\omega}}_j - \boldsymbol{\omega}_j^*\|_1 = O_{\mathbb{P}}(\lambda_j |\mathcal{S}_j|)$ and $\|\hat{\boldsymbol{\omega}}_j - \boldsymbol{\omega}_j^*\|_2 = O_{\mathbb{P}}(\lambda_j \sqrt{|\mathcal{S}_j|})$ uniformly for all $j \in [d]$. Similar to the proof of Theorem 2.4 in van de Geer et al. [2014], we obtain $|\hat{\nu}_j^2 - \nu_j^2| \xrightarrow{\mathbb{P}} 0$ uniformly for $j \in [d]$. Putting all these pieces together, we obtain (E.2). □

E.3 Proof of Theorem 3.1

Proof of Theorem 3.1. We first derive a Gaussian approximation result for $n^{-1/2} \boldsymbol{\Theta} \mathbf{U}^\top \mathcal{E}$. To this end, we shall apply Theorem 2.1 in Chernozhukov et al. [2017] and verify the conditions therein. By Assumption 2.4, we have $\lambda_{\max}(\boldsymbol{\Sigma}) \leq \|\boldsymbol{\Sigma}\|_1 \leq 1/\kappa$. Hence $\min_{j \in [d]} \Theta_{jj} \geq \lambda_{\min}(\boldsymbol{\Theta}) \geq \kappa > 0$. Combined with Assumption 2.1, for each $j \in [d]$, $\{\boldsymbol{\Theta}_j^\top \mathbf{u}_t \varepsilon_t\}_{t=1}^n$ are i.i.d. zero-mean sub-exponential random variables with $\text{Cov}(\boldsymbol{\Theta}_j^\top \mathbf{u}_t \varepsilon_t) = \sigma^2 \Theta_{jj} \geq \sigma^2 \kappa$ and

$$\max_{j \in [d]} \|\boldsymbol{\Theta}_j^\top \mathbf{u}_t \varepsilon_t\|_{\psi_1} \leq c_1 \max_{j \in [d]} \|\boldsymbol{\Theta}_j^\top \mathbf{u}_t\|_{\psi_2} \leq \frac{c_0 c_1}{\kappa}.$$

Then, by Theorem 2.2 given in Chernozhukov et al. [2020],

$$\rho^{\natural} := \sup_{x>0} \left| \mathbb{P} \left(n^{-1/2} \|\boldsymbol{\Theta} \mathbf{U}^\top \mathcal{E}\|_\infty \leq x \right) - \mathbb{P}(\|\mathbf{Z}\|_\infty \leq x) \right| \rightarrow 0.$$

Next, we will quantify the difference between our test statistics and $n^{-1} \boldsymbol{\Theta} \mathbf{U}^\top \mathcal{E}$. We first introduce Lemma E.3 below.

Lemma E.3. Assume that $\mathcal{V}_{n,d} \|\boldsymbol{\varphi}^*\|_2 \lesssim \sqrt{n \log d}$. Then, under Assumption 3.1 and conditions of Theorem 2.2, we have

$$\begin{aligned} \|n(\tilde{\boldsymbol{\beta}}_\lambda - \boldsymbol{\beta}^*) - \hat{\boldsymbol{\Theta}} \hat{\mathbf{U}}^\top \mathcal{E}\|_\infty &= O_{\mathbb{P}} \left(\Lambda_{\max} |\mathcal{S}_*| \sqrt{n \log d} + \mathcal{V}_{n,d} \|\boldsymbol{\Theta}\|_\infty \|\boldsymbol{\varphi}^*\|_2 \right), \\ \|\hat{\boldsymbol{\Theta}} \hat{\mathbf{U}}^\top \mathcal{E} - \boldsymbol{\Theta} \mathbf{U}^\top \mathcal{E}\|_\infty &= O_{\mathbb{P}} \left(\|\boldsymbol{\Theta}\|_\infty \sqrt{\log d + \frac{n}{d}} + \Delta_\infty \sqrt{n \log d} \right). \end{aligned}$$

By Lemma E.3, we have $\sqrt{n}\|\tilde{\beta}_\lambda - \beta^\star - n^{-1}\Theta U^\top \mathcal{E}\|_\infty = O_{\mathbb{P}}(\Delta_{n,d})$, where

$$\Delta_{n,d} = \Lambda_{\max}|\mathcal{S}_\star|\sqrt{\log d} + \frac{\mathcal{V}_{n,d}\|\Theta\|_\infty\|\varphi^\star\|_2}{\sqrt{n}} + \|\Theta\|_\infty\sqrt{\frac{\log d}{n} + \frac{1}{d}} + \Delta_\infty\sqrt{\log d}.$$

It follows from (3.5) that $\Delta_{n,d}\sqrt{\log d} \rightarrow 0$. Taking $\tilde{\Delta}_{n,d} = \sqrt{\Delta_{n,d}}/(\log d)^{1/4}$, then we have

$$\tilde{\Delta}_{n,d}\sqrt{\log d} \rightarrow 0 \quad \text{and} \quad \mathbb{P}\left(\sqrt{n}\|\tilde{\beta}_\lambda - \beta^\star - n^{-1}\Theta U^\top \mathcal{E}\|_\infty > \tilde{\Delta}_{n,d}\right) \rightarrow 0.$$

Consequently, by Lemma in Chernozhukov et al. [2017] and Lemma E.3, we obtain

$$\begin{aligned} & \sup_{x>0} \left| \mathbb{P}\left(\sqrt{n}\|\tilde{\beta}_\lambda - \beta\|_\infty \leq x\right) - \mathbb{P}(n^{-1/2}\|\Theta U^\top \mathcal{E}\|_\infty \leq x) \right| \\ & \leq \mathbb{P}\left(\sqrt{n}\|\tilde{\beta}_\lambda - \beta^\star - n^{-1}\Theta U^\top \mathcal{E}\|_\infty > \tilde{\Delta}_{n,d}\right) + 2\rho^\natural + \sup_{x>0} \mathbb{P}(x < \|\mathbf{Z}\|_\infty \leq x + \tilde{\Delta}_{n,d}) \\ & \leq \mathbb{P}\left(\sqrt{n}\|\tilde{\beta}_\lambda - \beta^\star - n^{-1}\Theta U^\top \mathcal{E}\|_\infty > \tilde{\Delta}_{n,d}\right) + 2\rho^\natural + C\tilde{\Delta}_{n,d}\sqrt{\log d} \rightarrow 0. \end{aligned}$$

□

E.4 Proof of Theorem 3.2

Proof of Theorem 3.2. By Theorem 3.1, it suffices to prove that

$$\rho^\star := \sup_{x>0} \left| \mathbb{P}(\|\mathbf{Z}\|_\infty \leq x) - \mathbb{P}^\star(\hat{L} \leq x) \right| \xrightarrow{\mathbb{P}} 0.$$

Note that $n^{-1/2}\hat{\Theta}\hat{U}^\top \xi$ is a zero-mean Gaussian random vector with covariance matrix

$$\text{Cov}^\star\left(\frac{1}{\sqrt{n}}\hat{\Theta}\hat{U}^\top \xi\right) = \hat{\sigma}^2\hat{\Theta}\tilde{\Sigma}\hat{\Theta}^\top.$$

By (3.7) and Assumptions 3.1–3.2, we have

$$\begin{aligned} \|\hat{\sigma}^2\hat{\Theta}\tilde{\Sigma}\hat{\Theta}^\top - \sigma^2\Theta\|_{\max} & \leq \hat{\sigma}^2\|\hat{\Theta}\tilde{\Sigma} - \mathbf{I}_d\|_{\max}\|\hat{\Theta}\|_\infty + \hat{\sigma}^2\|\hat{\Theta} - \Theta\|_{\max} + |\hat{\sigma}^2 - \sigma^2|\|\Theta\|_{\max} \\ & = O_{\mathbb{P}}(\Lambda_{\max}\|\Theta\|_\infty + \Delta_{\max} + \Delta_\sigma) = O_{\mathbb{P}}\left(\frac{1}{\log d}\right). \end{aligned}$$

Then it follows from Lemma 2.1 in Chernozhukov et al. [2020] that $\rho^\star \xrightarrow{\mathbb{P}} 0$.

□

E.5 Proof of Lemma E.3

Proof of Lemma E.3. By (3.3) and the triangle inequality,

$$\|n(\tilde{\beta}_\lambda - \beta^*) - \hat{\Theta}\hat{U}^\top \mathcal{E}\|_\infty \leq \|\hat{\Theta}\hat{U}^\top F\varphi^*\|_\infty + n\|(I_d - \hat{\Theta}\tilde{\Sigma})(\hat{\beta}_\lambda - \beta^*)\|_\infty.$$

Since $\mathcal{V}_{n,d}\|\varphi^*\|_2 \lesssim \sqrt{n \log d}$, it follows from Theorem 2.2 that $\|\hat{\beta}_\lambda - \beta^*\|_1 = O_{\mathbb{P}}(|\mathcal{S}_*|\sqrt{(\log d)/n})$. Combining this with Assumption 3.1, we obtain

$$\|(I - \hat{\Theta}\tilde{\Sigma})(\hat{\beta}_\lambda - \beta^*)\|_\infty \leq \|I_d - \hat{\Theta}\tilde{\Sigma}\|_{\max}\|\hat{\beta}_\lambda - \beta^*\|_1 = O_{\mathbb{P}}\left(\Lambda_{\max}|\mathcal{S}_*|\sqrt{\frac{\log d}{n}}\right).$$

By Assumption 3.1, we have $\|\hat{\Theta}\|_\infty \leq \|\hat{\Theta} - \Theta\|_\infty + \|\Theta\|_\infty = O_{\mathbb{P}}(\|\Theta\|_\infty)$. Then, it follows from Lemma C.3 that

$$\|\hat{\Theta}\hat{U}^\top F\varphi^*\|_\infty \leq \|\hat{\Theta}\|_\infty\|\hat{U}^\top F\varphi^*\|_\infty = O_{\mathbb{P}}(\mathcal{V}_{n,d}\|\Theta\|_\infty\|\varphi^*\|_2).$$

By Lemma C.4 and Assumption 2.1, we have $\|U^\top \mathcal{E}\|_\infty^2 = O_{\mathbb{P}}(n \log d)$ and

$$\begin{aligned} \|\hat{\Theta}\hat{U}^\top \mathcal{E} - \Theta U^\top \mathcal{E}\|_\infty &\leq \|\hat{\Theta}\|_\infty\|(\hat{U} - U)^\top \mathcal{E}\|_\infty + \|\hat{\Theta} - \Theta\|_\infty\|U^\top \mathcal{E}\|_\infty \\ &= O_{\mathbb{P}}\left(\|\Theta\|_\infty\sqrt{\log d + \frac{n}{d}} + \Delta_\infty\sqrt{n \log d}\right). \end{aligned}$$

□

F Proof of Results in Section 4

F.1 Proof of Proposition 4.1

Proof. For each $\ell \in [d]$, let $\beta_{\ell,M}^* \in \mathbb{R}$ denote the corresponding population version of the marginal least square estimator $\hat{\beta}_{\ell,M}$, that is,

$$\beta_{\ell,M}^* = \arg \min_{\beta \in \mathbb{R}} \mathbb{E}(Y - u_\ell \beta)^2 = \frac{\Sigma_\ell^\top \beta^*}{\Sigma_{\ell\ell}}.$$

For each $\ell \in [d]$, we have

$$\begin{aligned} \hat{\beta}_{\ell,M} - \beta_{\ell,M}^* &= \frac{1}{\hat{U}_\ell^\top \hat{U}_\ell} \left\{ \hat{U}_\ell^\top (\tilde{Y} - \hat{U} \beta^*) + \left(\hat{U}_\ell^\top \hat{U} - m e_\ell^\top \Sigma \right) \beta^* + \left(\hat{U}_\ell^\top \hat{U}_\ell - m \Sigma_{\ell\ell} \right) \beta_{\ell,M}^* \right\} \\ &=: \delta_{\ell,1} + \delta_{\ell,2} + \delta_{\ell,3}. \end{aligned}$$

Let $\mathbf{U}_\ell = (u_{1\ell}, \dots, u_{m\ell})^\top \in \mathbb{R}^m$ denote the ℓ -th column of the design matrix $\mathbf{U}$. By Assumption 2.1, $u_{1\ell}, \dots, u_{m\ell}$ are i.i.d. sub-Gaussian random variables with $\|u_{1\ell}\|_{\psi_2} \leq c_0$. Hence, for any $x > 0$, by Bernstein's inequality,

$$\mathbb{P} \left(\max_{\ell \in [d]} |\mathbf{U}_\ell^\top \mathbf{U}_\ell - m \Sigma_{\ell\ell}| > x \right) \leq 2d \exp \left\{ -c \min \left(\frac{x^2}{mc_0^4}, \frac{x}{c_0^2} \right) \right\}.$$

Here $c > 0$ is an absolute constant. Combined with Lemma 2.1, we obtain that

$$\max_{\ell \in [d]} |\hat{\mathbf{U}}_\ell^\top \hat{\mathbf{U}}_\ell - m \Sigma_{\ell\ell}| \leq \max_{\ell \in [d]} |\hat{\mathbf{U}}_\ell^\top \hat{\mathbf{U}}_\ell - \mathbf{U}_\ell^\top \mathbf{U}_\ell| + \max_{\ell \in [d]} |\mathbf{U}_\ell^\top \mathbf{U}_\ell - m \Sigma_{\ell\ell}| = O_{\mathbb{P}} \left(\sqrt{m \log d} \right).$$

Since $\log d = o(m)$, we have

$$\mathbb{P}(\mathcal{E}_U) := \mathbb{P} \left(\min_{\ell \in [d]} \hat{\mathbf{U}}_\ell^\top \hat{\mathbf{U}}_\ell \geq \frac{m}{2} \min_{\ell \in [d]} \Sigma_{\ell\ell} \right) \rightarrow 1.$$

Under $\mathcal{E}_U$, it follows from Lemma C.5 that

$$\max_{\ell \in [d]} |\delta_{\ell,1}| \leq \frac{\|\hat{\mathbf{U}}^\top (\tilde{\mathbf{Y}} - \hat{\mathbf{U}} \boldsymbol{\beta}^*)\|_\infty}{\min_{\ell \in [d]} \hat{\mathbf{U}}_\ell^\top \hat{\mathbf{U}}_\ell} \leq \frac{2 \|\hat{\mathbf{U}}^\top (\tilde{\mathbf{Y}} - \hat{\mathbf{U}} \boldsymbol{\beta}^*)\|_\infty}{m \min_{\ell \in [d]} \Sigma_{\ell\ell}} = O_{\mathbb{P}} \left(\frac{\mathcal{V}_{m,d}}{m} \|\boldsymbol{\varphi}^*\|_2 + \sqrt{\frac{\log d}{m}} \right).$$

Similarly, by Lemma C.6, we have

$$\begin{aligned} \max_{\ell \in [d]} |\delta_{\ell,2}| &\leq \frac{2 \|\hat{\mathbf{U}}^\top \hat{\mathbf{U}} - \mathbf{U}^\top \mathbf{U}\|_{\max} \|\boldsymbol{\beta}^*\|_1 + 2 \|\mathbf{U}^\top \mathbf{U} \boldsymbol{\beta}^* - m \Sigma \boldsymbol{\beta}^*\|_\infty}{m \min_{\ell \in [d]} \Sigma_{\ell\ell}} \\ &= O_{\mathbb{P}} \left(\|\boldsymbol{\beta}^*\|_1 \frac{\log d}{m} + \|\boldsymbol{\beta}^*\|_2 \sqrt{\frac{\log d}{m}} \right), \end{aligned}$$

and

$$\max_{\ell \in [d]} |\delta_{\ell,3}| \leq \frac{\max_{\ell \in [d]} |\hat{\mathbf{U}}_\ell^\top \hat{\mathbf{U}}_\ell - m \Sigma_{\ell\ell}| |\boldsymbol{\beta}_{\ell,M}^*|}{m \min_{\ell \in [d]} \Sigma_{\ell\ell}} = O_{\mathbb{P}} \left(\|\boldsymbol{\beta}^*\|_\infty \sqrt{\frac{\log d}{m}} \right).$$

Putting all these pieces together, it follows from (4.3) that $\max_{\ell \in [d]} |\hat{\beta}_{\ell,M} - \beta_{\ell,M}^*| = o(\phi)$, which implies that

$$\mathbb{P}(\mathcal{E}_\phi) := \mathbb{P} \left(\max_{\ell \in \mathcal{S}_*} |\hat{\beta}_{\ell,M} - \beta_{\ell,M}^*| \leq \bar{c}\phi \right) \rightarrow 1, \text{ as } m \rightarrow \infty.$$

Under $\mathcal{E}_\phi$, by (4.3),

$$\min_{\ell \in \mathcal{S}_*} |\hat{\beta}_{\ell,M}| \geq \min_{\ell \in \mathcal{S}_*} |\beta_{\ell,M}^*| - \max_{\ell \in \mathcal{S}_*} |\hat{\beta}_{\ell,M} - \beta_{\ell,M}^*| \geq (1 + \bar{c})\phi - \bar{c}\phi = \phi.$$

Consequently, we obtain

$$\mathbb{P}\left(\mathcal{S}_\star \subset \hat{\mathcal{S}}_\phi\right) \geq \mathbb{P}\left(\min_{\ell \in \mathcal{S}_\star} |\hat{\beta}_{\ell,M}| > \phi\right) \geq \mathbb{P}(\mathcal{E}_\phi) \rightarrow 1.$$

Under $\mathcal{E}_\phi$, elementary calculations show that

$$\begin{aligned} |\hat{\mathcal{S}}_\phi| &= \sum_{\ell=1}^d \mathbb{I}\left\{|\hat{\beta}_{\ell,M}| > \phi\right\} \leq \sum_{\ell=1}^d \mathbb{I}\left\{\max_{\ell \in [d]} |\hat{\beta}_{\ell,M} - \beta_{\ell,M}^\star| + |\beta_{\ell,M}^\star| > \phi\right\} \\ &\leq \sum_{\ell=1}^d \mathbb{I}\left\{|\beta_{\ell,M}^\star| > (1 - \bar{c})\phi\right\} \leq \sum_{\ell=1}^d \frac{|\beta_{\ell,M}^\star|^2}{(1 - \bar{c})^2 \phi^2} \\ &\leq \frac{\|\Sigma \beta^\star\|_2^2}{\lambda_{\min}^2(\Sigma)(1 - \bar{c})^2 \phi^2} = \frac{c_\diamond^2 m^{2\kappa} \|\Sigma \beta^\star\|_2^2}{\lambda_{\min}^2(\Sigma)(1 - \bar{c})^2}. \end{aligned}$$

□

F.2 Proof of Theorem 4.2

Proof. Before proceeding to the proof of the theorem, we first introduce the following Lemma.

Lemma F.1. *Let $\tilde{\mathcal{S}}$ be an estimator for $\mathcal{S}_\star$ such that $\mathbb{P}(|\tilde{\mathcal{S}}| \leq s_n) \rightarrow 1$ and $\mathbb{P}(\mathcal{S}_\star \subset \tilde{\mathcal{S}}) \rightarrow 1$. Moreover, we assume that $\tilde{\mathcal{S}}$ is independent of $(\mathbf{X}, \mathbf{Y})$ and*

$$s_n \left(\frac{\log d}{n} + \frac{1}{d} \right) \rightarrow 0. \quad (\text{F.1})$$

Then, under Assumptions 2.1–2.5, we have that as $n \rightarrow \infty$,

$$\tilde{\rho}_K = \sup_{x>0} \left| \mathbb{P}\left(\frac{\mathcal{E}^\top \mathbf{P}_{\tilde{\mathcal{S}}} \mathcal{E}}{\sigma^2} \leq x\right) - \mathbb{P}(\chi_K^2 \leq x) \right| \rightarrow 0, \text{ where } \mathbf{P}_{\tilde{\mathcal{S}}} = \mathbf{P}_{\hat{\mathbf{F}}} + \mathbf{P}_{\hat{\mathbf{U}}_{\tilde{\mathcal{S}}}} - \mathbf{P}_{\mathbf{X}_{\tilde{\mathcal{S}}}}.$$

By Lemma F.1, we have

$$\rho_K = \sup_{x>0} \left| \mathbb{P}\left(\frac{Q_n^{(2)}}{\sigma^2} \leq x\right) - \mathbb{P}(\chi_K^2 \leq x) \right| \rightarrow 0 \quad (\text{F.2})$$

Hence, it suffices to prove that

$$\rho_\sigma = \sup_{x>0} \left| \mathbb{P}\left(\frac{Q_n^{(2)}}{\hat{\sigma}^2} \leq x\right) - \mathbb{P}\left(\frac{Q_n^{(2)}}{\sigma^2} \leq x\right) \right| \rightarrow 0.$$

After conducting some elementary calculations, we obtain

$$\rho_\sigma \leq \mathbb{P}\left(\frac{Q_n^{(2)}}{\sigma^2} \left| \frac{\sigma^2}{\hat{\sigma}^2} - 1 \right| > \sqrt{\Delta_\sigma}\right) + \sup_{x>0} \mathbb{P}\left(x < \chi_K^2 \leq x + \sqrt{\Delta_\sigma}\right) + 2\rho_K =: \delta_1 + \delta_2 + 2\rho_K.$$

As $\Delta_\sigma \rightarrow 0$, we have $\delta_2 \rightarrow 0$. Moreover, by Assumption 3.3, we have $|\sigma^2/\hat{\sigma}^2 - 1| \leq |\sigma/\hat{\sigma} - 1| = O_{\mathbb{P}}(\Delta_\sigma)$. Combined with (F.2), we obtain $\delta_1 \rightarrow 0$. Thus, we claim our conclusion of Theorem 4.2. □

F.3 Proof of Lemma F.1

Proof. As $\tilde{\mathcal{S}}$ is independent of $(\mathbf{X}, \mathbf{Y})$, it follows that

$$\tilde{\rho}_K \leq \mathbb{P}(|\tilde{\mathcal{S}}| > s_n) + \sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \sup_{x > 0} \left| \mathbb{P} \left(\frac{\mathcal{E}^\top \mathbf{P}_{\mathcal{S}} \mathcal{E}}{\sigma^2} \leq x \right) - \mathbb{P}(\chi_K^2 \leq x) \right|.$$

Since $\mathbb{P}(|\tilde{\mathcal{S}}| > s_n) \rightarrow 0$ as $n \rightarrow \infty$, it suffices to prove that

$$\rho^\diamond := \sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \sup_{x > 0} \left| \mathbb{P} \left(\frac{\mathcal{E}^\top \mathbf{P}_{\mathcal{S}} \mathcal{E}}{\sigma^2} \leq x \right) - \mathbb{P}(\chi_K^2 \leq x) \right| \rightarrow 0.$$

For any set $\mathcal{S} \subset [d]$ such that $|\mathcal{S}| \leq s_n$ and $\mathcal{S}_\star \subset \mathcal{S}$, we define

$$\mathcal{A}_{\mathcal{S}} = \left\{ \lambda_{\min}(\hat{\mathbf{U}}_{\mathcal{S}}^\top \hat{\mathbf{U}}_{\mathcal{S}}) \geq \frac{n\lambda_0}{2} \right\}.$$

Then by Theorem 1.1 given in Bentkus [2005], it follows that

$$\sup_{x > 0} \left| \mathbb{P} \left(\frac{\mathcal{E}^\top \mathbf{P}_{\mathcal{S}} \mathcal{E}}{\sigma^2} \leq x \right) - \mathbb{P}(\chi_K^2 \leq x) \right| \leq \mathbb{P}(\mathcal{A}_{\mathcal{S}}^c) + C \mathbb{E}|\varepsilon_t|^3 K^{5/4} \mathbb{E} \left(\max_{t \in [n]} \sqrt{\mathbf{P}_{\mathcal{S}, tt}} \right) \mathbb{I}\{\mathcal{A}_{\mathcal{S}}\}, \quad (\text{F.3})$$

where $C > 0$ is an absolute constant. Now we proceed to prove (F.3). Recall that $\mathbf{P}_{\mathcal{S}} = \mathbf{P}_{\hat{\mathbf{F}}} + \mathbf{P}_{\hat{\mathbf{U}}_{\tilde{\mathcal{S}}}} - \mathbf{P}_{\mathbf{X}_{\mathcal{S}}}$ is a projection matrix onto the linear space generated by the columns of $\tilde{\mathbf{F}} = (\mathbf{I}_n - \mathbf{P}_{\mathbf{X}_{\mathcal{S}}})\tilde{\mathbf{F}} \in \mathbb{R}^{n \times K}$. Hence, under $\mathcal{A}_{\mathcal{S}}$, we can write

$$\mathbf{P}_{\mathcal{S}} = \tilde{\mathbf{F}}(\tilde{\mathbf{F}}^\top \tilde{\mathbf{F}})^{-1} \tilde{\mathbf{F}}^\top = \tilde{\mathbf{F}}(\tilde{\mathbf{F}}^\top \tilde{\mathbf{F}})^{-1/2} (\tilde{\mathbf{F}}^\top \tilde{\mathbf{F}})^{-1/2} \tilde{\mathbf{F}}^\top =: \mathbf{W} \mathbf{W}^\top,$$

where $\mathbf{W} = (\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_n)^\top \in \mathbb{R}^{n \times K}$. Therefore,

$$\mathcal{E}^\top \mathbf{P}_{\mathcal{S}} \mathcal{E} = \mathcal{E}^\top \mathbf{W} \mathbf{W}^\top \mathcal{E} = \left\| \sum_{t=1}^n \mathbf{w}_t \varepsilon_t \right\|_2^2,$$

which further implies that for any $x > 0$,

$$\mathbb{P}(\mathcal{E}^\top \mathbf{P}_{\mathcal{S}} \mathcal{E} \leq x) = \mathbb{P} \left(\sum_{t=1}^n \mathbf{w}_t \varepsilon_t \in \mathcal{B}_K(\sqrt{x}) \right),$$

where $\mathcal{B}_K(r)$ denotes the K -dimensional ball centered at the origin with radius $r > 0$. Observe that

$$\text{Cov} \left(\sum_{t=1}^n \mathbf{w}_t \varepsilon_t \right) = \sum_{t=1}^n \mathbf{w}_t \mathbf{w}_t^\top = \mathbf{W}^\top \mathbf{W} = (\tilde{\mathbf{F}}^\top \tilde{\mathbf{F}})^{-1/2} \tilde{\mathbf{F}}^\top \tilde{\mathbf{F}} (\tilde{\mathbf{F}}^\top \tilde{\mathbf{F}})^{-1/2} = \mathbf{I}_K.$$

To apply Theorem 1.1 in [Bentkus \[2005\]](#), let $\mathcal{Z} \in \mathbb{R}^K$ be a zero-mean Gaussian random vector with covariance matrix $\text{Cov}(\mathcal{Z}) = \mathbf{I}_K$, then $\|\mathcal{Z}\|^2 \sim \chi_K^2$ and

$$\begin{aligned}
& \sup_{x>0} \left| \mathbb{P}(\mathcal{E}^\top \mathbf{P}_S \mathcal{E} \leq x | \mathbf{X}) - \mathbb{P}(\chi_K^2 \leq x) \right| \\
&= \sup_{r>0} \left| \mathbb{P} \left(\sum_{t=1}^n \mathbf{w}_t \varepsilon_t \in \mathcal{B}_K(r) | \mathbf{X} \right) - \mathbb{P}(\mathcal{Z} \in \mathcal{B}_K(r)) \right| \\
&\leq cK^{1/4} \sum_{t=1}^n \mathbb{E} \|\mathbf{w}_t \varepsilon_t\|_2^3 = cK^{1/4} \sum_{t=1}^n \|\mathbf{w}_t\|_2^3 \mathbb{E} |\varepsilon_t|^3 \\
&\leq cK^{1/4} \mathbb{E} |\varepsilon|^3 \max_{t \in [n]} \|\mathbf{w}_t\|_2 \sum_{t=1}^n \|\mathbf{w}_t\|_2^2.
\end{aligned}$$

Observe that the diagonal elements of the matrix $\mathbf{P}_S$ are $\|\mathbf{w}_1\|_2^2, \dots, \|\mathbf{w}_n\|_2^2$. Hence

$$\sum_{t=1}^n \|\mathbf{w}_t\|_2^2 = \text{tr}(\mathbf{P}_S) = K \quad \text{and} \quad \max_{t \in [n]} \|\mathbf{w}_t\|_2 = \max_{t \in [n]} \sqrt{\mathbf{P}_{S,tt}}.$$

Consequently, we obtain

$$\sup_{x>0} \left| \mathbb{P}(\mathcal{E}^\top \mathbf{P}_S \mathcal{E} \leq x | \mathbf{X}) - \mathbb{P}(\chi_K^2 \leq x) \right| \leq cK^{5/4} \mathbb{E} |\varepsilon|^3 \max_{t \in [n]} \sqrt{\mathbf{P}_{S,tt}},$$

and

$$\sup_{x>0} \left| \mathbb{P}(\mathcal{E}^\top \mathbf{P}_S \mathcal{E} \leq x) - \mathbb{P}(\chi_K^2 \leq x) \right| \leq cK^{5/4} \mathbb{E} |\varepsilon|^3 \mathbb{E} \left(\max_{t \in [n]} \sqrt{\mathbf{P}_{S,tt}} \mathbb{I}\{\mathcal{A}_S\} \right) + \mathbb{P}(\mathcal{A}_S^c). \quad (\text{F.4})$$

Recall that both $\mathbf{P}_{\hat{\mathbf{F}}} + \mathbf{P}_{\hat{\mathbf{U}}_S} - \mathbf{P}_{\mathbf{X}_S}$ and $\mathbf{P}_{\mathbf{X}_S}$ are orthogonal projection matrices. Hence

$$\max_{t \in [n]} \mathbf{P}_{S,tt} = \max_{t \in [n]} \left(\mathbf{P}_{\hat{\mathbf{F}}} + \mathbf{P}_{\hat{\mathbf{U}}_S} - \mathbf{P}_{\mathbf{X}_S} \right)_{tt} \leq \max_{t \in [n]} \mathbf{P}_{\hat{\mathbf{F}},tt} + \max_{t \in [n]} \mathbf{P}_{\hat{\mathbf{U}}_S,tt}.$$

By Lemma 2.1, it follows that

$$\max_{t \in [n]} \mathbf{P}_{\hat{\mathbf{F}},tt} = \max_{t \in [n]} \frac{1}{n} \|\hat{\mathbf{f}}_t\|_2^2 \lesssim \max_{t \in [n]} \frac{1}{n} \|\hat{\mathbf{f}}_t - \mathbf{H} \mathbf{f}_t\|_2^2 + \max_{t \in [n]} \frac{1}{n} \|\mathbf{H}\|_2^2 \|\mathbf{f}_t\|_2^2 = O_{\mathbb{P}} \left(\frac{\log n}{n} \right). \quad (\text{F.5})$$

Now we bound $\max_{t \in [n]} \mathbf{P}_{\hat{\mathbf{U}}_S,tt}$. Under event $\mathcal{A}_S$, we have

$$\max_{t \in [n]} \mathbf{P}_{\hat{\mathbf{U}}_S,tt} \leq \frac{1}{\lambda_{\min}(\hat{\mathbf{U}}_S^\top \hat{\mathbf{U}}_S)} \max_{t \in [n]} \sum_{\ell \in \mathcal{S}} |\hat{u}_{t\ell}|^2 \leq \frac{2}{n\lambda_0} \max_{t \in [n]} \sum_{\ell \in \mathcal{S}} |\hat{u}_{t\ell}|^2.$$

Recall that $\hat{\mathbf{U}} = (\mathbf{I}_n - \hat{\mathbf{P}})\mathbf{X} = (\mathbf{I}_n - \hat{\mathbf{P}})\mathbf{F}\mathbf{B}^\top + (\mathbf{I}_n - \hat{\mathbf{P}})\mathbf{U}$ and $\hat{\mathbf{F}}^\top \hat{\mathbf{U}} = \mathbf{O}$. Hence, under the event $\mathcal{E}_\lambda$, we have

$$\hat{\mathbf{U}} = (\mathbf{I}_n - \hat{\mathbf{P}})(\mathbf{F}\mathbf{H}^\top - \hat{\mathbf{F}})\mathbf{H}^{-\top}\mathbf{B}^\top - \hat{\mathbf{P}}(\mathbf{U} - \hat{\mathbf{U}}) + \mathbf{U}.$$

Consequently, it follows that

$$\begin{aligned} \max_{t \in [n]} \sum_{\ell \in \mathcal{S}} |\hat{u}_{t\ell}|^2 &\lesssim \max_{t \in [n]} \sum_{\ell \in \mathcal{S}} |\mathbf{e}_t^\top (\mathbf{I}_n - \hat{\mathbf{P}})(\mathbf{F}\mathbf{H}^\top - \hat{\mathbf{F}})\mathbf{H}^{-\top}\mathbf{b}_\ell|^2 \\ &\quad + \max_{t \in [n]} \sum_{\ell \in \mathcal{S}} |\mathbf{e}_t^\top \hat{\mathbf{P}}(\mathbf{U}_\ell - \hat{\mathbf{U}}_\ell)|^2 + \max_{t \in [n]} \sum_{\ell \in \mathcal{S}} |u_{t\ell}|^2 \\ &=: \Delta_1^{\natural} + \Delta_2^{\natural} + \Delta_3^{\natural}. \end{aligned}$$

By Lemma 2.1, we have

$$\begin{aligned} \Delta_1^{\natural} &\leq \|\mathbf{F}\mathbf{H}^\top - \hat{\mathbf{F}}\|_{\mathbb{F}}^2 \|\mathbf{H}^{-\top}\|_2^2 \sum_{\ell \in \mathcal{S}} \|\mathbf{b}_\ell\|_2^2 = O_{\mathbb{P}} \left(\frac{|\mathcal{S}|}{n} + \frac{|\mathcal{S}|n}{d} \right), \\ \Delta_2^{\natural} &\leq \frac{|\mathcal{S}|}{n} \max_{t \in [n]} \|\hat{\mathbf{f}}_t\|_2^2 \max_{\ell \in [d]} \sum_{s=1}^n |\hat{u}_{s\ell} - u_{s\ell}|^2 = O_{\mathbb{P}} \left\{ |\mathcal{S}| \log n \left(\frac{\log d}{n} + \frac{1}{d} \right) \right\}, \end{aligned}$$

and $\Delta_3^{\natural} = O_{\mathbb{P}}(|\mathcal{S}| + \log n)$. Therefore, under event $\mathcal{A}_{\mathcal{S}}$, we have

$$\max_{t \in [n]} \mathbf{P}_{\hat{\mathbf{U}}_{\mathcal{S}}, tt} \leq \frac{2}{n\lambda_0} \max_{t \in [n]} \sum_{\ell \in \mathcal{S}} |\hat{u}_{t\ell}|^2 = O_{\mathbb{P}} \left(\frac{s_n}{n \wedge d} + \frac{\log n}{n} \right).$$

Combined with (F.5), we obtain

$$\max_{t \in [n]} \mathbf{P}_{\mathcal{S}, tt} = O_{\mathbb{P}}(\varpi_n), \text{ where } \varpi_n = \frac{s_n}{n \wedge d} + \frac{\log n}{n}.$$

Note that $\varpi_n \rightarrow 0$ by (F.1). As $\mathbf{P}_{\mathcal{S}}$ is an orthogonal matrix, we have $\max_{t \in [n]} \mathbf{P}_{\mathcal{S}, tt} \leq 1$ and consequently

$$\sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \mathbb{E} \left(\max_{t \in [n]} \mathbf{P}_{\mathcal{S}, tt} \right) \mathbb{I}\{\mathcal{A}_{\mathcal{S}}\} \leq \sqrt{\varpi_n} + \sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \mathbb{P} \left(\max_{t \in [n]} \mathbf{P}_{\mathcal{S}, tt} > \sqrt{\varpi_n} \right) \rightarrow 0.$$

In view of (F.3), we obtain $\rho^\diamond \rightarrow 0$ once we prove that

$$\sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \mathbb{P}(\mathcal{A}_{\mathcal{S}}^c) = \sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \mathbb{P} \left\{ \lambda_{\min}(\hat{\mathbf{U}}_{\mathcal{S}}^\top \hat{\mathbf{U}}_{\mathcal{S}}) \leq \frac{n\lambda_0}{2} \right\} \rightarrow 0. \quad (\text{F.6})$$

Next, we will prove (F.6). Recall that $\tilde{\Sigma} = n^{-1} \hat{\mathbf{U}}^\top \hat{\mathbf{U}}$ and $\hat{\Sigma} = n^{-1} \mathbf{U}^\top \mathbf{U}$. By Weyl's theorem on eigenvalues, it follows that

$$\sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \mathbb{P} \left\{ \lambda_{\min}(\hat{\mathbf{U}}_{\mathcal{S}}^\top \hat{\mathbf{U}}_{\mathcal{S}}) \leq \frac{n\lambda_0}{2} \right\} \leq \sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \mathbb{P} \left(\|\tilde{\Sigma}_{\mathcal{S}\mathcal{S}} - \hat{\Sigma}_{\mathcal{S}\mathcal{S}}\| > \frac{\lambda_0}{4} \right)$$

$$+ \sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \mathbb{P} \left(\|\tilde{\Sigma}_{\mathcal{S}\mathcal{S}} - \hat{\Sigma}_{\mathcal{S}\mathcal{S}}\| > \frac{\lambda_0}{4} \right) =: \pi_1 + \pi_2.$$

We first bound π_1 . By Lemma C.6 and (F.1),

$$\sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} \|\tilde{\Sigma}_{\mathcal{S}\mathcal{S}} - \hat{\Sigma}_{\mathcal{S}\mathcal{S}}\| \leq \sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} |\mathcal{S}| \|\tilde{\Sigma} - \hat{\Sigma}\|_{\max} = O_{\mathbb{P}} \left(\frac{s_n \log d}{n} + \frac{s_n}{d} \right) = o_{\mathbb{P}}(1). \quad (\text{F.7})$$

Hence $\pi_1 \rightarrow 0$ as $n \rightarrow \infty$. Now we bound π_2 . For any subset $\mathcal{S} \subset [d]$ such that $|\mathcal{S}| \leq s_n$, by Lemma 5.4 in Vershynin [2012] and Assumption 2.1,

$$\begin{aligned} \mathbb{P} \left(\|\hat{\Sigma}_{\mathcal{S}\mathcal{S}} - \Sigma_{\mathcal{S}\mathcal{S}}\| > \frac{\lambda_0}{4} \right) &\leq 9^{|\mathcal{S}|} \sup_{\mathbf{v} \in \mathbb{S}^{d-1}: \mathbf{v}_{\mathcal{S}^c} = 0} \mathbb{P} \left(\left| \frac{1}{n} \sum_{t=1}^n (\mathbf{u}_t^\top \mathbf{v})^2 - \mathbf{v}^\top \Sigma \mathbf{v} \right| > \frac{\lambda_0}{8} \right) \\ &\leq 2 \exp \{ |\mathcal{S}| \log 9 - C \min(n \lambda_0^2, n \lambda_0) \}. \end{aligned} \quad (\text{F.8})$$

Consequently, by (F.1), we obtain

$$\begin{aligned} \pi_2 &\leq \sup_{\mathcal{S} \subset [d]: |\mathcal{S}| \leq s_n} 2 \exp (|\mathcal{S}| \log 9 - C n \min\{\lambda_0^2, \lambda_0\}) \\ &\leq 2 \exp (s_n \log 9 - C n \min\{\lambda_0^2, \lambda_0\}) \rightarrow 0. \end{aligned}$$

□

Lemma F.2. Assume that

$$\|\varphi^*\|_2 \left(\sqrt{n/d} + 1/\sqrt{n} \right) \rightarrow 0. \quad (\text{F.9})$$

Define

$$\mathcal{H}(\alpha, \theta) = \left\{ \varphi \in \mathbb{R}^K : \frac{n \|\varphi\|_2^2}{1 + K s_n \Upsilon^2 / \lambda_{\min}(\Sigma)} \geq \sigma^2 (2 + \delta) (\chi_{K, 1-\alpha}^2 + \chi_{K, 1-\theta}^2) \right\}$$

for some $\delta > 0$. Then, under the conditions of Lemma F.1, we have

$$\inf_{\varphi^* \in \mathcal{H}(\alpha, \theta)} \mathbb{P} \left(\frac{\|P_{\hat{\mathcal{S}}} Y\|_2^2}{\sigma^2} > \chi_{K, 1-\alpha}^2 \right) \geq 1 - \theta. \quad (\text{F.10})$$

Proof of Lemma F.2. Recall that $Y = F\varphi^* + X\beta^* + \mathcal{E} = Y = F\varphi^* + X_{\hat{\mathcal{S}}}\beta_{\hat{\mathcal{S}}}^* + \mathcal{E}$. Hence $P_{\hat{\mathcal{S}}}Y = P_{\hat{\mathcal{S}}}\mathcal{E} + P_{\hat{\mathcal{S}}}F\varphi^*$. Without loss of generality, we assume that $\sigma = 1$. By the Cauchy-Schwarz inequality,

$$\|P_{\hat{\mathcal{S}}}Y\|_2^2 \geq \frac{1}{2} \|P_{\hat{\mathcal{S}}}F\varphi^*\|_2^2 - \|P_{\hat{\mathcal{S}}}\mathcal{E}\|_2^2.$$

Therefore, for any $\varphi^\star \in \mathcal{H}(\alpha, \theta)$, we have

$$\mathbb{P}(\|P_{\tilde{\mathcal{S}}} \mathbf{Y}\|_2^2 > \chi_{K,1-\alpha}^2) \geq \mathbb{P}\left(\frac{1}{2}\|P_{\tilde{\mathcal{S}}} \mathbf{F} \varphi^\star\|_2^2 > \chi_{K,1-\alpha}^2 + \|P_{\tilde{\mathcal{S}}} \mathcal{E}\|_2^2\right).$$

We first bound $\|P_{\tilde{\mathcal{S}}} \mathbf{F} \varphi^\star\|_2^2$. By Lemma 2.1 and the Cauchy-Schwarz inequality, for $\epsilon > 0$,

$$\|P_{\tilde{\mathcal{S}}} \mathbf{F} \varphi^\star\|_2^2 \geq (1 - \epsilon) \|P_{\tilde{\mathcal{S}}} \hat{\mathbf{F}} \mathbf{H} \varphi^\star\|_2^2 - \left(\frac{1}{\epsilon} - 1\right) \|P_{\tilde{\mathcal{S}}} (\mathbf{F} - \hat{\mathbf{F}} \mathbf{H}) \varphi^\star\|_2^2.$$

As $P_{\tilde{\mathcal{S}}}$ is a projection matrix, it follows from Lemma 2.1 and (F.9) that

$$\|P_{\tilde{\mathcal{S}}} (\mathbf{F} - \hat{\mathbf{F}} \mathbf{H}) \varphi^\star\|_2 \leq \|\mathbf{F} - \hat{\mathbf{F}} \mathbf{H}\|_{\mathbb{F}} \|\varphi^\star\|_2 = O_{\mathbb{P}}\left(\|\varphi^\star\|_2 \sqrt{n/d} + \|\varphi^\star\|_2 \sqrt{1/n}\right) = o_{\mathbb{P}}(1).$$

Recall that $P_{\tilde{\mathcal{S}}} = P_{\hat{\mathbf{F}}} + P_{\hat{\mathbf{U}}_{\tilde{\mathcal{S}}}} - P_{\mathbf{X}_{\tilde{\mathcal{S}}}}$ and $\tilde{\Sigma} = n^{-1} \hat{\mathbf{U}}^\top \hat{\mathbf{U}}$. Hence

$$\begin{aligned} \|P_{\tilde{\mathcal{S}}} \hat{\mathbf{F}} \mathbf{H} \varphi^\star\|_2^2 &= n(\mathbf{H} \varphi^\star)^\top \left\{ \mathbf{I}_K - \hat{\mathbf{B}}_{\tilde{\mathcal{S}}}^\top \left(\hat{\mathbf{B}}_{\tilde{\mathcal{S}}} \hat{\mathbf{B}}_{\tilde{\mathcal{S}}}^\top + \tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}} \right)^{-1} \hat{\mathbf{B}}_{\tilde{\mathcal{S}}} \right\} (\mathbf{H} \varphi^\star) \\ &= n(\mathbf{H} \varphi^\star)^\top \left(\mathbf{I}_K + \hat{\mathbf{B}}_{\tilde{\mathcal{S}}}^\top \tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}^{-1} \hat{\mathbf{B}}_{\tilde{\mathcal{S}}} \right)^{-1} (\mathbf{H} \varphi^\star) \\ &= n \varphi^{\star\top} \left\{ (\mathbf{H}^\top \mathbf{H})^{-1} + \mathbf{H}^{-1} \hat{\mathbf{B}}_{\tilde{\mathcal{S}}}^\top \tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}^{-1} \hat{\mathbf{B}}_{\tilde{\mathcal{S}}} \mathbf{H}^{-\top} \right\}^{-1} \varphi^\star =: n \varphi^{\star\top} \hat{\mathbb{W}}^{-1} \varphi^\star. \end{aligned}$$

Let $\mathbb{W} = \mathbf{I}_K + \mathbf{B}_{\tilde{\mathcal{S}}}^\top \Sigma_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}^{-1} \mathbf{B}_{\tilde{\mathcal{S}}}$. We first upper bound $\|\hat{\mathbb{W}} - \mathbb{W}\|_2$. For any $\mathbf{h} \in \mathbb{R}^{|\tilde{\mathcal{S}}|}$ with $\|\mathbf{h}\|_2 = 1$, we decompose $\mathbf{h}^\top (\hat{\mathbb{W}} - \mathbb{W}) \mathbf{h} = \Psi_1 + \Psi_2 + \Psi_3 + \Psi_4$, where

$$\begin{aligned} \Psi_1 &= \mathbf{h}^\top (\mathbf{H}^\top \mathbf{H})^{-1} (\mathbf{I}_K - \mathbf{H}^\top \mathbf{H}) \mathbf{h}, \\ \Psi_2 &= \mathbf{h}^\top \mathbf{H}^{-1} \left(\hat{\mathbf{B}}_{\tilde{\mathcal{S}}} - \mathbf{B}_{\tilde{\mathcal{S}}} \mathbf{H}^\top \right)^\top \tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}^{-1} \hat{\mathbf{B}}_{\tilde{\mathcal{S}}} \mathbf{H}^{-\top} \mathbf{h}, \\ \Psi_3 &= \mathbf{h}^\top \mathbf{B}_{\tilde{\mathcal{S}}}^\top \tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}^{-1} \left(\Sigma_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}} - \tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}} \right) \Sigma_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}^{-1} \hat{\mathbf{B}}_{\tilde{\mathcal{S}}} \mathbf{H}^{-\top} \mathbf{h}, \\ \Psi_4 &= \mathbf{h}^\top \mathbf{B}_{\tilde{\mathcal{S}}}^\top \Sigma_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}^{-1} \left(\hat{\mathbf{B}}_{\tilde{\mathcal{S}}} - \mathbf{B}_{\tilde{\mathcal{S}}} \mathbf{H}^\top \right) \mathbf{H}^{-\top} \mathbf{h}. \end{aligned}$$

By Lemma 2.1, we have $\|\mathbf{H}^\top \mathbf{H} - \mathbf{I}_K\|_{\mathbb{F}} = o_{\mathbb{P}}(1)$ and $\mathbb{P}\{\lambda_{\min}(\mathbf{H}^\top \mathbf{H}) \geq 1/2\} \rightarrow 1$. Therefore

$$\Psi_1 \leq \frac{\|\mathbf{H}^\top \mathbf{H} - \mathbf{I}_K\|_{\mathbb{F}}}{\lambda_{\min}(\mathbf{H}^\top \mathbf{H})} = o_{\mathbb{P}}(1)$$

and $\mathbb{P}(\|\mathbf{H}^{-\top} \mathbf{h}\|_2^2 \leq 2) \rightarrow 1$. Since $\mathbb{P}(|\tilde{\mathcal{S}}| \leq s_n) \rightarrow 1$, it follows that

$$\left\| \left(\hat{\mathbf{B}}_{\tilde{\mathcal{S}}} - \mathbf{B}_{\tilde{\mathcal{S}}} \mathbf{H}^\top \right) \mathbf{H}^{-\top} \mathbf{h} \right\|^2 \leq |\tilde{\mathcal{S}}| \max_{j \in [d]} \|\hat{\mathbf{b}}_j - \mathbf{H} \mathbf{b}_j\|^2 \|\mathbf{H}^{-\top} \mathbf{h}\|^2 = O_{\mathbb{P}} \left\{ s_n \left(\frac{\log d}{n} + \frac{1}{d} \right) \right\}.$$

Combining this with (F.1) and the fact that $\|\mathbf{B}_{\tilde{\mathcal{S}}}\mathbf{h}\|^2 \leq K|\tilde{\mathcal{S}}|\mathcal{K}^2$, we obtain

$$\Psi_2 \leq \frac{\|\hat{\mathbf{B}}_{\tilde{\mathcal{S}}}\mathbf{H}^{-\top}\mathbf{h}\|}{\lambda_{\min}(\tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}})} \left\| \left(\hat{\mathbf{B}}_{\tilde{\mathcal{S}}} - \mathbf{B}_{\tilde{\mathcal{S}}}\mathbf{H}^{\top} \right) \mathbf{H}^{-\top}\mathbf{h} \right\| = o_{\mathbb{P}}(s_n^{1/2}).$$

Similarly, by (F.7) and (F.8), we have $\|\tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}} - \Sigma_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}\|_2 = o_{\mathbb{P}}(1)$,

$$\Psi_3 \leq \frac{\|\mathbf{B}_{\tilde{\mathcal{S}}}\mathbf{h}\| \|\tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}} - \Sigma_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}}\|_2 \|\hat{\mathbf{B}}_{\tilde{\mathcal{S}}}\mathbf{H}^{-\top}\mathbf{h}\|}{\lambda_{\min}(\tilde{\Sigma}_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}})\lambda_{\min}(\Sigma_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}})} = o_{\mathbb{P}}(s_n^{1/2}).$$

and

$$\Psi_4 \leq \frac{\|\mathbf{B}_{\tilde{\mathcal{S}}}\mathbf{h}\|}{\lambda_{\min}(\Sigma_{\tilde{\mathcal{S}}\tilde{\mathcal{S}}})} \left\| \left(\hat{\mathbf{B}}_{\tilde{\mathcal{S}}} - \mathbf{B}_{\tilde{\mathcal{S}}}\mathbf{H}^{\top} \right) \mathbf{H}^{-\top}\mathbf{h} \right\| = o_{\mathbb{P}}(s_n^{1/2}).$$

Consequently, we obtain $\|\hat{\mathbb{W}} - \mathbb{W}\|_2 = o_{\mathbb{P}}(\sqrt{s_n})$. Combining this with the fact that

$$\mathbb{P} \left\{ \lambda_{\max}(\mathbb{W}) \leq 1 + \frac{Ks_n\|\mathbf{B}\|_{\max}^2}{\lambda_{\min}(\Sigma)} \right\} \rightarrow 1,$$

we obtain

$$\mathbb{P} \left\{ \|\hat{\mathbb{W}} - \mathbb{W}\|_2 \leq \epsilon^{\diamond} \left(1 + \frac{Ks_n\|\mathbf{B}\|_{\max}^2}{\lambda_{\min}(\Sigma)} \right) \right\} \rightarrow 1.$$

Note that

$$\|P_{\tilde{\mathcal{S}}}\hat{\mathbf{F}}\mathbf{H}\varphi^{\star}\|^2 \geq \frac{n\|\varphi^{\star}\|^2}{\lambda_{\max}(\hat{\mathbb{W}})} \geq \frac{n\|\varphi^{\star}\|^2}{\lambda_{\max}(\mathbb{W}) + \|\hat{\mathbb{W}} - \mathbb{W}\|_2}.$$

Hence

$$\mathbb{P} \left\{ \|P_{\tilde{\mathcal{S}}}\hat{\mathbf{F}}\mathbf{H}\varphi^{\star}\|^2 \geq \frac{n\|\varphi^{\star}\|^2}{(1 + \epsilon^{\diamond})(1 + Ks_n\|\mathbf{B}\|_{\max}^2/\lambda_{\min}(\Sigma))} \right\} \rightarrow 1.$$

Putting all these pieces together, we obtain (F.10). $\square$

F.4 Application to multi-modal sparse regression model

Dataset with multiple types are now frequently collected for some mutual experimental subjects. This data structure is called as multimodal data and is becoming more and more popular in many fields. Factor analysis is commonly used in integrative analysis of multimodal data, and is particularly useful to overcome the curse of high dimensionality and high correlations. Recently Li and Li [2021] study sparse linear multi-modal regression model with factor structures. However, the hypothesis testing

problem on whether sparse regression for a given modal (a certain group of modal) is adequate hasn't been investigated.

Thus, in this subsection, we extend our results given in the last subsection to the multi-modal sparse regression model. Our observations are $(\mathbf{Y}, \mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_M)$ in which we assume covariate $\mathbf{X}_i$ is generated from i -th group (modal). Moreover, for every $i \in [M]$, $\mathbf{X}_i$ is assumed to have its own factor structure $\mathbf{X}_i = \mathbf{F}_i \mathbf{B}_i^\top + \mathbf{U}_i$. We next consider the hypothesis test as follows:

$$H_0 : \mathbf{Y} = \sum_{i=1}^L \mathbf{X}_i \boldsymbol{\beta}_i^* + \sum_{i=L+1}^M \mathbf{X}_i \boldsymbol{\beta}_i^* + \mathcal{E} \quad \text{versus} \quad H_1 : \mathbf{Y} = \sum_{i=1}^L (\mathbf{F}_i \boldsymbol{\gamma}_i^* + \mathbf{U}_i \boldsymbol{\beta}_i^*) + \sum_{i=L+1}^M \mathbf{X}_i \boldsymbol{\beta}_i^* + \mathcal{E}.$$

Here we aim at simultaneously testing whether the sparse regression is adequate for the first L modals. L is assumed to be any fixed number in $[M]$. The hypothesis testing problem in the previous section is a special case to this with $L = M = 1$.

In order to proceed the hypothesis testing procedure, similar to section 4, we separate our dataset into two parts with size m and $n - m$ respectively. We let $\mathbf{X}_i^{(j)}$ and $\mathbf{Y}^{(j)}$, $j \in \{1, 2\}$ denote the i -th modal and our response respectively in the j -th part of our splitted data. Next, we decompose every $\mathbf{X}_i^{(1)}$ into $(\hat{\mathbf{F}}_i^{(1)}, \hat{\mathbf{U}}_i^{(1)})$ and then use sure screening to select $(\hat{S}_1, \hat{S}_2, \hat{S}_3, \dots, \hat{S}_M)$ by using $(\mathbf{Y}^{(1)}, \hat{\mathbf{U}}_1^{(1)}, \hat{\mathbf{U}}_2^{(1)}, \dots, \hat{\mathbf{U}}_L^{(1)}, \hat{\mathbf{U}}_{L+1}^{(1)}, \dots, \hat{\mathbf{U}}_M^{(1)})$. Here we also make an assumption that

Assumption F.1. Here we assume

$$\mathbb{P} \left(\text{for all } i \in [M], S_i^* \subset \hat{S}_i \text{ and } \sum_{i=1}^M |S_i| \leq s_n^M \right) \rightarrow 1. \quad (\text{F.11})$$

Next, we also impose an assumption on the covariance of our factors and idiosyncratic components

Assumption F.2. We assume

$$\text{Cov}(\mathbf{f}_{1,t}, \mathbf{f}_{2,t}, \dots, \mathbf{f}_{M,t}) = I_K, \quad \text{Cov}(\mathbf{u}_{1,t}, \mathbf{u}_{2,t}, \dots, \mathbf{u}_{M,t}) = \Sigma_u,$$

with $0 < \lambda_{\min}(\Sigma_u), \lambda_{\max}(\Sigma_u) \leq C$ with C being a absolute constant. Moreover, $(\mathbf{f}_{1,t}, \dots, \mathbf{f}_{M,t}), t \in [n]$ is assumed to be uncorrelated with $(\mathbf{u}_{1,t}, \dots, \mathbf{u}_{M,t})$, where $\mathbf{f}_{i,t}$ and $\mathbf{u}_{i,t}$ are the t -th row of $\mathbf{F}_i$ and $\mathbf{U}_i$ respectively.

Next, we use the second part of the data to construct our test statistic. For simplicity we denote

$$\mathbf{Z}_1 = (\mathbf{X}_{\hat{S}_1}^{(2)}, \dots, \mathbf{X}_{\hat{S}_L}^{(2)}, \hat{\mathbf{F}}_{L+1}^{(2)}, \hat{\mathbf{U}}_{\hat{S}_{L+1}}^{(2)}, \dots, \hat{\mathbf{F}}_M^{(2)}, \hat{\mathbf{U}}_{\hat{S}_M}^{(2)})$$

and

$$\mathbf{Z}_2 = (\hat{\mathbf{F}}_1^{(2)}, \hat{\mathbf{U}}_{\hat{S}_1}^{(2)}, \dots, \hat{\mathbf{F}}_L^{(2)}, \hat{\mathbf{U}}_{\hat{S}_L}^{(2)}, \hat{\mathbf{F}}_{L+1}^{(2)}, \hat{\mathbf{U}}_{\hat{S}_{L+1}}^{(2)}, \dots, \hat{\mathbf{F}}_M^{(2)}, \hat{\mathbf{U}}_{\hat{S}_M}^{(2)}).$$

The test statistic is given by

$$Q_{n,M}^{(2)} = \|(\mathbf{I}_{n-m} - \mathbf{P}_{\mathbf{Z}_1}) \mathbf{Y}^{(2)}\|_2^2 - \|(\mathbf{I}_{n-m} - \mathbf{P}_{\mathbf{Z}_2}) \mathbf{Y}^{(2)}\|_2^2 = \|(\mathbf{P}_{\mathbf{Z}_2} - \mathbf{P}_{\mathbf{Z}_1}) \mathbf{Y}^{(2)}\|_2^2.$$

We finally summarize the asymptotic behaviors of our test statistic in the following Corollary F.3.

Corollary F.3. Let Assumptions 2.1–2.5 and Assumption F.1 hold with

$$s_n^M \left(\frac{\log d}{n} + \frac{1}{d} \right) \rightarrow 0 \text{ and } \Delta_\sigma \rightarrow 0,$$

where Δ_σ is given in Assumption 3.3. Then we have

$$\sup_{x>0} \left| \mathbb{P} \left(Q_{n,M}^{(2)} / \hat{\sigma}^2 \leq x | H_0 \right) - \mathbb{P} \left(\chi_{K^*}^2 \leq x \right) \right| \rightarrow 0, \text{ as } n \rightarrow \infty,$$

in which $K^* = \sum_{i=1}^L K_i$. In addition, we define $\varphi = (\varphi_1, \varphi_2, \dots, \varphi_L)$ and

$$\mathcal{D}^*(\alpha, \theta) = \left\{ \varphi \in \mathbb{R}^K : \frac{n \|\varphi\|^2}{1 + K^* s_n^M \Upsilon^2 / \lambda_{\min}(\Sigma_u)} \geq \sigma^2 (2 + \delta) (\chi_{K^*, 1-\alpha}^2 + \chi_{K^*, 1-\theta}^2) \right\},$$

where $\delta > 0$ is some constant. Then when $s_n^M = o(n^{1/6})$, we have

$$\inf_{\varphi^* \in \mathcal{D}^*(\alpha, \theta)} \mathbb{P}(\psi_\alpha^* = 1 | H_1) \geq 1 - \theta, \quad (\text{F.12})$$

in which

$$\psi_\alpha^* = \mathbb{I} \left\{ Q_{n,M}^{(2)} / \hat{\sigma}^2 \geq \chi_{K^*, 1-\alpha}^2 \right\}.$$

F.5 Proof of Corollary F.3

Proof. Here we assume $(\tilde{S}_1, \dots, \tilde{S}_M)$ are independent with $(\mathbf{X}_1, \dots, \mathbf{X}_M, \mathbf{Y})$.

First, $\mathbf{P}_{\mathbf{Z}_2} - \mathbf{P}_{\mathbf{Z}_1}$ is a projection matrix onto the space of

$$(\mathbf{I} - \mathbf{P}_{\mathbf{Z}_1})\mathbf{Z}_2.$$

Recall that

$$\mathbf{Z}_2 = (\hat{\mathbf{F}}_1^{(2)}, \hat{\mathbf{U}}_{\hat{S}_1}^{(2)}, \dots, \hat{\mathbf{F}}_L^{(2)}, \hat{\mathbf{U}}_{\hat{S}_L}^{(2)}, \hat{\mathbf{F}}_{L+1}^{(2)}, \hat{\mathbf{U}}_{\hat{S}_{L+1}}^{(2)}, \dots, \hat{\mathbf{F}}_M^{(2)}, \hat{\mathbf{U}}_{\hat{S}_M}^{(2)}).$$

Note that for each $1 \leq l \leq L$,

$$(\mathbf{I} - \mathbf{P}_{\mathbf{Z}_1})(\hat{\mathbf{F}}_l^{(2)}, \hat{\mathbf{U}}_{\hat{S}_l}^{(2)}) = (\mathbf{I} - \mathbf{P}_{\mathbf{Z}_1})(\hat{\mathbf{F}}_l^{(2)}, \mathbf{X}_{\hat{S}_l}^{(2)} - \hat{\mathbf{F}}_l^{(2)} \hat{\mathbf{B}}_{\hat{S}_l}^{(2)}).$$

For any vector ϕ , if we there exists an $\tilde{\phi}$ such that $\mathbf{Z}_1 \phi = \mathbf{Z}_2 \tilde{\phi}$, we obtain $\text{col}(\mathbf{Z}_1) \subset \text{col}(\mathbf{Z}_2)$. Thus, we have $\mathbf{P}_{S_M^*} := \mathbf{P}_{\mathbf{Z}_2} - \mathbf{P}_{\mathbf{Z}_1}$ is a projection matrix onto the column space spanned by $\tilde{\mathbf{F}}^* := (\mathbf{I}_n - \mathbf{P}_{\mathbf{Z}_1})(\hat{\mathbf{F}}_1, \dots, \hat{\mathbf{F}}_L) \in \mathbb{R}^{n \times K^*}$, in which $K^* = \sum_{i=1}^L K_i$. $\tilde{\mathbf{F}}^* := (\mathbf{I}_n - \mathbf{P}_{\mathbf{Z}_1})(\hat{\mathbf{F}}_1, \dots, \hat{\mathbf{F}}_L) \in \mathbb{R}^{n \times K^*}$.

Similar with the proof of Lemma F.1, we are able to write $\mathbf{P}_{S_M^*} = \mathbf{W}^* \mathbf{W}^{*\top}$, and we also obtain

$$\sup_{x>0} \left| \mathbb{P}(\mathcal{E}^\top \mathbf{P}_{S_M^*} \mathcal{E} \leq x | \mathbf{X}) - \mathbb{P}(\chi_K^{*2} \leq x) \right| \leq cK^{*1/4} \mathbb{E}|\varepsilon_t|^3 \max_{t \in [n]} \|\mathbf{w}_t^*\| \sum_{t=1}^n \|\mathbf{w}_t^*\|^2$$

with $\max_{t \in [n]} \|\mathbf{w}_t^*\| = \max_{t \in [n]} \sqrt{\mathbf{P}_{S_M^*, tt}}$.

Recall we have $\mathbf{P}_{S_M^*} := \mathbf{P}_{\mathbf{Z}_2} - \mathbf{P}_{\mathbf{Z}_1}$ and

$$\mathbf{Z}_2 = (\hat{\mathbf{F}}_1^{(2)}, \hat{\mathbf{U}}_{\hat{S}_1}^{(2)}, \dots, \hat{\mathbf{F}}_L^{(2)}, \hat{\mathbf{U}}_{\hat{S}_L}^{(2)}, \hat{\mathbf{F}}_{L+1}^{(2)}, \hat{\mathbf{U}}_{\hat{S}_{L+1}}^{(2)}, \dots, \hat{\mathbf{F}}_M^{(2)}, \hat{\mathbf{U}}_{\hat{S}_M}^{(2)}).$$

We first denote the event

$$\mathcal{A}_{\mathbf{Z}_2} = \left\{ \lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2) \geq \frac{n\lambda_0}{2} \right\},$$

in which λ_0 is a absolute constant. We know from (F.4), in order to prove Corollary F.3, we only need to prove (i). $\mathbb{E}[\max_{t \in [n]} \sqrt{\mathbf{P}_{S_M^*, tt}} \mathbb{I}_{\mathcal{A}_{\mathbf{Z}_2}}] \rightarrow 0$ and (ii). $\mathbb{P}(\mathcal{A}_{\mathbf{Z}_2}^c) \rightarrow 0$.

Next, we prove the term (i).

Since $\mathbf{P}_{\mathbf{Z}_1}$ is a projection matrix with $\mathbf{P}_{\mathbf{Z}_1, tt} > 0$ for any $t \in [n]$, we have

$$\max_{t \in [n]} \sqrt{\mathbf{P}_{S_M^*, tt}} \leq \max_{t \in [n]} \sqrt{\mathbf{P}_{\mathbf{Z}_2, tt}}.$$

The next step is prove that $\mathbb{E}[\max_{t \in [n]} \sqrt{\mathbf{P}_{\mathbf{Z}_2, tt}} \mathbb{I}_{\mathcal{A}_{\mathbf{Z}_2}}] \rightarrow 0$. We first prove that

$$\begin{aligned} \max_{t \in [n]} \mathbf{P}_{\mathbf{Z}_2, tt} &\leq \frac{1}{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)} \max_{t \in [n]} \left(\sum_{i=1}^M |\mathbf{f}_t^{(i)}|^2 + \sum_{i=1}^M \sum_{\ell \in \hat{S}_i} |\hat{u}_{t\ell}|^2 \right) \\ &\leq \frac{1}{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)} \max_{t \in [n]} \sum_{i=1}^M |\mathbf{f}_t^{(i)}|^2 + \max_{t \in [n]} \sum_{i=1}^M \sum_{\ell \in \hat{S}_i} |\hat{u}_{t\ell}|^2 \end{aligned}$$

Leveraging similar techniques given in section F.3, we have

$$\mathbb{E} \left[\max_{t \in [n]} \sqrt{\mathbf{P}_{S_M^*, tt}} \mathbb{I}_{\mathcal{A}_{\mathbf{Z}_2}} \right] \leq \mathbb{E} \left[\max_{t \in [n]} \sqrt{\mathbf{P}_{\mathbf{Z}_2, tt}} \mathbb{I}_{\mathcal{A}_{\mathbf{Z}_2}} \right] \rightarrow 0$$

Next, using certain concentration inequalities given in section F.3, we prove that there exists an positive constant λ_0 such that

$$\mathbb{P}(\mathcal{A}_{\mathbf{Z}_2}^c) := \mathbb{P} \left(\frac{\lambda_{\min}(\mathbf{Z}_2^\top \mathbf{Z}_2)}{2} \leq \frac{n\lambda_0}{2} \right) \rightarrow 0.$$

Thus, we claim our conclusion for the first part of Corollary F.3.

We now prove (F.12). The proof details are the almost the same with the proof of Lemma F.2, so we only describe the outline here. We let

$$\mathcal{Z}_2 = (\hat{\mathbf{F}}_1^{(2)}, \dots, \hat{\mathbf{F}}_M^{(2)}, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_1}^{(2)}, \dots, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_M}^{(2)}).$$

Note that $\mathbf{P}_{\mathcal{Z}_2} = \mathbf{P}_{\mathcal{Z}_2}$. Hence it is equivalent to bound $\max_{t \in [n]} \mathbf{P}_{\mathcal{Z}_2, tt}$. Let

$$\mathcal{Z}_1 = (\mathbf{X}_{\hat{\mathcal{S}}_1}, \dots, \mathbf{X}_{\hat{\mathcal{S}}_L}, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_{L+1}}, \dots, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_M}, \hat{\mathbf{F}}_{L+1}, \dots, \hat{\mathbf{F}}_M).$$

Then $\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1} = \mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1}$, which is a projection matrix onto the column space of $(\mathbf{I} - \mathbf{P}_{\mathcal{Z}_1})(\hat{\mathbf{F}}_1, \dots, \hat{\mathbf{F}}_L) =: (\mathbf{I} - \mathbf{P}_{\mathcal{Z}_1})\hat{\mathbb{F}}$. Here we denote $\hat{\mathbb{F}} := (\hat{\mathbf{F}}_1, \dots, \hat{\mathbf{F}}_L)$.

Under the alternative hypothesis, $\mathbf{Y} = \sum_{m=1}^L (\mathbf{F}_m \boldsymbol{\varphi}_m^* + \mathbf{X}_m \boldsymbol{\beta}_m^*) + \sum_{m=L+1}^M (\mathbf{F}_m \boldsymbol{\varphi}_m^* + \mathbf{X}_m \boldsymbol{\beta}_m^*) + \mathcal{E}$. We have,

$$\begin{aligned} \|(\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1}) \sum_{m=L+1}^M (\mathbf{F}_m \boldsymbol{\varphi}_m^* + \mathbf{X}_m \boldsymbol{\beta}_m^*)\|^2 &= \|(\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1}) \sum_{m=L+1}^M \mathbf{F}_m \boldsymbol{\varphi}_m^*\|^2 \\ &= o(1) + \|(\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1}) \sum_{m=L+1}^M \hat{\mathbf{F}}_m \mathbf{H}_m \boldsymbol{\varphi}_m^*\|^2 = o(1) \end{aligned}$$

and

$$\|(\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1})\mathcal{E}\|^2 \sim \chi_{K^*}^2,$$

where $K^* = \sum_{i=1}^L K_i$. According to the definition of $\mathcal{Z}_1$, we obtain

$$\begin{aligned} \mathcal{Z}_1 &= (\mathbf{X}_{\hat{\mathcal{S}}_1}, \dots, \mathbf{X}_{\hat{\mathcal{S}}_L}, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_{L+1}}, \dots, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_M}, \hat{\mathbf{F}}_{L+1}, \dots, \hat{\mathbf{F}}_M) \\ &= (\hat{\mathbf{U}}_{\hat{\mathcal{S}}_1}, \dots, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_L}, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_{L+1}}, \dots, \hat{\mathbf{U}}_{\hat{\mathcal{S}}_M}, \hat{\mathbf{F}}_{L+1}, \dots, \hat{\mathbf{F}}_M) \\ &\quad + (\hat{\mathbf{F}}_1, \dots, \hat{\mathbf{F}}_L) \begin{pmatrix} \hat{\mathbf{B}}_{1, \hat{\mathcal{S}}_1}^\top & \cdots & \cdots & \\ & 0 & & \mathbf{O} \\ & & \hat{\mathbf{B}}_{L, \hat{\mathcal{S}}_L}^\top & \end{pmatrix} =: \hat{\mathbf{U}} + \hat{\mathbb{F}} \hat{\mathbb{B}}^\top \end{aligned}$$

In addition, we let $\mathbb{F} := (\mathbf{F}_1, \dots, \mathbf{F}_L)$ and $\mathbb{H} = (\mathbf{H}_1, \dots, \mathbf{H}_L)$. We obtain

$$\begin{aligned} \|(\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1}) \sum_{m=1}^L (\mathbf{F}_m \boldsymbol{\varphi}_m^* + \mathbf{X}_m \boldsymbol{\beta}_m^*)\|^2 &= \|(\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1}) \sum_{m=1}^L \mathbf{F}_m \boldsymbol{\varphi}_m^*\|^2 = \|(\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1})\mathbb{F} \boldsymbol{\varphi}^*\|^2 \\ &= o(1) + \|(\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1})\hat{\mathbb{F}} \mathbb{H} \boldsymbol{\varphi}^*\|^2 = o(1) + (\mathbb{H} \boldsymbol{\varphi}^*)^\top \hat{\mathbb{F}}^\top (\mathbf{P}_{\mathcal{Z}_2} - \mathbf{P}_{\mathcal{Z}_1}) \hat{\mathbb{F}} (\mathbb{H} \boldsymbol{\varphi}^*) \\ &= o(1) + (\mathbb{H} \boldsymbol{\varphi}^*)^\top \hat{\mathbb{F}}^\top (\mathbf{I} - \mathbf{P}_{\mathcal{Z}_1}) \hat{\mathbb{F}} (\mathbb{H} \boldsymbol{\varphi}^*) \\ &= (\mathbb{H} \boldsymbol{\varphi}^*)^\top \hat{\mathbb{F}}^\top (\mathbf{I} - \hat{\mathbb{F}} \hat{\mathbb{B}}^\top (\hat{\mathbb{B}} \hat{\mathbb{F}}^\top \hat{\mathbb{B}}^\top + \hat{\mathbf{U}}^\top \hat{\mathbf{U}})^{-1} \hat{\mathbb{B}} \hat{\mathbb{F}}^\top) \hat{\mathbb{F}} (\mathbb{H} \boldsymbol{\varphi}^*) \end{aligned}$$

$$\begin{aligned}
&= (\mathbf{H}\boldsymbol{\varphi}^*)^\top \hat{\mathbf{F}}^\top (\mathbf{I} + \hat{\mathbf{F}}\hat{\mathbf{B}}^\top (\hat{\mathbf{U}}^\top \hat{\mathbf{U}})^{-1} \hat{\mathbf{B}}\hat{\mathbf{F}}^\top)^{-1} \hat{\mathbf{F}}(\mathbf{H}\boldsymbol{\varphi}^*) \\
&\geq \frac{\|\hat{\mathbf{F}}(\mathbf{H}\boldsymbol{\varphi}^*)\|^2}{1 + \lambda_{\max}(\hat{\mathbf{F}}\hat{\mathbf{B}}^\top (\hat{\mathbf{U}}^\top \hat{\mathbf{U}})^{-1} \hat{\mathbf{B}}\hat{\mathbf{F}}^\top)} \geq \frac{\lambda_{\min}(\hat{\mathbf{F}}^\top \hat{\mathbf{F}}) \|\boldsymbol{\varphi}^*\|^2}{1 + \frac{\lambda_{\max}(\hat{\mathbf{B}}^\top \hat{\mathbf{B}}) \lambda_{\max}(\hat{\mathbf{F}}\hat{\mathbf{F}}^\top)}{\lambda_{\min}(\hat{\mathbf{U}}^\top \hat{\mathbf{U}})}}.
\end{aligned}$$

The remaining proof steps follow the same procedure of Lemma F.2. □

References

- S. Basu and G. Michailidis. Regularized estimation in sparse high-dimensional time series models. *The Annals of Statistics*, 43(4):1535–1567, 2015.
- V. Bentkus. A lyapunov-type bound in R^d . *Theory of Probability & Its Applications*, 49(2):311–323, 2005.
- P. J. Bickel, Y. Ritov, and A. B. Tsybakov. Simultaneous analysis of lasso and Dantzig selector. *Ann. Statist.*, 37(4):1705–1732, 2009.
- V. Chernozhukov, D. Chetverikov, and K. Kato. Central limit theorems and bootstrap in high dimensions. *Ann. Probab.*, 45(4):2309–2352, 2017.
- V. Chernozhukov, D. Chetverikov, and Y. Koike. Nearly optimal central limit theorem and bootstrap approximations in high dimensions. *arXiv preprint arXiv:2012.09513*, 2020.
- Y. Deshpande, A. Javanmard, and M. Mehrabi. Online debiasing for adaptively collected high-dimensional data with applications to time series analysis. *Journal of the American Statistical Association*, pages 1–14, 2021.
- J. Fan, Y. Ke, and K. Wang. Factor-adjusted regularized model selection. *J. Econometrics*, 216(1):71–85, 2020.
- D. Giannone, M. Lenza, and G. E. Primiceri. Economic predictions with big data: The illusion of sparsity. *Econometrica*, 89(5):2409–2437, 2021.
- F. Han, H. Lu, and H. Liu. A direct estimation of high dimensional stationary vector autoregressions. *Journal of Machine Learning Research*, 2015.
- A. Kneip and P. Sarda. Factor models and variable selection in high-dimensional regression analysis. *The Annals of Statistics*, 39(5):2410–2447, 2011.
- Q. Li and L. Li. Integrative factor regression and its inference for multimodal data analysis. *J. Amer. Statist. Assoc.*, 113(521):1–15, 2021.

- J. Lin and G. Michailidis. System identification of high-dimensional linear dynamical systems with serially correlated output noise components. *IEEE Transactions on Signal Processing*, 68:5573–5587, 2020.
- P.-L. Loh and M. J. Wainwright. High-dimensional regression with noisy and missing data: Provable guarantees with nonconvexity. *Ann. Statist.*, 40(3):1637–1664, 2012.
- R. P. Masini, M. C. Medeiros, and E. F. Mendes. Regularized estimation of high-dimensional vector autoregressions with weakly dependent innovations. *Journal of Time Series Analysis*, 43(4):532–557, 2022.
- Q. Sun, W.-X. Zhou, and J. Fan. Adaptive Huber regression. *Journal of the American Statistical Association*, 115(529):254–265, 2020.
- R. Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society: Series B (Methodological)*, 58(1):267–288, 1996.
- S. van de Geer, P. Bühlmann, Y. Ritov, and R. Dezeure. On asymptotically optimal confidence regions and tests for high-dimensional models. *Ann. Statist.*, 42(3):1166–1202, 2014.
- R. Vershynin. Introduction to the non-asymptotic analysis of random matrices. In *Compressed sensing*, pages 210–268. Cambridge Univ. Press, Cambridge, 2012.
- K. C. Wong, Z. Li, and A. Tewari. Lasso guarantees for β -mixing heavy-tailed time series. *The Annals of Statistics*, 48(2):1124–1142, 2020.
- W.-B. Wu and Y. N. Wu. Performance bounds for parameter estimates of high-dimensional linear models with correlated errors. *Electronic Journal of Statistics*, 10(1):352–379, 2016.
- Y. Yan, Y. Chen, and J. Fan. Inference for heteroskedastic pca with missing data. *arXiv preprint arXiv:2107.12365*, 2021.
- P. Zhao and B. Yu. On model selection consistency of Lasso. *J. Mach. Learn. Res.*, 7:2541–2563, 2006.
- H. Zou and T. Hastie. Regularization and variable selection via the elastic net. *Journal of the royal statistical society: series B (statistical methodology)*, 67(2):301–320, 2005.