

# Statistical Inference for High-dimensional Partially Linear Models via Debiased Rank Lasso

**Songshan Yang**

YANGSS@RUC.EDU.CN

*Center for Applied Statistics, BRAIN and Institute of Statistics and Big Data  
Renmin University of China  
Beijing, 100872, China*

**Delin Zhao\***

DLZHAO@FZU.EDU.CN

*School of Mathematics and Statistics  
Fuzhou University  
Fuzhou, 350108, China*

**Runze Li**

RZLI@PSU.EDU

*Department of Statistics  
Pennsylvania State University  
University Park, PA 16802-2111, USA*

**Editor:** Ji Zhu

## Abstract

This paper aims to develop tuning-free and robust regularized methods for partially linear models based on partial residual methods. In order to preserve the near-oracle rate of rank Lasso, we construct a new estimator via a data-splitting procedure and name the resulting estimator as data-splitting (DS) rank Lasso, which is based on partial **prediction** residuals. Thus, the proposed estimation procedure is distinguished from the traditional partial residual based methods, which are based on the model-fitting residuals. We show that the DS rank Lasso enjoys several appealing features including tuning-free property and robust property to heavy-tailed random errors, and possesses near-oracle rate. We further propose statistical inference procedures for individual coefficients using the debiased Lasso techniques. The resulting estimator is termed the debiased DS rank Lasso estimator. This paper reveals an interesting theoretical finding: the DS rank Lasso estimate achieves near-oracle rate and the proposed debiased DS rank Lasso estimator does not suffer asymptotic efficiency loss due to the use of data-splitting. We further establish a simultaneous inference procedure based on multiplier bootstrap. We conduct Monte Carlo studies to assess finite sample performance of the proposed procedures, and illustrate the proposed methodology via an empirical analysis of a real-world data set.

**Keywords:** Data splitting, heavy-tailed error, tuning-free regularization, variable selection

## 1. Introduction

Statistical inference for high-dimensional linear and generalized linear models receives increasing attention in the recent literature. See, for example, Chapter 7 of Fan et al. (2020a) and references therein for a comprehensive account of this topic. Zhang and Zhang (2014)

---

\*: Corresponding author.

and van de Geer et al. (2014) proposed debiased and desparsified Lasso to construct inference procedures for low-dimensional regression coefficients based on the Lasso estimator. A decorrelated score method was proposed by Ning and Liu (2017) to derive an inference procedure for low-dimensional regression coefficients based on penalized M-estimators. To have better finite sample performance, bootstrap-assisted procedures were used in Zhang and Cheng (2017) and Dezeure et al. (2017) to conduct simultaneous inference based on the debiased Lasso estimator. Inference on low-dimensional regression coefficients in high-dimensional linear models has been considered by Belloni et al. (2014), Cai and Guo (2017), Zhu and Bradic (2018a), Zhu and Bradic (2018b), Neykov et al. (2018), Javanmard and Montanari (2018), Chang et al. (2021) and among many others. To address the limitation that existing inference methods are not robust against heavy-tailed errors, several important works have proposed distinct approaches: Feng et al. (2013) developed a rank-based score test, Fan et al. (2020b) introduced a debiased rank Lasso estimator, and Cai et al. (2025) proposed a debiased convoluted rank Lasso estimator. This paper aims to develop robust statistical inference procedures for high-dimensional partially linear models.

Let  $Y$  be a response along with a  $p$ -dimensional covariate vector  $\mathbf{x}$  and a univariate covariate  $Z$ . The partially linear model was first introduced by Engle et al. (1986) and is defined by

$$Y = \mathbf{x}^T \boldsymbol{\beta}_0 + g_0(Z) + \epsilon, \quad (1)$$

where  $g_0(\cdot)$  is a nonparametric baseline function, and  $\epsilon$  is a random error. Compared with linear models and generalized linear models, there is much less work on inference for partially linear models. Zhu (2017) conducted a nonasymptotic analysis, for a two-step projection based Lasso estimator under the setting of  $p \gg n$ , and Zhu et al. (2019) extended the debiased Lasso for linear models to partially linear models by using partial residual techniques. Both works are based on penalized least squares with the Lasso penalty. It is well-known that least squares is not robust to outliers and heavy-tailed random errors. This motivates us to consider penalized rank regression (Wang et al., 2020) to achieve robustness. Furthermore, rank Lasso enjoys the tuning-free property for high-dimensional linear models.

In this paper, we first aim to develop a tuning-free and robust regularized regression method for partially linear models based on partial residual methods. In order to preserve the near-oracle rate of rank Lasso under mild conditions on nonparametric estimators, we construct a new estimator via a data-splitting procedure. Specifically, we randomly partition the entire sample into two subsamples with equal sizes, and employ the first-half subsample to estimate the nonparametric functions, and then we obtain a rank Lasso estimate based on the partial **prediction** residuals of the second-half subsample. See Section 2 for the definition of partial prediction residuals. We further switch the roles of the two subsamples and obtain another rank Lasso estimate. Thus, a natural estimator of  $\boldsymbol{\beta}_0$  is the average of the two estimators from the data-splitting procedure. This estimator is termed data-splitting (DS) rank Lasso. We show that the DS rank Lasso enjoys several appealing properties including the tuning-free property and robustness to heavy-tailed random errors. We systematically study its asymptotic properties and establish its near-oracle rate.

Another goal of this paper is to develop robust statistical inference procedures for  $\boldsymbol{\beta}_0$ . We apply the debiased Lasso technique to the proposed DS rank Lasso and propose a new

estimator, termed the debiased DS rank Lasso estimator. We first establish the asymptotic normality of the debiased DS rank Lasso estimator. The asymptotic normality reveals an interesting theoretical finding: the proposed debiased DS rank Lasso estimator does not suffer asymptotic efficiency loss despite the use of data-splitting. Based on the asymptotic normality of the estimator, we can construct robust confidence intervals for either individual coefficients or any finite-dimensional subvector of  $\beta_0$ . We further develop a simultaneous inference procedure based on the multiplier bootstrap. We show that the two proposed inference procedures are robust to heavy-tailed random errors. Furthermore, a comprehensive numerical study demonstrates the effectiveness of the proposed estimation and inference procedures.

This paper is organized as follows. In Section 2, we propose the DS rank Lasso estimator for high-dimensional partially linear models and establish its theoretical properties. In Section 3, we propose a debiased DS rank Lasso estimator for statistical inference for individual coefficients in Section 3.1 and then develop a simultaneous inference procedure in Section 3.2. Numerical comparisons are presented in Section 4. Some discussion and concluding remarks are given in Section 5. All technical proofs are included in the Appendix.

## 2. Rank Lasso for Partially Linear Models

It is assumed throughout this paper that  $\mathbf{m}_{\mathbf{x}}(z) = E(\mathbf{x} \mid Z = z)$  is well-defined and finite. Furthermore, we focus on settings where  $m_0(z) = E(Y \mid Z = z)$  is also well-defined and finite. Based on model (1), it follows that

$$E(Y \mid Z) = E(\mathbf{x} \mid Z)^T \beta_0 + g_0(Z), \quad (2)$$

since it is assumed that  $E(\epsilon \mid Z) = 0$ . It follows from (2) that

$$Y - m_0(Z) = \{\mathbf{x} - \mathbf{m}_{\mathbf{x}}(Z)\}^T \beta_0 + \epsilon.$$

This leads to the partial residual method for estimating  $\beta_0$ . Specifically, denote  $\tilde{Y} = Y - m_0(Z)$  and  $\tilde{\mathbf{x}} = \mathbf{x} - \mathbf{m}_{\mathbf{x}}(Z)$ . If both  $m_0(Z)$  and  $\mathbf{m}_{\mathbf{x}}(Z)$  were known, we could directly apply rank Lasso proposed in Wang et al. (2020) for

$$\tilde{Y} = \tilde{\mathbf{x}}^T \beta_0 + \epsilon.$$

In practice, both  $E(Y \mid Z)$  and  $E(\mathbf{x} \mid Z)$  are unknown. Intuitively, we may obtain estimates of  $m_0(\cdot)$  and  $\mathbf{m}_{\mathbf{x}}(\cdot)$  by regressing  $Y$  and  $\mathbf{x}$  on  $Z$ , respectively. In particular, one may use nonparametric regression techniques to estimate  $m_0(\cdot)$  and  $\mathbf{m}_{\mathbf{x}}(\cdot)$  so that the forms of the regression functions are flexible. Once we obtain the estimated partial residuals from the nonparametric regression, we can apply the rank Lasso to these residuals. To allow flexibility in choosing the nonparametric estimators, we only assume the convergence rates of these estimators without specifying their exact forms. However, overfitting in the nonparametric regression may lead to poor performance of the rank Lasso if it uses the same data as the nonparametric regression. We propose using data-splitting to deal with this challenge.

### 2.1 Data-Splitting Rank Lasso

We randomly split the data into two equal halves,  $\mathcal{D}_1$  and  $\mathcal{D}_2$ . For simplicity, assume the sample size  $N$  is even, let  $n = N/2$ , and denote  $\mathcal{D}_1 = \{\mathbf{x}_i^{(1)}, Y_i^{(1)}, Z_i^{(1)}, i = 1, \dots, n\}$  and

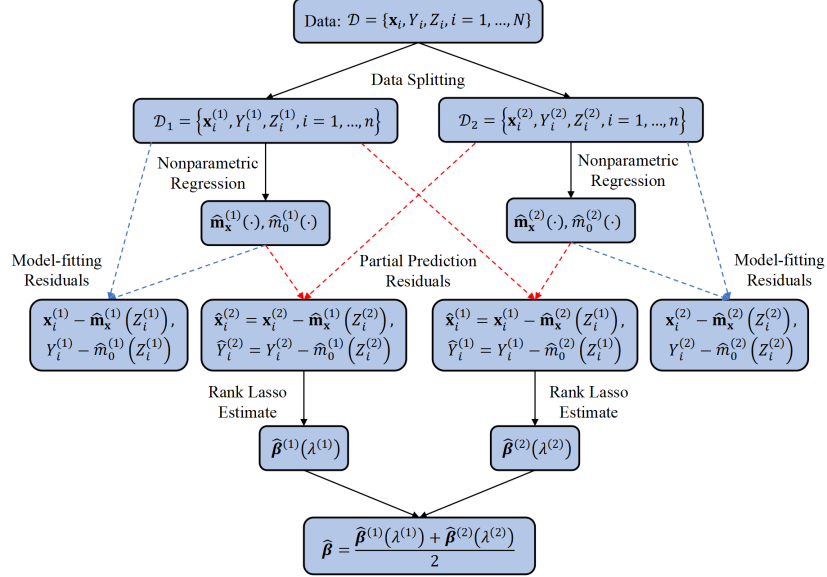


Figure 1: Flowchart of the procedure for constructing the DS rank Lasso estimator.

$\mathcal{D}_2 = \{\mathbf{x}_i^{(2)}, Y_i^{(2)}, Z_i^{(2)}, i = 1, \dots, n\}$ . The strategy of the data-splitting rank Lasso can be described as follows. We use the first half of the data to estimate  $m_0(\cdot)$  and  $\mathbf{m}_x(\cdot)$ , obtaining the nonparametric estimators  $\hat{m}_0^{(1)}(\cdot)$  and  $\hat{\mathbf{m}}_x^{(1)}(\cdot)$ . Next, we use the second half of the data to calculate the partial prediction residuals  $\hat{Y}_i^{(2)} = Y_i^{(2)} - \hat{m}_0^{(1)}(Z_i^{(2)})$  and  $\hat{\mathbf{x}}_i^{(2)} = \mathbf{x}_i^{(2)} - \hat{\mathbf{m}}_x^{(1)}(Z_i^{(2)})$ ,  $i = 1, \dots, n$ . Here, we introduce the term “**partial prediction residuals**” to distinguish these residuals  $\{\hat{Y}_i^{(2)}, \hat{\mathbf{x}}_i^{(2)}\}_{i=1}^n$  from the model-fitting residuals  $\{Y_i^{(1)} - \hat{m}_0^{(1)}(Z_i^{(1)}), \mathbf{x}_i^{(1)} - \hat{\mathbf{m}}_x^{(1)}(Z_i^{(1)})\}_{i=1}^n$ . These partial prediction residuals are then used to construct the rank Lasso estimator  $\hat{\beta}^{(1)}(\lambda^{(1)})$ , defined as

$$\hat{\beta}^{(1)}(\lambda^{(1)}) = \arg \min_{\beta \in \mathbb{R}^p} \left[ \{n(n-1)\}^{-1} \sum_{i \neq j} |(\hat{Y}_i^{(2)} - \beta^T \hat{\mathbf{x}}_i^{(2)}) - (\hat{Y}_j^{(2)} - \beta^T \hat{\mathbf{x}}_j^{(2)})| + \lambda^{(1)} \|\beta\|_1 \right],$$

where  $\lambda^{(1)}$  is a tuning parameter to be determined by a data-driven method. Similarly, we may switch the roles of  $\mathcal{D}_1$  and  $\mathcal{D}_2$  and obtain the second rank Lasso estimator  $\hat{\beta}^{(2)}(\lambda^{(2)})$ . The data-splitting (DS) rank Lasso estimator is then defined as

$$\hat{\beta} = \frac{\hat{\beta}^{(1)}(\lambda^{(1)}) + \hat{\beta}^{(2)}(\lambda^{(2)})}{2},$$

where we use  $\lambda^{(1)}$  and  $\lambda^{(2)}$  to allow flexibility in choosing the tuning parameters. Figure 1 illustrates the overall procedure for constructing the DS rank Lasso estimator.

We next show that the proposed DS rank Lasso enjoys the tuning-free property (Wang et al., 2020). Denote  $g(z, \beta_0) = E(Y | Z = z) - \beta_0^T E(\mathbf{x} | Z = z) = m_0(z) - \beta_0^T \mathbf{m}_x(z) = g_0(z)$

and  $\hat{g}^{(1)}(z, \beta_0) = \hat{m}_0^{(1)}(z) - \beta_0^\top \hat{\mathbf{m}}_{\mathbf{x}}^{(1)}(z)$ . Further define  $e_i^{(2)} = \epsilon_i^{(2)} - \{\hat{g}^{(1)}(Z_i^{(2)}, \beta_0) - g_0(Z_i^{(2)})\}$ . Thus,  $\hat{m}_0^{(1)}(\cdot)$ ,  $\hat{\mathbf{m}}_{\mathbf{x}}^{(1)}(\cdot)$  and  $\hat{g}^{(1)}(\cdot, \beta_0)$  are independent of  $\{\epsilon_i^{(2)}, i = 1, \dots, n\}$ . This is critical for DS rank Lasso to retain the tuning-free property of the rank Lasso for linear regression models. Denote  $\gamma = \beta - \beta_0$ . By the definition of  $\hat{Y}_i^{(2)}$  and  $\hat{\mathbf{x}}_i^{(2)}$ ,

$$\begin{aligned} & \{n(n-1)\}^{-1} \sum_{i \neq j} |(\hat{Y}_i^{(2)} - \beta^\top \hat{\mathbf{x}}_i^{(2)}) - (\hat{Y}_j^{(2)} - \beta^\top \hat{\mathbf{x}}_j^{(2)})| \\ &= \{n(n-1)\}^{-1} \sum_{i \neq j} |(e_i^{(2)} - e_j^{(2)}) - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma| =: Q^{(2)}(\gamma). \end{aligned}$$

Because  $Q^{(2)}(\gamma)$  is piecewise linear in  $\gamma$ , its gradient is defined everywhere excluding the linear boundaries. The gradient of  $Q^{(2)}(\gamma)$  with respect to  $\gamma$  is

$$\begin{aligned} \mathbf{S}^{(2)}(\gamma) &= -\{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) \text{sign}\{(e_i^{(2)} - e_j^{(2)}) - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\} \\ &= -2\{n(n-1)\}^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(2)} \sum_{j=1, j \neq i}^n \text{sign}\{(e_i^{(2)} - e_j^{(2)}) - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\}. \end{aligned}$$

Note that  $\sum_{j=1, j \neq i}^n \text{sign}\{(e_i^{(2)} - e_j^{(2)}) - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\}$  depends on the rank of  $e_i^{(2)} - \gamma^\top \hat{\mathbf{x}}_i^{(2)}$  over  $\{e_j^{(2)} - \gamma^\top \hat{\mathbf{x}}_j^{(2)}, j = 1, \dots, n\}$ . Motivated by the tuning parameter selection framework in Wang et al. (2020), we select  $\lambda^{(1)}$  to satisfy

$$\text{pr}\{\lambda^{(1)} \geq c \|\mathbf{S}^{(2)}(\mathbf{0})\|_\infty\} \geq 1 - \alpha_0$$

for constants  $c > 1$  and  $\alpha_0 > 0$ . Under mild conditions on nonparametric estimators,  $\mathbf{S}^{(2)}(\mathbf{0})$  and  $\mathbf{S}_0^{(2)}$  (defined below) are asymptotically equivalent with a difference of  $o_p(\sqrt{\log p/n})$ :

$$\mathbf{S}_0^{(2)} = -2\{n(n-1)\}^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(2)} \sum_{j=1, j \neq i}^n \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)}),$$

where  $\{\sum_{j=1, j \neq i}^n \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)})\}_{i=1}^n$  follows a uniform distribution over permutations of  $2\{1, \dots, n\} - (n+1) = \{-n+1, -n+3, \dots, n-3, n-1\}$ . Following Wang et al. (2020), let  $G_{(1)}^{-1}(1 - \alpha_0)$  denote the  $(1 - \alpha_0)$ -quantile of the conditional distribution of  $\|\mathbf{S}_0^{(2)}\|_\infty$  given  $\hat{\mathbf{x}}_1^{(2)}, \dots, \hat{\mathbf{x}}_n^{(2)}$ , where  $\|\cdot\|_\infty$  is the infinity norm. The quantile  $G_{(1)}^{-1}(1 - \alpha_0)$  can be easily simulated. We then set the tuning parameter  $\lambda^{(1)}$  for  $\hat{\beta}^{(1)}(\lambda^{(1)})$  as

$$\hat{\lambda}^{(1)} = c \left\{ G_{(1)}^{-1}(1 - \alpha_0) + c_0 \sqrt{\log p/n} \right\}, \quad (3)$$

where  $c > 1$ , and  $c_0, \alpha_0 > 0$  are pre-specified constants. By swapping the roles of  $\mathcal{D}_1$  and  $\mathcal{D}_2$ , we analogously obtain  $\hat{\lambda}^{(2)}$  for  $\hat{\beta}^{(2)}(\lambda^{(2)})$ . The final DS rank Lasso estimator is

$$\hat{\beta} = \frac{\hat{\beta}^{(1)}(\hat{\lambda}^{(1)}) + \hat{\beta}^{(2)}(\hat{\lambda}^{(2)})}{2}.$$

## 2.2 Rates of Convergence

Let  $q = \|\beta_0\|_0$ , where  $\|\cdot\|_0$  denotes the  $\ell_0$  norm. Define  $\mathcal{B}$  as the index set of the non-zero components of  $\beta_0$ , and introduce the following cone set:

$$\Gamma = \{\gamma \in \mathbb{R}^p : \|\gamma_{\mathcal{B}^c}\|_1 \leq \bar{c}\|\gamma_{\mathcal{B}}\|_1\},$$

where  $\bar{c} = (c+1)/(c-1)$  and  $c$  is the constant defined in (3). We denote the cumulative distribution function and probability density function of  $\epsilon_i - \epsilon_j$  (for  $i \neq j$ ) as  $F^*(\cdot)$  and  $f^*(\cdot)$ , respectively.

To establish the theoretical properties of the DS Rank Lasso estimator, we impose the following regularity conditions.

(C1) There exist sequences  $r_n > 0$  and  $\rho_n = o(1)$  such that for  $l \in \{1, 2\}$ ,

$$\text{pr} \left\{ \max_{k \in \{0, 1, \dots, p\}} \sup_{z \in \mathcal{Z}} |\hat{m}_k^{(l)}(z) - m_k(z)| \geq r_n \right\} \leq \rho_n.$$

Here,  $\hat{m}_k^{(l)}(z)$  and  $m_k(z)$  are the  $k$ -th components of  $\hat{\mathbf{m}}_{\mathbf{x}}^{(l)}(z)$  and  $\mathbf{m}_{\mathbf{x}}(z)$ , respectively, where  $k \in \{1, \dots, p\}$ . The sequence  $r_n$  satisfies  $(\|\beta_0\|_1 \vee 1)r_n = o(1)$  and  $(\|\beta_0\|_1 \vee 1)r_n^2 = o(\sqrt{\log p/n})$ , where  $\|\cdot\|_1$  is the  $\ell_1$  norm of a vector, and  $\mathcal{Z}$  is the support of  $Z_i$ .

(C2) We assume that  $\tilde{\mathbf{x}}_i = \mathbf{x}_i - \mathbf{m}_{\mathbf{x}}(Z_i)$  is sub-Gaussian with parameter  $\sigma$ . That is, for any unit vector  $\alpha \in \mathbb{R}^p$ , the random variable  $\alpha^T \tilde{\mathbf{x}}_i$  is sub-Gaussian with parameter at most  $\sigma$ .

(C3) The population covariance matrix  $\Sigma = E(\tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i^T)$  is positive definite.

(C4) There exist constants  $\delta_0 > 0$  and  $\kappa_l, \kappa_u > 0$  such that  $\inf_{t \in [-\delta_0, \delta_0]} f^*(t) \geq \kappa_l$  and  $\sup_{t \in \mathbb{R}} f^*(t) \leq \kappa_u$ . There exists some  $L_1 > 0$  such that  $|f^*(t_1) - f^*(t_2)| \leq L_1|t_1 - t_2|$  for any  $t_1, t_2 \in \mathbb{R}$ .

Condition (C1) plays a similar role to Assumption 5 in Zhu et al. (2019). It imposes requirements on the model parameters  $(n, p, q, \|\beta_0\|_1)$  and the uniform convergence rate  $r_n$  for the nonparametric estimators  $\hat{m}_k(\cdot)$ . Although we focus on univariate  $Z_i$  in this paper, our results apply to  $\mathbf{z}_i \in \mathbb{R}^d$  provided that Condition (C1) holds for  $r_n$ . For standard nonparametric regression estimators on a compact domain,  $r_n$  is typically of order  $O[\{\log(p \vee n)/n\}^r]$  under Condition (C2), after applying a union bound to the  $(p+1)$  estimators. For example, under certain regularity conditions (Masry, 1996; Li and Racine, 2007), local constant and local linear kernel estimators (Fan and Gijbels, 1996) can achieve a rate of  $r_n = O[\{\log(p \vee n)/n\}^{2/5}]$  when  $m_0(\cdot)$  and  $\mathbf{m}_{\mathbf{x}}(\cdot)$  are twice continuously differentiable and  $E(|\tilde{Y}_i|^s) < +\infty$  for some  $s > 2$ . Consequently, Condition (C1) is satisfied for a wide range of  $(n, p, q, \|\beta_0\|_1)$ . Notably, this moment condition on  $\tilde{Y}_i$  is much weaker than the sub-Gaussian assumption imposed in Zhu (2017) and Zhu et al. (2019). This condition is only required in the first-step for estimating the nonparametric functions and is not needed in the second-step for estimating  $\beta_0$ . The uniform bound required in Condition (C1) can be relaxed to an  $\ell_2$  bound, that is,  $\text{pr}(\max_{k \in \{0, 1, \dots, p\}} E[\{\hat{m}_k^{(1)}(Z_i^{(2)}) - m_k(Z_i^{(2)})\}^2 \mid \mathcal{D}_1] \geq r_n^2) \leq \rho_n$ . We assume the uniform bound to simplify the presentation and avoid technical complications.

Condition (C2) assumes that the design matrix is sub-Gaussian, a common assumption in high-dimensional regression (Fan et al., 2020a).

Condition (C3) is a standard assumption for high-dimensional design matrices and is relatively mild as it concerns the population covariance (van de Geer et al., 2014; Javanmard and Montanari, 2014; Zhang and Zhang, 2014; Ning and Liu, 2017; Zhang and Cheng, 2017). Moreover, we only require  $\Sigma$  to be positive definite, without imposing uniform boundedness on  $\lambda_{\min}^{-1}(\Sigma)$  or  $\|\Theta\|_1$ . Here,  $\Theta = \Sigma^{-1}$ , and  $\|\Theta\|_1 = \max_k \|\theta_k\|_1$  is the  $\ell_1$  norm of  $\Theta$ , where  $\theta_k$  is the  $k$ -th column of  $\Theta$ . Note that since  $\Sigma$  is symmetric, we have  $\|\Theta\|_1 \geq \lambda_{\min}^{-1}(\Sigma)$ . In our theoretical analysis, we allow both  $\|\Theta\|_1$  and  $\lambda_{\min}^{-1}(\Sigma)$  to diverge at a certain rate as  $n$  increases.

The assumptions on random errors in Condition (C4) are substantially weaker than those required in existing literature for high-dimensional partially linear models. For example, Zhu (2017) and Zhu et al. (2019) impose sub-Gaussianity on random errors, thereby excluding heavy-tailed distributions such as  $t$ -distributions. In general, Condition (C4) is not restrictive and can be satisfied by a wide range of density functions. Two examples are provided in Section A.6 of the Appendix.

To establish the theoretical properties of  $\hat{\beta}^{(l)}(\hat{\lambda}^{(l)})$  for  $l \in \{1, 2\}$ , we first introduce the following lemma:

**Lemma 1** *Under Conditions (C1), (C2) and (C4), and on the event  $\left\{ \max_{k \in \{0, 1, \dots, p\}} \sup_{z \in \mathcal{Z}} |\hat{m}_k^{(1)}(z) - m_k(z)| \leq r_n \right\}$ , we have the following results:*

- (i) *There exists some  $C_n > 0$  such that  $\text{pr} \left\{ \hat{\beta}^{(1)}(\hat{\lambda}^{(1)}) - \beta_0 \in \Gamma \mid \mathcal{D}_1 \right\} \geq 1 - \alpha_0 - 4p^{-(C_n^2 - 1)}$ , where  $C_n \rightarrow \infty$  as  $n \rightarrow \infty$ .*
- (ii) *For any  $C > 1$ , if  $\alpha_0 > 2p^{-(C^2 - 1)}$ , then  $\text{pr} \left( \hat{\lambda}^{(1)} \leq c_1 \sqrt{\log p/n} \mid \mathcal{D}_1 \right) \geq 1 - p \exp(-n)$ , where  $c_1 = c\{c_0 + 4\sqrt{2}C(2\sqrt{3}\sigma + r_n)\}$ .*

*These results also hold when the roles of  $\mathcal{D}_1$  and  $\mathcal{D}_2$  are switched.*

Part (i) of Lemma 1 indicates that  $\hat{\beta}^{(l)}(\hat{\lambda}^{(l)}) - \beta_0$  belongs to the cone set  $\Gamma$  with high probability. Part (ii) implies that  $\hat{\lambda}^{(l)}$  has an upper bound of  $c_1 \sqrt{\log p/n}$  with probability tending to one. Then, the following theorem establishes the  $\ell_1$  and  $\ell_2$  convergence rates of  $\hat{\beta}^{(l)}(\hat{\lambda}^{(l)})$ .

**Theorem 2** *Under Conditions (C1)–(C4), for any  $C > 1$ , if  $\alpha_0 \geq 2p^{-(C^2 - 1)}$ , then  $\hat{\beta}^{(l)}(\hat{\lambda}^{(l)})$  satisfies*

$$\begin{aligned} \|\hat{\beta}^{(l)}(\hat{\lambda}^{(l)}) - \beta_0\|_2 &\leq \delta \sqrt{q \log p/n}, \\ \|\hat{\beta}^{(l)}(\hat{\lambda}^{(l)}) - \beta_0\|_1 &\leq (1 + \bar{c}) \delta q \sqrt{\log p/n}, \end{aligned}$$

*with probability at least  $1 - 8p^{-(C^2 - 1)} - c_2 \exp(-c_3 n[\{\lambda_{\min}(\Sigma)/q\}^2/\sigma^4 \wedge 1] + \log p) - \rho_n - \alpha_0 - \alpha_n$  for some universal constants  $c_2, c_3 > 0$ , provided that  $2\{2(C\sigma L_n + r_n)(1 + \bar{c})q\delta \sqrt{\log p/n} + 2(\|\beta_0\|_1 \vee 1)r_n\} \leq \delta_0$  and  $16\kappa_u(\|\beta_0\|_1 \vee 1)r_n^2 \leq c_0 \sqrt{\log p/n}$ , where  $\delta = 2(1 + \bar{c})[c_1 + 32C\{4\sigma(\kappa_u \delta_0 \vee 1) + r_n\}]/\{\kappa_l \lambda_{\min}(\Sigma)\}$ ,  $L_n = \sqrt{\log(p \vee n)}$ ,  $\alpha_n = 4p^{-(C_n^2 - 1)}$ , with  $c_1$  and  $C_n$  defined in Lemma 1. Here,  $\|\cdot\|_2$  denotes the  $\ell_2$  norm.*

Theorem 2 demonstrates that under mild conditions on nonparametric estimators and error distributions, the DS rank Lasso estimator achieves near-oracle rates comparable to those of the rank Lasso estimator in linear models (Wang et al., 2020), as if the unknown functions  $m_0(\cdot)$  and  $\mathbf{m}_\mathbf{x}(\cdot)$  were known. Although Zhu (2017) obtained similar convergence rates for a two-step projection based Lasso estimator, their analysis requires the sub-Gaussian error assumption, which is substantially more restrictive.

**Remark 3** *Data-splitting allows the rank Lasso to achieve near-oracle rates under mild conditions on nonparametric estimators. Suppose we use the full sample to construct both the nonparametric estimators and the rank Lasso estimator. We define  $\mathbf{S}(\mathbf{0})$  as the full-sample counterpart of  $\mathbf{S}^{(l)}(\mathbf{0})$ . Following the tuning parameter selection framework in Wang et al. (2020), we choose  $\lambda$  such that  $\lambda \geq \|\mathbf{S}(\mathbf{0})\|_\infty$  with high probability. However, if overfitting occurs in the nonparametric estimators, the quantity  $\|\mathbf{S}(\mathbf{0})\|_\infty$  may no longer scale as  $O_p(\sqrt{\log p/N})$ . Consequently, without data splitting, the  $\ell_2$  convergence rate of the rank Lasso estimator could fail to achieve the near-oracle rate of  $O_p(\sqrt{q \log p/N})$ . A toy example illustrating this phenomenon is provided in Section A.7 of the Appendix.*

### 3. Robust Statistical Inference

In this section, we introduce a one-step update to the DS rank Lasso estimator for statistical inference in high-dimensional partially linear models. We define the debiased estimator as

$$\begin{aligned} \hat{\beta}^{d,(l)} &= \hat{\beta}^{(l)} + \hat{\Theta}^\top \{n(n-1)\}^{-1} \sum_{i \neq j} \left[ (\hat{\mathbf{x}}_i^{(l)} - \hat{\mathbf{x}}_j^{(l)}) \right. \\ &\quad \left. \times \text{sign} \left\{ (\hat{Y}_i^{(l)} - \hat{Y}_j^{(l)}) - (\hat{\mathbf{x}}_i^{(l)} - \hat{\mathbf{x}}_j^{(l)})^\top \hat{\beta}^{(l)} \right\} \right] / \{4\hat{f}^*(0)\} \\ &= \hat{\beta}^{(l)} - \hat{\Theta}^\top \mathbf{S}^{(l)}(\hat{\beta}^{(l)} - \beta_0) / \{4\hat{f}^*(0)\}, \end{aligned} \quad (4)$$

for each  $l \in \{1, 2\}$ , where  $\hat{\beta}^{(l)} = \hat{\beta}^{(l)}(\hat{\lambda}^{(l)})$ , and  $\hat{\Theta}$  and  $\hat{f}^*(0)$  are estimators of  $\Theta$  and  $f^*(0)$ , respectively.

**Estimator of  $\Theta$ .** For simplicity, we directly apply a modified version of the nodewise regression method proposed by van de Geer et al. (2014), as suggested in Zhu et al. (2019), to estimate  $\Theta$ . As shown in the proofs of Zhu et al. (2019, Theorems 1 and A.1), this method effectively controls the error bound  $\|\hat{\Theta} - \Theta\|_1 = \max_k \|\hat{\theta}_k - \theta_k\|_1$  under mild conditions, where  $\hat{\theta}_k$  and  $\theta_k$  are the  $k$ -th columns of  $\hat{\Theta}$  and  $\Theta$ , respectively. However, it is important to note that any estimator capable of providing reasonable control over  $\|\hat{\Theta} - \Theta\|_1$  can be employed for this task.

**Estimator of  $f^*(0)$ .** Based on the observation that  $f^*(0) = \int_{-\infty}^{\infty} f^2(t) dt = E\{f(\epsilon_i)\}$ , we estimate  $f^*(0)$  by

$$\hat{f}^{*,(l)}(0) = n^{-1} \sum_{i=1}^n \hat{f}_{-i}^{(l)}(\hat{\epsilon}_i^{(l)}), \quad (5)$$

where  $\hat{f}_{-i}^{(l)}(\hat{\epsilon}_i^{(l)}) = (n-1)^{-1} \sum_{j=1, j \neq i}^n K((\hat{\epsilon}_i^{(l)} - \hat{\epsilon}_j^{(l)})/h)/h$  and  $\hat{\epsilon}_i^{(l)} = \hat{Y}_i^{(l)} - (\hat{\mathbf{x}}_i^{(l)})^\top \hat{\beta}^{(l)}$ . Here,  $K(\cdot)$  is a kernel function, and  $h > 0$  is the bandwidth. The final estimator for  $f^*(0)$  is defined as  $\hat{f}^*(0) = (\hat{f}^{*,(1)}(0) + \hat{f}^{*,(2)}(0))/2$ . For this estimator, we impose the following assumption:

(C5) The kernel function  $K : \mathbb{R} \rightarrow [0, \infty)$  satisfies (i)  $K(t) = K(-t)$  for any  $t \in \mathbb{R}$ , (ii)  $\int_{-\infty}^{\infty} K(t) dt = 1$ , (iii)  $\int_{-\infty}^{\infty} t^2 K(t) dt < \infty$ , (iv)  $\int_{-\infty}^{\infty} K^2(t) dt < \infty$ , and (v) there exists some  $L_2 > 0$  such that  $|K(t_1) - K(t_2)| \leq L_2 |t_1 - t_2|$  for any  $t_1, t_2 \in \mathbb{R}$ . Furthermore, there exists some  $L_3 > 0$  such that the first derivative of  $f^*$ , denoted  $(f^*)'$ , satisfies  $|(f^*)'(t_1) - (f^*)'(t_2)| \leq L_3 |t_1 - t_2|$  for any  $t_1, t_2 \in \mathbb{R}$ .

Condition (C5) is a standard assumption in kernel density estimation.

The rationale for the proposed debiased estimators stems from the following linear approximation

$$\sqrt{n} \hat{\Theta}^T \mathbf{S}^{(l)} (\hat{\beta}^{(l)} - \beta_0) / \{4\hat{f}^*(0)\} \approx \sqrt{n} \Theta^T \mathbf{S}^{(l)} (\mathbf{0}) / \{4f^*(0)\} + \sqrt{n} (\hat{\beta}^{(l)} - \beta_0).$$

While Hettmansperger and McKean (2010) and Fan et al. (2020b) derived similar expressions for fixed-dimensional linear models, we extend their results to partially linear models and allow the dimensionality  $p$  to diverge. Let  $F(\cdot)$  denote the distribution function of  $\epsilon_i$ . It can further be shown that the two terms  $-\sqrt{n} \theta_k^T \mathbf{S}^{(l)} (\mathbf{0}) / \{4f^*(0)\}$  and  $n^{-1/2} \sum_{i=1}^n \theta_k^T \tilde{\mathbf{x}}_i^{(l)} \{F(\epsilon_i^{(l)}) - 1/2\} / f^*(0)$  are asymptotically equivalent. The following lemma formally states these results.

**Lemma 4** *Under Conditions (C1)–(C5),  $\|\Theta\|_1 L_n q \sqrt{\log p/n} = o(1)$ ,  $\|\hat{\Theta} - \Theta\|_1 = o_p(1)$ ,  $\|\hat{\beta}^{(l)}\|_0 = O_p(q)$ , and  $h \asymp \{\|\Theta\|_1 q \log(p)/n\}^{1/4} \vee n^{-1/5}$ ,  $\hat{\beta}^{d,(l)}$  satisfies*

$$\hat{\beta}^{d,(l)} - \beta_0 = -\Theta^T \mathbf{S}^{(l)} (\mathbf{0}) / \{4f^*(0)\} + \mathbf{R}_{n,1}^{(l)},$$

with

$$\begin{aligned} \|\mathbf{R}_{n,1}^{(l)}\|_{\infty} &= O_p \left[ \|\Theta\|_1 \sqrt{q \log p/n} \left\{ \sqrt{q} \|\hat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 L_n \sqrt{q \log p/n} \right. \right. \\ &\quad \left. \left. \vee \|\Theta\|_1 (\|\beta_0\|_1 \vee 1) r_n \right\} \vee \|\Theta\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{1/4} (q/n)^{3/4} \right]. \end{aligned}$$

Furthermore, we have

$$-\Theta^T \mathbf{S}^{(l)} (\mathbf{0}) / \{4f^*(0)\} = n^{-1} \sum_{i=1}^n \Theta^T \tilde{\mathbf{x}}_i^{(l)} \{F(\epsilon_i^{(l)}) - 1/2\} / f^*(0) + \mathbf{R}_{n,2}^{(l)},$$

with

$$\|\mathbf{R}_{n,2}^{(l)}\|_{\infty} = O_p \left[ \|\Theta\|_1 \left\{ (\|\beta_0\|_1 \vee 1) r_n^2 \vee \sqrt{(\|\beta_0\|_1 \vee 1) r_n \log p/n} \vee L_n \log p/n \right\} \right].$$

If we were to use only one of the estimators,  $\hat{\beta}^{d,(1)}$  or  $\hat{\beta}^{d,(2)}$ , for statistical inference, this would result in a loss of efficiency, since each estimator asymptotically uses only half of the data. However, full efficiency can be recovered by averaging  $\hat{\beta}^{d,(1)}$  and  $\hat{\beta}^{d,(2)}$ . The intuition is that, by Lemma 4,  $\hat{\beta}^{d,(l)} - \beta_0$  for  $l \in \{1, 2\}$  decomposes into a leading term,  $n^{-1} \sum_{i=1}^n \Theta^T \tilde{\mathbf{x}}_i^{(l)} \{F(\epsilon_i^{(l)}) - 1/2\} / f^*(0)$ , plus a remainder term that is asymptotically negligible. The two leading terms are independent and have mean zero, so  $\hat{\beta}^{d,(1)}$  and  $\hat{\beta}^{d,(2)}$  are asymptotically independent and unbiased. Therefore, we define the final debiased DS rank Lasso estimator as

$$\hat{\beta}^d = \frac{\hat{\beta}^{d,(1)} + \hat{\beta}^{d,(2)}}{2}, \quad (6)$$

and in the following section, we show that this debiased estimator achieves the same asymptotic variance as a full-sample oracle estimator. This means that data-splitting does not lead to any asymptotic efficiency loss in our setting.

### 3.1 Inference for Individual Coefficients

In Theorem 5, we present the theoretical properties of  $\hat{\beta}_k^d$ , the  $k$ -th component of the debiased estimator  $\hat{\beta}^d$ .

**Theorem 5** *Under the conditions of Lemma 4, if in addition  $\|\Theta\|_1^3 L_n^3 \{\log(p)\}^{1/2} q^{3/2} n^{-1/2} = o(1)$ ,  $\|\Theta\|_1^2 (\|\beta_0\|_1 \vee 1) r_n (\sqrt{q \log p} \vee \log p) = o(1)$ ,  $\sqrt{n} \|\Theta\|_1 (\|\beta_0\|_1 \vee 1) r_n^2 = o(1)$ , and  $\|\Theta\|_1 q \sqrt{\log p} \|\hat{\Theta} - \Theta\|_1 = o_p(1)$ , the debiased DS rank Lasso estimator  $\hat{\beta}_k^d$  satisfies*

$$\sqrt{N}(\hat{\beta}_k^d - \beta_{0,k})/\hat{\omega}_k \xrightarrow{d} \mathcal{N}(0, 1),$$

for each  $k \in \{1, \dots, p\}$ , where

$$\hat{\omega}_k^2 = \hat{\theta}_k^T \left( \frac{\hat{\Sigma}^{(1)} + \hat{\Sigma}^{(2)}}{2} \right) \hat{\theta}_k / [12 \{\hat{f}^*(0)\}^2],$$

with  $\hat{\Sigma}^{(l)} = n^{-1} \sum_{i=1}^n (\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)})(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)})^T$  and  $\bar{\mathbf{x}}^{(l)} = n^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(l)}$ .

**Remark 6** *Without data-splitting, some remainder terms in  $\sqrt{N}(\hat{\beta}_k^d - \beta_{0,k})$  may not vanish if overfitting occurs in the nonparametric estimators. This can lead to poor performance of  $\hat{\beta}_k^d$ . A toy example, as used in Remark 3, is provided in Appendix A.7 to illustrate this issue.*

Based on Theorem 5, we construct the  $(1 - \alpha)$ -confidence interval for  $\beta_{0,k}$  as follows:

$$\left[ \hat{\beta}_k^d - \Phi^{-1}(1 - \alpha/2) \hat{\omega}_k / \sqrt{N}, \quad \hat{\beta}_k^d + \Phi^{-1}(1 - \alpha/2) \hat{\omega}_k / \sqrt{N} \right],$$

where  $\Phi^{-1}(1 - \alpha/2)$  denotes the  $(1 - \alpha/2)$ -quantile of the standard normal distribution.

Theorem 5 implies that under mild conditions, the asymptotic distribution of  $\hat{\beta}_k^d$  behaves as if the unknown functions  $m_0(\cdot)$  and  $\mathbf{m}_{\mathbf{x}}(\cdot)$  were known. To elaborate, if  $m_0(\cdot)$  and  $\mathbf{m}_{\mathbf{x}}(\cdot)$  were known, we could define the debiased full-sample (FS) oracle rank Lasso estimator as:

$$\tilde{\beta}^d = \tilde{\beta} + \tilde{\Theta}^T \{N(N-1)\}^{-1} \sum_{i \neq j} (\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j) \text{sign}\{(\tilde{Y}_i - \tilde{Y}_j) - (\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)^T \tilde{\beta}\} / \{4 \tilde{f}^*(0)\},$$

where  $\tilde{\beta}$  and  $\tilde{\Theta}$  are, respectively, the ordinary rank Lasso estimator (Wang et al., 2020) and the nodewise Lasso estimator (van de Geer et al., 2014) for  $\beta_0$  and  $\Theta$ , based on  $\{\tilde{\mathbf{x}}_i, \tilde{Y}_i, i = 1, \dots, N\}$ , and  $\tilde{f}^*(0) = \{N(N-1)\}^{-1} \sum_{i \neq j}^N K[\{(\tilde{Y}_i - \tilde{Y}_j) - (\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)^T \tilde{\beta}\}/h]/h$ . Following arguments similar to those in the proofs of Lemma 4 and Theorem 5, it can be shown that

$$\sqrt{N}(\tilde{\beta}_k^d - \beta_{0,k})/\tilde{\omega}_k \xrightarrow{d} \mathcal{N}(0, 1),$$

where  $\tilde{\omega}_k^2 = \tilde{\boldsymbol{\theta}}_k^T \tilde{\boldsymbol{\Sigma}} \tilde{\boldsymbol{\theta}}_k / [12\{\tilde{f}^*(0)\}^2]$ ,  $\tilde{\boldsymbol{\theta}}_k$  is the  $k$ -th column of  $\tilde{\boldsymbol{\Theta}}$ ,  $\tilde{\boldsymbol{\Sigma}} = N^{-1} \sum_{i=1}^N (\tilde{\mathbf{x}}_i - \tilde{\bar{\mathbf{x}}})(\tilde{\mathbf{x}}_i - \tilde{\bar{\mathbf{x}}})^T$  and  $\tilde{\bar{\mathbf{x}}} = N^{-1} \sum_{i=1}^N \tilde{\mathbf{x}}_i$ . Both  $\tilde{\omega}_k^2/\omega_k^2$  and  $\tilde{\omega}_k^2/\omega_k^2$  converge in probability to 1, where  $\omega_k^2 = \Theta_{k,k}/[12\{f^*(0)\}^2]$ . This implies that the proposed debiased DS rank Lasso estimator achieves the same asymptotic variance as the debiased FS oracle rank Lasso estimator. The intuition is that, by Lemma 4,  $\sqrt{N}(\tilde{\boldsymbol{\beta}}^d - \boldsymbol{\beta}_0)$  decomposes into a leading term,  $N^{-1/2} \sum_{i=1}^N \boldsymbol{\Theta}^T \tilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\}/f^*(0)$ , plus remainder terms; similarly,  $\sqrt{N}(\hat{\boldsymbol{\beta}}^d - \boldsymbol{\beta}_0)$  decomposes into the same leading term plus its own remainder terms. These remainder terms can be shown to vanish asymptotically under the conditions of Theorem 5. In Section B.1 of the Appendix, we provide simulation studies comparing inference based on the debiased DS rank Lasso estimator with that based on the debiased FS oracle rank Lasso estimator. The results demonstrate that data splitting introduces only minor efficiency loss in finite samples when the sample size is moderately large.

As shown in Zhu et al. (2019, Theorem 1), the debiased two-step projection-based Lasso estimator  $\hat{\boldsymbol{\beta}}_{\text{Lasso}}^d$  satisfies

$$\sqrt{n}(\hat{\boldsymbol{\beta}}_{k,\text{Lasso}}^d - \beta_{0,k})/\hat{\omega}_{k,\text{Lasso}} \xrightarrow{d} \mathcal{N}(0, 1),$$

under the sub-Gaussian assumption for  $\epsilon_i$ , where  $\hat{\omega}_{k,\text{Lasso}}^2/\omega_{k,\text{Lasso}}^2 \xrightarrow{p} 1$  for  $\omega_{k,\text{Lasso}}^2 = \Theta_{k,k}E(\epsilon_i^2)$ . Consequently, the asymptotic relative efficiency (ARE) of our estimator  $\hat{\boldsymbol{\beta}}_k^d$  relative to  $\hat{\boldsymbol{\beta}}_{k,\text{Lasso}}^d$  is  $12\{f^*(0)\}^2 E(\epsilon_i^2)$ . This ARE can be substantially greater than one for heavy-tailed error distributions (Hettmansperger and McKean, 2010; Wang et al., 2020). For example, when  $\epsilon_i$  follows a mixture-normal distribution with density  $f_\epsilon(x) = (1 - \varepsilon)\phi(x) + \varepsilon\phi(x/3)/3$  and  $\varepsilon$  varies in  $\{0.00, 0.01, 0.03, 0.05, 0.10, 0.15\}$ , where  $\phi(\cdot)$  is the density of the standard normal distribution, the corresponding ARE values are  $\{0.955, 1.009, 1.108, 1.196, 1.373, 1.497\}$  (Hettmansperger and McKean, 2010, Table 1.7.1). This indicates that the proposed debiased DS rank Lasso estimator offers greater robustness against heavy-tailed errors compared to the debiased two-step projection-based Lasso estimator.

It is worth mentioning that the initial DS rank Lasso estimator can be replaced by its ANOPE subsampled version, as suggested by Fan et al. (2020b). Specifically, the loss function  $Q^{(l)}(\boldsymbol{\gamma})$  can be replaced by

$$Q^{(l),\text{sub}}(\boldsymbol{\gamma}) = (M_n|\mathcal{S}|)^{-1} \sum_{\pi \in \mathcal{S}} \sum_{i=1}^{M_n} |(e_{\pi(i)}^{(2)} - e_{\pi(M_n+i)}^{(2)}) - (\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^T \boldsymbol{\gamma}|,$$

where  $\mathcal{S}$  is sampled with replacement from permutations of  $\{1, \dots, n\}$ ,  $|\mathcal{S}|$  denotes its cardinality, and  $M_n = \lfloor n/2 \rfloor$  is the largest integer smaller than or equal to  $n/2$ . This subsampling method can further reduce computational complexity. Our simulation studies in Section 4.1 demonstrate that ANOPE only results in a negligible loss of statistical efficiency.

In this paper, we randomly split the data into two equal halves to construct a two-split debiased estimator. It is also possible to partition the data into  $K \geq 3$  equally sized subsamples. For  $l \in \{1, \dots, K\}$ , we compute the nonparametric estimators using all subsamples except the  $l$ -th one and then obtain the corresponding debiased estimator similar to (4). The final debiased estimator is then obtained by averaging the  $K$  debiased estimators. Similar to the two-split case, the error of the  $K$ -split debiased estimator can be decomposed into an  $O_p(N^{-1/2})$  asymptotically normal term and a remainder term. For

fixed  $K$ , the remainder term is asymptotically negligible, and the  $K$ -split debiased estimator shares the same asymptotic normality as  $\tilde{\beta}^d$ . Although theoretically valid, this approach is less computationally attractive, as it requires performing  $K$  nonparametric regressions without improving asymptotic efficiency compared to the two-split estimator.

Alternatively, one may consider unbalanced splitting instead of balanced splitting. In such cases, the final debiased estimator is a weighted average of the subsample-specific estimators, with weights proportional to their respective sample sizes. Theoretical analysis for unbalanced splitting follows similar reasoning as in the balanced case, and the resulting estimator exhibits analogous asymptotic properties. Since unbalanced splitting likewise does not improve asymptotic efficiency, we adopt the two-split debiased estimator with balanced splitting in this article for simplicity.

### 3.2 Simultaneous Inference

In high-dimensional settings, researchers may be interested in conducting a simultaneous hypothesis test for a set of coefficients indexed by  $\mathcal{G} \subseteq \{1, \dots, p\}$ :

$$H_{0,\mathcal{G}} : \beta_{0,k} = \beta_k^*, \text{ for all } k \in \mathcal{G},$$

versus

$$H_{a,\mathcal{G}} : \beta_{0,k} \neq \beta_k^*, \text{ for some } k \in \mathcal{G},$$

where  $\{\beta_k^*\}_{k \in \mathcal{G}}$  are prespecified values. We allow the cardinality of  $\mathcal{G}$ , denoted by  $|\mathcal{G}|$ , to grow as fast as  $p$ , which itself may grow exponentially with the sample size. To test  $H_{0,\mathcal{G}}$ , we use two test statistics: the non-studentized statistic  $T_{\mathcal{G}} = \max_{k \in \mathcal{G}} \sqrt{N} |\hat{\beta}_k^d - \beta_k^*|$  and the studentized statistic  $\bar{T}_{\mathcal{G}} = \max_{k \in \mathcal{G}} \sqrt{N} |\hat{\beta}_k^d - \beta_k^*| / \hat{\omega}_k$ , where  $\hat{\omega}_k^2$  is defined in Theorem 5. Inspired by Zhang and Cheng (2017) and Zhu et al. (2019), and based on the asymptotic expansions in Lemma 4, we propose two multiplier bootstrap statistics:

$$\begin{aligned} W_{\mathcal{G}} &= \max_{k \in \mathcal{G}} \left| \sum_{l=1}^2 \sum_{i=1}^n \hat{\theta}_k^T(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)}) u_i^{(l)} \right| / [\sqrt{12N} \{\hat{f}^*(0)\}], \\ \bar{W}_{\mathcal{G}} &= \max_{k \in \mathcal{G}} \left| \sum_{l=1}^2 \sum_{i=1}^n \hat{\theta}_k^T(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)}) u_i^{(l)} \right| / [\sqrt{12N} \{\hat{f}^*(0)\} \hat{\omega}_k], \end{aligned}$$

where  $\bar{\mathbf{x}}^{(l)} = n^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(l)}$ , and  $u_1^{(1)}, \dots, u_n^{(1)}, u_1^{(2)}, \dots, u_n^{(2)}$  are independent standard normal random variables, independent of  $\{\mathbf{x}_i, Y_i, Z_i\}_{i=1}^N$ . Then, the corresponding bootstrap critical values are given by  $c_{\mathcal{G}}(\alpha) = \inf\{t \in \mathbb{R} : \text{pr}[W_{\mathcal{G}} \leq t \mid \{\mathbf{x}_i, Y_i, Z_i\}_{i=1}^N] \geq 1 - \alpha\}$  and  $\bar{c}_{\mathcal{G}}(\alpha) = \inf\{t \in \mathbb{R} : \text{pr}[\bar{W}_{\mathcal{G}} \leq t \mid \{\mathbf{x}_i, Y_i, Z_i\}_{i=1}^N] \geq 1 - \alpha\}$  for any  $\alpha \in (0, 1)$ . We reject  $H_{0,\mathcal{G}}$  if  $T_{\mathcal{G}} > c_{\mathcal{G}}(\alpha)$ , or  $\bar{T}_{\mathcal{G}} > \bar{c}_{\mathcal{G}}(\alpha)$  at significance level  $\alpha$ . Furthermore, we can construct two simultaneous confidence intervals for  $\{\beta_{0,k}\}_{k \in \mathcal{G}}$  as  $[\hat{\beta}_k^d - c_{\mathcal{G}}(\alpha)/\sqrt{N}, \hat{\beta}_k^d + c_{\mathcal{G}}(\alpha)/\sqrt{N}]$ , and  $[\hat{\beta}_k^d - \hat{\omega}_k \bar{c}_{\mathcal{G}}(\alpha)/\sqrt{N}, \hat{\beta}_k^d + \hat{\omega}_k \bar{c}_{\mathcal{G}}(\alpha)/\sqrt{N}]$ . Note that the first simultaneous confidence intervals have equal lengths across all  $k \in \mathcal{G}$ , while the second can have varying lengths.

The validity of these bootstrap-based methods is established by the following theorem.

**Theorem 7** *Under the conditions of Lemma 4, if further  $\|\Theta\|_1^3 L_n^3 \{\log(p)\}^{3/2} q^{3/2} n^{-1/2} = o(1)$ ,  $\|\Theta\|_1^2 (\|\beta_0\|_1 \vee 1) r_n [\sqrt{q} \log p \vee \{\log(p)\}^2] = o(1)$ ,  $\|\Theta\|_1 \sqrt{n \log p} (\|\beta_0\|_1 \vee 1) r_n^2 = o(1)$ ,  $\|\Theta\|_1^6 L_n^{14}/n \leq M n^{-\iota}$  for some  $M, \iota > 0$ , and there exists some  $\zeta_n \rightarrow +\infty$  such that  $\zeta_n [\|\Theta\|_1 q \log p \vee \{\log(p)\}^2] \|\widehat{\Theta} - \Theta\|_1 = o_p(1)$ , we have*

$$\sup_{\alpha \in (0,1)} \left| \Pr \left\{ \max_{k \in \mathcal{G}} \sqrt{N} |\widehat{\beta}_k^d - \beta_{0,k}| / \widehat{\omega}_k > \bar{c}_{\mathcal{G}}(\alpha) \right\} - \alpha \right| = o(1).$$

*If in addition  $\max_{k \in \{1, \dots, p\}} \Theta_{k,k} < +\infty$ , we have*

$$\sup_{\alpha \in (0,1)} \left| \Pr \left\{ \max_{k \in \mathcal{G}} \sqrt{N} |\widehat{\beta}_k^d - \beta_{0,k}| > c_{\mathcal{G}}(\alpha) \right\} - \alpha \right| = o(1).$$

Theorem 7 shows that the coverage probabilities of the two types of simultaneous confidence intervals converge to the nominal level  $1 - \alpha$ . It also implies that under  $H_{0,\mathcal{G}}$ ,  $\sup_{\alpha \in (0,1)} |\Pr\{T_{\mathcal{G}} > c_{\mathcal{G}}(\alpha)\} - \alpha| = o(1)$  and  $\sup_{\alpha \in (0,1)} |\Pr\{\bar{T}_{\mathcal{G}} > \bar{c}_{\mathcal{G}}(\alpha)\} - \alpha| = o(1)$ , indicating that the bootstrap-based tests have correct size under the null hypothesis. Furthermore, Theorem 7 enables power analysis of the two tests, analogous to Zhang and Cheng (2017, Theorem 2.3 and Theorem 2.4).

## 4. Numerical Studies

This section presents numerical studies to evaluate the performance of the proposed methods. Specifically, we demonstrate the effectiveness of the proposed estimation and inference methods through simulation studies in Section 4.1. In Section 4.2, we examine the performance of the proposed methods in solving real-world problems.

### 4.1 Simulation Studies

In the simulation study, covariates  $\mathbf{x}_i$  ( $i = 1, \dots, N$ ) are generated as follows. First,  $\mathbf{x}_{0,i}$  is sampled from a  $p$ -dimensional normal distribution with mean  $\mathbf{0}$  and covariance matrix  $\Sigma = (\Sigma_{k,k'})_{p \times p}$ , where  $\Sigma_{k,k'} = 0.5^{|k-k'|}$ . Then,  $Z_i$  is generated from a uniform distribution on the interval  $[0, 2]$ . Finally, the covariates  $X_{i,k}$  are defined as:  $X_{i,1} = X_{0,i,1} + 3Z_i$ ,  $X_{i,2} = X_{0,i,2} + 3Z_i^2$ ,  $X_{i,3} = X_{0,i,3} - 3Z_i$ , and  $X_{i,k} = X_{0,i,k}$  for  $4 \leq k \leq p$ , for all  $1 \leq i \leq N$ .

Following Example 1 in Wang et al. (2020), we define the active set as  $\mathcal{B} = \{1, 2, 3\}$  and set the regression coefficient  $\beta_0$  to be  $(\sqrt{3}, \sqrt{3}, \sqrt{3}, 0, \dots, 0)^T$ . For the nonlinear nuisance function, we use  $g_0(z) = 1.5 \sin(2\pi z)$  in this section. The sensitivity of the proposed inference methods to different forms of  $g_0(z)$  is further examined in Section B.3 of the Appendix. To evaluate the performance of the proposed methods, we consider the following error distributions for  $\epsilon_i$ :

- (i) Normal distribution:  $\epsilon_i \sim \mathcal{N}(0, 2^2)$ ;
- (ii) Standard log-normal (LogN) distribution:  $\log(\epsilon_i) \sim \mathcal{N}(0, 1)$ ;
- (iii) Mixture-Normal (MN) distribution:  $\epsilon_i \sim 0.95\mathcal{N}(0, 1) + 0.05\mathcal{N}(0, 10^2)$ ;
- (iv) Scaled  $t$ -distribution:  $\epsilon_i/\sqrt{2} \sim t(2)$ , where  $t(2)$  denotes the  $t$ -distribution with 2 degrees of freedom;

(v) Standard Cauchy distribution:  $\epsilon_i \sim \text{Cauchy}(0, 1)$ .

We then simulate  $Y_i$  based on model (1). We set the sample size  $N \in \{300, 450\}$  and the dimensionality  $p = 500$ . We consider both the DS rank Lasso estimator (denoted as DSR-Lasso) and its ANOPE (Fan et al., 2020b) subsampled version (denoted as DSRLasso<sup>sub</sup>). To estimate  $m_0(\cdot)$ , we use local linear regression (Fan and Gijbels, 1996). For  $\mathbf{m}_\mathbf{x}(\cdot)$ , where  $p$  estimators are required, we apply local constant regression due to its faster computational speed. The sensitivity of the proposed methods to the choice of nonparametric estimators is investigated in Section B.4 of the Appendix.

In Example 1, we compare the proposed DSRLasso estimator with the two-step projection based Lasso estimator (denoted as TSLasso; Zhu, 2017). In Examples 2 and 3, we further compare the performance of two debiased procedures: the proposed debiased DSR-Lasso and the debiased TSLasso (Zhu et al., 2019), with Example 2 focusing on inference for individual coefficients and Example 3 on simultaneous inference. For the TSLasso-based methods, we follow Zhu et al. (2019) and use smoothing splines for the first-step nonparametric projection. In these examples, we observe that although the moment condition on  $\tilde{Y}_i$  required by local regression fails to hold under  $t(2)$  and Cauchy errors, the proposed DSRLasso and debiased DSRLasso estimators still exhibit reasonable performance. This contrasts sharply with the degraded performance of the TSLasso and debiased TSLasso estimators. This discrepancy arises because the moment condition is necessary only for the first-step local regression, whereas the second-step penalized rank regression does not require it. In contrast, the sub-Gaussian assumption on  $\tilde{Y}_i$  is crucial for both the first-step regularized nonparametric least squares and the second-step penalized least squares in the TSLasso methods (Zhu, 2017; Zhu et al., 2019).

For high-dimensional linear models, aside from the debiased rank Lasso estimator (Fan et al., 2020b), another robust inference method that does not rely on sub-Gaussian or moment conditions for random errors is the one based on the debiased Lasso-penalized quantile, particularly median, estimator (Bradic and Kolar, 2017). However, to the best of our knowledge, median-based inference methods have not been developed for high-dimensional partially linear models. Moreover, median-based methods have been shown to incur significant efficiency loss for light-tailed errors (e.g., normal errors) compared to rank-based methods (Hettmansperger and McKean, 2010; Wang et al., 2020; Cai et al., 2025). In Section B.5 of the Appendix, we show this efficiency loss persists when the debiased Lasso-penalized median estimator is extended to high-dimensional partially linear models.

Following the recommendations of Wang et al. (2020) and Fan et al. (2020b), we set the constants  $c = 1.01$  and  $\alpha_0 = 0.10$  in (3). Additionally, we set the constant  $c_0 = 0.01$ . These choices lead to satisfactory performance in our practice. Furthermore, a sensitivity analysis for these constants is conducted in Section B.2 of the Appendix.

**Example 1.** The comparison between the proposed DSRLasso estimator and the TSLasso estimator is conducted using the following four metrics: (i)  $\ell_1$  error  $\|\hat{\beta} - \beta_0\|_1$ ; (ii)  $\ell_2$  error  $\|\hat{\beta} - \beta_0\|_2$ ; (iii) number of false positive variables (FP)  $\sum_{k \in \mathcal{B}^c} I(\hat{\beta}_k \neq 0)$ ; and (iv) number of false negative variables (FN)  $\sum_{k \in \mathcal{B}} I(\hat{\beta}_k = 0)$ . The results are summarized in Table 1, which presents the averages and standard errors of these metrics based on 200 simulations.

| Error                 | $N$ | Method                  | $\ell_1$ Error | $\ell_2$ Error | FP           | FN           |
|-----------------------|-----|-------------------------|----------------|----------------|--------------|--------------|
| $\mathcal{N}(0, 2^2)$ | 300 | TSLasso                 | 1.141(0.050)   | 0.450(0.008)   | 10.77(0.871) | 0.000(0.000) |
|                       |     | DSRLasso                | 1.212(0.015)   | 0.776(0.010)   | 0.225(0.035) | 0.000(0.000) |
|                       |     | DSRLasso <sup>sub</sup> | 1.225(0.016)   | 0.783(0.011)   | 0.360(0.043) | 0.000(0.000) |
|                       | 450 | TSLasso                 | 0.926(0.040)   | 0.351(0.006)   | 11.46(0.828) | 0.000(0.000) |
|                       |     | DSRLasso                | 0.976(0.012)   | 0.632(0.008)   | 0.200(0.032) | 0.000(0.000) |
|                       |     | DSRLasso <sup>sub</sup> | 0.986(0.013)   | 0.637(0.008)   | 0.400(0.047) | 0.000(0.000) |
| LogN                  | 300 | TSLasso                 | 1.195(0.057)   | 0.464(0.011)   | 11.15(0.775) | 0.000(0.000) |
|                       |     | DSRLasso                | 0.789(0.012)   | 0.507(0.008)   | 0.250(0.037) | 0.000(0.000) |
|                       |     | DSRLasso <sup>sub</sup> | 0.799(0.013)   | 0.512(0.008)   | 0.450(0.047) | 0.000(0.000) |
|                       | 450 | TSLasso                 | 0.928(0.037)   | 0.372(0.008)   | 10.34(0.756) | 0.000(0.000) |
|                       |     | DSRLasso                | 0.610(0.009)   | 0.394(0.006)   | 0.260(0.034) | 0.000(0.000) |
|                       |     | DSRLasso <sup>sub</sup> | 0.615(0.010)   | 0.397(0.006)   | 0.480(0.045) | 0.000(0.000) |
| MN                    | 300 | TSLasso                 | 1.306(0.057)   | 0.530(0.012)   | 10.32(0.783) | 0.000(0.000) |
|                       |     | DSRLasso                | 0.840(0.013)   | 0.539(0.009)   | 0.260(0.036) | 0.000(0.000) |
|                       |     | DSRLasso <sup>sub</sup> | 0.845(0.014)   | 0.540(0.009)   | 0.420(0.047) | 0.000(0.000) |
|                       | 450 | TSLasso                 | 1.137(0.044)   | 0.438(0.009)   | 11.97(0.777) | 0.000(0.000) |
|                       |     | DSRLasso                | 0.636(0.009)   | 0.412(0.006)   | 0.200(0.029) | 0.000(0.000) |
|                       |     | DSRLasso <sup>sub</sup> | 0.641(0.009)   | 0.413(0.006)   | 0.360(0.043) | 0.000(0.000) |
| $t(2)$                | 300 | TSLasso                 | 1.998(0.080)   | 0.854(0.029)   | 9.450(0.664) | 0.025(0.013) |
|                       |     | DSRLasso                | 1.401(0.020)   | 0.886(0.013)   | 0.280(0.038) | 0.000(0.000) |
|                       |     | DSRLasso <sup>sub</sup> | 1.415(0.020)   | 0.892(0.013)   | 0.485(0.050) | 0.000(0.000) |
|                       | 450 | TSLasso                 | 1.710(0.070)   | 0.692(0.025)   | 10.76(0.726) | 0.030(0.016) |
|                       |     | DSRLasso                | 1.051(0.016)   | 0.666(0.010)   | 0.275(0.042) | 0.000(0.000) |
|                       |     | DSRLasso <sup>sub</sup> | 1.061(0.016)   | 0.670(0.010)   | 0.475(0.050) | 0.000(0.000) |
| Cauchy                | 300 | TSLasso                 | 5.331(0.181)   | 2.542(0.049)   | 4.995(0.551) | 1.685(0.090) |
|                       |     | DSRLasso                | 2.397(0.062)   | 1.451(0.033)   | 0.545(0.058) | 0.055(0.025) |
|                       |     | DSRLasso <sup>sub</sup> | 2.433(0.064)   | 1.462(0.033)   | 0.815(0.072) | 0.045(0.023) |
|                       | 450 | TSLasso                 | 5.185(0.130)   | 2.515(0.047)   | 4.910(0.506) | 1.680(0.088) |
|                       |     | DSRLasso                | 1.907(0.066)   | 1.170(0.035)   | 0.525(0.054) | 0.010(0.007) |
|                       |     | DSRLasso <sup>sub</sup> | 1.943(0.078)   | 1.186(0.046)   | 0.820(0.067) | 0.010(0.007) |

Table 1: Simulation Results for Example 1. Averages (and Standard Errors) of  $\ell_1$  Error,  $\ell_2$  Error, # of False Positive Variables (FP) and # of False Negative Variables (FN) Based on 200 Simulations.

We have the following observations. (1) DSRLasso yields considerably fewer false positives than TSLasso across all error distributions. (2) When random errors follow the normal distribution, DSRLasso produces comparable  $\ell_1$  errors to TSLasso, while under heavy-tailed distributions, such as mixture-normal, log-normal, and  $t(2)$ , DSRLasso yields significantly smaller  $\ell_1$  errors. (3) In the case of Cauchy-distributed errors, DSRLasso outperforms TSLasso across all four evaluation metrics. Notably, TSLasso fails to control false negatives: on average, it incorrectly shrinks more than half of the coefficients in the active set to zero. This leads to substantially larger  $\ell_1$  and  $\ell_2$  errors compared to DSRLasso. For example, when  $N = 450$ , the average  $\ell_1$  error,  $\ell_2$  error and FN for DSRLasso are 1.907, 1.170 and 0.010, respectively, while the corresponding values for TSLasso are 5.185, 2.515 and 1.680. (4) DSRLasso<sup>sub</sup> incurs only a negligible loss of statistical efficiency compared to DSRLasso. (5) Overall, DSRLasso demonstrates greater robustness to heavy-tailed random errors, which aligns with its weaker assumptions on error distributions and is consistent with earlier findings on the robustness of rank Lasso (Wang et al., 2020).

**Example 2.** We compare the performance of the proposed debiased DSRLasso and the debiased TSLasso in inferring individual coefficients, using the following metrics:

- (i) Coverage probability for  $\beta_{0,1}$  (CP<sub>1</sub>):  $\text{pr}(\beta_{0,1} \in \text{CI}_1)$ ;
- (ii) Coverage probability for  $\beta_{0,500}$  (CP<sub>500</sub>):  $\text{pr}(\beta_{0,500} \in \text{CI}_{500})$ ;
- (iii) Average coverage probability for  $\mathcal{B}$  (ACP <sub>$\mathcal{B}$</sub> ):  $\sum_{k \in \mathcal{B}} \text{pr}(\beta_{0,k} \in \text{CI}_k)/q$ ;
- (iv) Average coverage probability for  $\mathcal{B}^c$  (ACP <sub>$\mathcal{B}^c$</sub> ):  $\sum_{k \in \mathcal{B}^c} \text{pr}(0 \in \text{CI}_k)/(p - q)$ ;
- (v) Average rejection probability/power for  $\mathcal{B}$  (ARP <sub>$\mathcal{B}$</sub> ):  $1 - \sum_{k \in \mathcal{B}} \text{pr}(0 \in \text{CI}_k)/q$ ;
- (vi) Rejection and coverage probability for  $\beta_{0,1}$  (RCP<sub>1</sub>):  $\text{pr}(0 \notin \text{CI}_1, \beta_{0,1} \in \text{CI}_1)$ ;
- (vii) Average interval length for  $\mathcal{B}$  (AIL <sub>$\mathcal{B}$</sub> ):  $\sum_{k \in \mathcal{B}} \text{length}(\text{CI}_k)/q$ ;
- (viii) Average interval length for  $\mathcal{B}^c$  (AIL <sub>$\mathcal{B}^c$</sub> ):  $\sum_{k \in \mathcal{B}^c} \text{length}(\text{CI}_k)/(p - q)$ ,

where  $\text{CI}_k$  denotes the two-sided 95% confidence interval for  $\beta_{0,k}$ .

The results reported in Table 2 are based on 200 simulations. Both debiased DSRLasso and debiased TSLasso show coverage probabilities close to the nominal 95% level for  $\mathcal{B}^c$ . For  $\mathcal{B}$ , however, debiased DSRLasso exhibits coverage probabilities closer to 95% in most cases compared to debiased TSLasso. Furthermore, debiased DSRLasso demonstrates greater robustness to heavy-tailed errors. For heavy-tailed error distributions such as mixture-normal, log-normal,  $t(2)$ , and Cauchy, the confidence intervals of debiased DSRLasso are significantly shorter than those of debiased TSLasso. In particular, debiased DSRLasso consistently maintains interval lengths below 1.5 across all settings, whereas debiased TSLasso's intervals exceed 10 under Cauchy errors. When errors follow  $t(2)$  or Cauchy distributions, debiased DSRLasso achieves significantly higher power in detecting nonzero coefficients. For example, under Cauchy errors with  $N = 450$ , debiased DSRLasso attains values of 0.992 and 0.940 for metrics (v) and (vi), respectively, whereas debiased TSLasso yields substantially lower values of 0.438 and 0.360. This implies that under heavy-tailed errors, the excessively wide confidence intervals of debiased TSLasso can lead to the

| Error                 | $N$ | Method                  | $CP_1$ | $CP_{500}$ | $ACP_{\mathcal{B}}$ | $ACP_{\mathcal{B}^c}$ | $ARP_{\mathcal{B}}$ | $RCP_1$ | $AIL_{\mathcal{B}}$ | $AIL_{\mathcal{B}^c}$ |
|-----------------------|-----|-------------------------|--------|------------|---------------------|-----------------------|---------------------|---------|---------------------|-----------------------|
| $\mathcal{N}(0, 2^2)$ | 300 | TSLasso                 | 0.870  | 0.940      | 0.848               | 0.948                 | 1.000               | 0.870   | 0.454               | 0.457                 |
|                       |     | DSRLasso                | 0.915  | 0.925      | 0.898               | 0.947                 | 1.000               | 0.915   | 0.617               | 0.600                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.910  | 0.930      | 0.893               | 0.947                 | 1.000               | 0.910   | 0.618               | 0.601                 |
|                       | 450 | TSLasso                 | 0.925  | 0.940      | 0.928               | 0.949                 | 1.000               | 0.925   | 0.384               | 0.388                 |
|                       |     | DSRLasso                | 0.955  | 0.955      | 0.942               | 0.949                 | 1.000               | 0.955   | 0.483               | 0.476                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.960  | 0.950      | 0.943               | 0.949                 | 1.000               | 0.960   | 0.483               | 0.477                 |
| LogN                  | 300 | TSLasso                 | 0.885  | 0.945      | 0.872               | 0.949                 | 1.000               | 0.885   | 0.467               | 0.470                 |
|                       |     | DSRLasso                | 0.950  | 0.950      | 0.925               | 0.952                 | 1.000               | 0.950   | 0.486               | 0.473                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.965  | 0.950      | 0.933               | 0.952                 | 1.000               | 0.965   | 0.487               | 0.474                 |
|                       | 450 | TSLasso                 | 0.915  | 0.935      | 0.903               | 0.949                 | 1.000               | 0.915   | 0.406               | 0.410                 |
|                       |     | DSRLasso                | 0.960  | 0.935      | 0.953               | 0.953                 | 1.000               | 0.960   | 0.363               | 0.359                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.950  | 0.935      | 0.953               | 0.953                 | 1.000               | 0.950   | 0.364               | 0.359                 |
| MN                    | 300 | TSLasso                 | 0.830  | 0.955      | 0.858               | 0.948                 | 1.000               | 0.830   | 0.549               | 0.552                 |
|                       |     | DSRLasso                | 0.900  | 0.960      | 0.902               | 0.949                 | 1.000               | 0.900   | 0.435               | 0.424                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.900  | 0.955      | 0.903               | 0.949                 | 1.000               | 0.900   | 0.436               | 0.424                 |
|                       | 450 | TSLasso                 | 0.905  | 0.950      | 0.875               | 0.949                 | 1.000               | 0.905   | 0.464               | 0.470                 |
|                       |     | DSRLasso                | 0.945  | 0.950      | 0.937               | 0.948                 | 1.000               | 0.945   | 0.329               | 0.326                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.955  | 0.950      | 0.942               | 0.948                 | 1.000               | 0.955   | 0.330               | 0.326                 |
| $t(2)$                | 300 | TSLasso                 | 0.875  | 0.925      | 0.860               | 0.949                 | 0.985               | 0.870   | 0.936               | 0.942                 |
|                       |     | DSRLasso                | 0.940  | 0.940      | 0.918               | 0.951                 | 1.000               | 0.940   | 0.754               | 0.752                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.940  | 0.945      | 0.927               | 0.951                 | 1.000               | 0.940   | 0.755               | 0.753                 |
|                       | 450 | TSLasso                 | 0.905  | 0.960      | 0.903               | 0.948                 | 0.985               | 0.885   | 0.827               | 0.836                 |
|                       |     | DSRLasso                | 0.950  | 0.960      | 0.940               | 0.949                 | 1.000               | 0.950   | 0.555               | 0.560                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.965  | 0.960      | 0.943               | 0.949                 | 1.000               | 0.965   | 0.556               | 0.562                 |
| Cauchy                | 300 | TSLasso                 | 0.945  | 0.960      | 0.922               | 0.949                 | 0.445               | 0.425   | 10.30               | 10.34                 |
|                       |     | DSRLasso                | 0.950  | 0.925      | 0.955               | 0.947                 | 0.973               | 0.930   | 1.457               | 1.454                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.940  | 0.935      | 0.952               | 0.948                 | 0.973               | 0.920   | 1.461               | 1.458                 |
|                       | 450 | TSLasso                 | 0.915  | 0.965      | 0.913               | 0.949                 | 0.438               | 0.360   | 12.90               | 13.10                 |
|                       |     | DSRLasso                | 0.950  | 0.945      | 0.937               | 0.951                 | 0.992               | 0.940   | 1.104               | 1.116                 |
|                       |     | DSRLasso <sup>sub</sup> | 0.945  | 0.935      | 0.933               | 0.952                 | 0.987               | 0.930   | 1.137               | 1.150                 |

Table 2: Simulation Results for Example 2. Empirical Coverage Probabilities, Power and Lengths of Confidence Intervals Based on 200 Simulations.

omission of important variables. Finally, debiased DSRLasso<sup>sub</sup> and debiased DSRLasso show comparable performance.

**Example 3.** We compare the performance of the proposed debiased DSRLasso and the debiased TSLasso in simultaneous inference, using the following three metrics:

- (a) Coverage probability (CP):  $\text{pr}[\cap_{k \in \mathcal{G}} \{\beta_{0,k} \in \text{SCI}_k\}]$ ;
- (b) Rejection probability/power (RP):  $\text{pr}[\cup_{k \in \mathcal{G}} \{0 \notin \text{SCI}_k\}]$ ;
- (c) Average interval length (AIL):  $\sum_{k \in \mathcal{G}} \text{length}(\text{SCI}_k)/|\mathcal{G}|$ ,

where  $\{\text{SCI}_k\}_{k \in \mathcal{G}}$  denotes the 95% simultaneous confidence interval (SCI) for  $\{\beta_{0,k}\}_{k \in \mathcal{G}}$ .

We set  $\mathcal{G} = \{1, \dots, p\}$  and consider both non-studentized (NST) and studentized (ST) SCIs. The NST intervals have equal lengths across all  $k \in \mathcal{G}$ , whereas the ST intervals have varying lengths. Empirical values of the three metrics, estimated from 200 simulations, are summarized in Table 3. The overall pattern aligns with Table 2. Generally, SCIs based on the debiased DSRLasso exhibit coverage probabilities closer to the nominal 95% level than those based on the debiased TSLasso. Moreover, the debiased DSRLasso-based intervals are significantly shorter under heavy-tailed error distributions. When errors follow  $t(2)$  or Cauchy distributions, the debiased DSRLasso-based method has significantly better power performance than the debiased TSLasso-based method in testing the simultaneous hypothesis  $H_{0,\mathcal{G}} : \beta_{0,k} = 0$  for all  $k \in \mathcal{G}$ . Additionally, the performance of ST-based methods is comparable to NST-based methods, though the NST intervals tend to be slightly longer. Finally, DSRLasso<sup>sub</sup>-based and DSRLasso-based SCIs show similar performance.

## 4.2 A Real Example

In this section, we apply the proposed methods to analyze the breast cancer dataset GSE93601 from the GEO database (<https://www.ncbi.nlm.nih.gov/geo>). This dataset was analyzed in Wang et al. (2015) and Zhao et al. (2022). It includes clinical information, such as age, body mass index (BMI), and alcohol consumption, from participants of the Nurses' Health Study I and II who were diagnosed with invasive breast cancer, along with gene expression profiling of 602 breast tumors and 508 tumor-adjacent normal tissues. Among the genes studied, TP53 (tumor protein p53) plays a critical role in inducing apoptosis under various stress conditions, including oncogene activation or DNA damage. Loss of p53 tumor-suppressor activity, whether through mutation/deletion of the TP53 gene or inhibition of the p53 protein, promotes cancer development. Our objective is to identify genes whose expression levels are associated with TP53. Hence, we use TP53 expression as the response variable and the expression levels of other genes as explanatory variables.

We restrict our analysis to breast tumor tissues. Individuals with missing age, BMI, or alcohol consumption data are excluded, resulting in a final sample size of  $n = 596$ . Genes with a missing rate exceeding 10% are removed, and the remaining missing gene expression values are imputed using the “impute.knn” from the R package `impute`. Among the retained genes, the MDM2 (mouse double minute 2 homolog) protein is a known negative regulator of p53. By binding to p53, it suppresses its transcriptional activity, promotes nuclear export, and enhances degradation. Overexpression of MDM2 in various tumors inhibits p53, thereby facilitating uncontrolled cell proliferation (Chène, 2003). Therefore, we include MDM2 in

| Error                 | $N$ | Method | TSLasso |       |       | DSRLasso |       |       | DSRLasso <sup>sub</sup> |       |       |
|-----------------------|-----|--------|---------|-------|-------|----------|-------|-------|-------------------------|-------|-------|
|                       |     |        | CP      | RP    | AIL   | CP       | RP    | AIL   | CP                      | RP    | AIL   |
| $\mathcal{N}(0, 2^2)$ | 300 | NST    | 0.855   | 1.000 | 0.912 | 0.940    | 1.000 | 1.207 | 0.940                   | 1.000 | 1.209 |
|                       |     | ST     | 0.850   | 1.000 | 0.903 | 0.920    | 1.000 | 1.189 | 0.940                   | 1.000 | 1.191 |
|                       | 450 | NST    | 0.880   | 1.000 | 0.773 | 0.935    | 1.000 | 0.951 | 0.935                   | 1.000 | 0.952 |
|                       |     | ST     | 0.900   | 1.000 | 0.768 | 0.940    | 1.000 | 0.944 | 0.940                   | 1.000 | 0.944 |
| LogN                  | 300 | NST    | 0.840   | 1.000 | 0.938 | 0.955    | 1.000 | 1.014 | 0.960                   | 1.000 | 1.016 |
|                       |     | ST     | 0.855   | 1.000 | 0.929 | 0.955    | 1.000 | 0.997 | 0.950                   | 1.000 | 0.999 |
|                       | 450 | NST    | 0.920   | 1.000 | 0.817 | 0.955    | 1.000 | 0.713 | 0.960                   | 1.000 | 0.714 |
|                       |     | ST     | 0.910   | 1.000 | 0.812 | 0.960    | 1.000 | 0.706 | 0.960                   | 1.000 | 0.707 |
| MN                    | 300 | NST    | 0.870   | 1.000 | 1.103 | 0.900    | 1.000 | 0.851 | 0.905                   | 1.000 | 0.852 |
|                       |     | ST     | 0.850   | 1.000 | 1.093 | 0.940    | 1.000 | 0.838 | 0.945                   | 1.000 | 0.839 |
|                       | 450 | NST    | 0.870   | 1.000 | 0.934 | 0.955    | 1.000 | 0.651 | 0.950                   | 1.000 | 0.651 |
|                       |     | ST     | 0.880   | 1.000 | 0.929 | 0.945    | 1.000 | 0.646 | 0.945                   | 1.000 | 0.646 |
| $t(2)$                | 300 | NST    | 0.905   | 0.980 | 1.879 | 0.945    | 1.000 | 1.514 | 0.950                   | 1.000 | 1.516 |
|                       |     | ST     | 0.880   | 0.980 | 1.861 | 0.950    | 1.000 | 1.487 | 0.950                   | 1.000 | 1.489 |
|                       | 450 | NST    | 0.930   | 0.985 | 1.663 | 0.965    | 1.000 | 1.119 | 0.965                   | 1.000 | 1.121 |
|                       |     | ST     | 0.935   | 0.985 | 1.652 | 0.965    | 1.000 | 1.107 | 0.960                   | 1.000 | 1.109 |
| Cauchy                | 300 | NST    | 0.890   | 0.325 | 20.64 | 0.915    | 0.915 | 3.769 | 0.900                   | 0.910 | 3.777 |
|                       |     | ST     | 0.910   | 0.320 | 20.39 | 0.905    | 0.765 | 3.583 | 0.905                   | 0.775 | 3.590 |
|                       | 450 | NST    | 0.905   | 0.320 | 26.01 | 0.930    | 0.985 | 2.868 | 0.935                   | 0.980 | 2.900 |
|                       |     | ST     | 0.910   | 0.330 | 25.81 | 0.960    | 0.955 | 2.810 | 0.960                   | 0.945 | 2.841 |

Table 3: Simulation Results of Example 3. Empirical Coverage Probabilities, Power and Average Interval Lengths for Non-studentized (NST) and Studentized (ST) Methods Based on 200 Simulations.

our analysis and assume a functional relationship between MDM2 and TP53. For the remaining genes, we further filter the genes by including only those with expression values that fall within the top 800 in terms of absolute Spearman correlation with TP53.

Given that aging, excess body weight, and alcohol consumption are widely recognized as significant risk factors for many cancers (Renehan et al., 2008; Yoo et al., 2022; Montégut et al., 2024), we also include age, BMI, and alcohol intake as explanatory variables. We define two dummy variables for BMI and two for alcohol consumption. Specifically, for  $i = 1, \dots, 596$ , let  $Y_i = \text{TP53}_i$ ,  $Z_i = \text{MDM2}_i$ , and  $\mathbf{x}_i = \{X_{i,k}\}_{k=1}^{805} \in \mathbb{R}^{805}$ , where  $X_{i,1} = \text{Age}_i$ ,  $X_{i,2} = I(\text{BMI}_i > 25 \text{ kg/m}^2)$ ,  $X_{i,3} = I(\text{BMI}_i > 30 \text{ kg/m}^2)$ ,  $X_{i,4} = I(0 \text{ g/day} < \text{Alcohol}_i \leq 10 \text{ g/day})$ ,  $X_{i,5} = I(\text{Alcohol}_i > 10 \text{ g/day})$ , and  $\{X_{i,k}\}_{k=6}^{805}$  represents the expression values of the 800 selected genes for the  $i$ -th sample. We consider the following partially linear model:

$$Y_i = g_0(Z_i) + \sum_{k=1}^{805} \beta_{0,k} X_{i,k} + \epsilon_i, \quad i = 1, \dots, 596.$$

The kurtosis of TP53 expression values is 2.36, suggesting a potentially heavy-tailed distribution. Let  $p_k^{\text{TSLasso}}$  and  $p_k^{\text{DSRLasso}}$  denote the  $p$ -values from the debiased two-step projection-based Lasso estimator and the proposed debiased DS rank Lasso estimator, respectively, for testing the significance of  $\beta_{0,k}$ . At the 1% significance level, DSRLasso indicates a significant positive effect of aging ( $p_k^{\text{DSRLasso}} = 0.003$ ) consistent with previous findings that aging is a significant risk factor for cancer, whereas the Lasso-based method does not ( $p_k^{\text{TSLasso}} = 0.059$ ). A summary of the  $p$ -values for the regression coefficients of the genes from both methods is provided in Table 4.

Overall, the genes selected by the two methods show considerable overlap at the 5% significance level. In particular, the  $p$ -values for the regression coefficients of the genes  $\{\text{C17orf85}, \text{POLR2B}, \text{COL10A1}, \text{KIAA0368}, \text{PRODH2}\}$  are all below 0.01 in both methods, and these genes have been previously implicated in tumorigenesis and prognosis (Zacharakis et al., 2018; Zhang et al., 2019; Hu et al., 2025; Wang et al., 2025; Gong et al., 2025). Moreover, the proposed debiased DSRLasso method identifies several genes not detected by the debiased TSLasso method. Specifically, the  $p$ -values for  $\{\text{PAPOLA}, \text{HBP1}, \text{SMG7}, \text{EP300}, \text{WDR75}\}$  are all below 0.05 for DSRLasso, whereas they exceed 0.05 for TSLasso. These genes have been linked to p53 or breast cancer in existing literature (Yang et al., 2020; Komini et al., 2021; Moudry et al., 2022; Zhou et al., 2023; Gronkowska and Robaszkiewicz, 2024). This demonstrates that the proposed debiased DSRLasso method is effective in identifying genes that the debiased TSLasso method fails to detect.

We also apply the proposed methods to a worker wage dataset in Section B.6 of the Appendix.

## 5. Discussion

In this study, we propose the DS rank Lasso, a novel regularized method for high-dimensional partially linear models. We also develop a debiased DS rank Lasso estimator to enable inference for individual coefficients as well as simultaneous inference. All proposed methods exhibit robustness against heavy-tailed random errors. To preserve the near-oracle rate

|                                          | $p_k^{\text{TSLasso}} \in [0, 0.01)$                              | $p_k^{\text{TSLasso}} \in [0.01, 0.05)$                                                                                                                                                            | $p_k^{\text{TSLasso}} \in [0.05, 1]$                                                                                                       |
|------------------------------------------|-------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------|
| $p_k^{\text{DSRLasso}} \in [0, 0.01)$    | C17orf85(NCBP3),<br>POLR2B, COL10A1,<br>KIAA0368<br>PRODH2(NPHS1) | C13orf23, N-PAC<br>PIBF1(C13orf24),<br>STAU2, RBM39                                                                                                                                                | PAPOLA,<br>Q6ZU24(FER1L5)                                                                                                                  |
| $p_k^{\text{DSRLasso}} \in [0.01, 0.05)$ | BRWD1, SFPQ,<br>WDR57, TRAPPC1,<br>STAB2, ZC3H11A,<br>NCKAP1      | SET, NAT10, UBA3,<br>CNOT6, GOPC,<br>ATP2A2, TMEM131,<br>NUP188, SETD2,<br>WAC                                                                                                                     | WDR42A, ZZEF1,<br>HBP1, SPTLC2,<br>SMG7, WDR75,<br>NUPL1, EP300,<br>ZNF23(ZNF19),<br>NUP214, RPN1,<br>ZC3H7A, KIAA0494,<br>UBQLN1, SLC33A1 |
| $p_k^{\text{DSRLasso}} \in [0.05, 1]$    | GFM1, DNAH2,<br>DDX5                                              | WNK1, POLR2A,<br>BIRC6, AKAP11,<br>MYOM1, TUBA1B,<br>LOC645159(MLL3),<br>KPNB1, KIAA0196,<br>ELF1, SEC31A,<br>TRAM1, PSMD1,<br>NRAP, CEP350,<br>ANXA7, ZNF124,<br>LRRC16A(LRRC16),<br>GNPTAB, NRD1 | Other Genes                                                                                                                                |

Table 4: A Summary of the  $p$ -values for the Regression Coefficients of the 800 Genes.

of rank Lasso under mild conditions on nonparametric estimators, we incorporate data-splitting into the regularized regression method. The two-split procedure is also used for inference, where the proposed debiased DS rank Lasso estimator is constructed as an average of two debiased estimators. Furthermore, we demonstrate that data splitting asymptotically incurs no loss of statistical efficiency in the inference procedure.

## Acknowledgments

The authors would like to thank the editor and reviewers for their helpful comments. Yang's research is supported by the BRAIN of Renmin University of China and NSFC 12301389. Zhao's research is supported by the Natural Science Foundation of Fujian Province, China (No. 2026J008060) and the Startup Fund of Fuzhou University (No. 511839).

## Appendix A. Technical Details

In this appendix, we provide proofs for Lemma 1 and Lemma 4, and Theorem 2, Theorem 5, and Theorem 7 in the main text, assuming without loss of generality that each subsample has an equal size of  $n$ .

We define the event

$$\mathcal{N} = \left\{ \max_{k \in \{0,1,\dots,p\}} \sup_{z \in \mathcal{Z}} |\hat{m}_k^{(1)}(z) - m_k(z)| \leq r_n \right\}.$$

Consequently, we have  $\text{pr}(\mathcal{N}^c) \leq \rho_n$  by Condition (C1). For notational simplicity, we denote  $a_n = 2(\|\beta_0\|_1 \vee 1)r_n$ ,  $\hat{\Delta}_k^{(1)}(Z_i^{(2)}) = \hat{m}_k^{(1)}(Z_i^{(2)}) - m_k(Z_i^{(2)})$ ,  $\hat{\Delta}_g^{(1)}(Z_i^{(2)}) = \hat{g}(Z_i^{(2)}, \beta_0) - g_0(Z_i^{(2)})$ , and use  $E_*(\cdot)$  and  $\text{pr}_*(\cdot)$  to represent the conditional expectation and conditional probability given  $\mathcal{D}_1$ .

### A.1 Proof of Lemma 1

**Proof** We will prove the results only for  $l = 1$ , as the results for  $l = 2$  can be obtained directly by swapping the roles of the two subsamples.

For the first part of Lemma 1, we will derive a high probability upper bound for  $|\|\mathbf{S}^{(2)}(\mathbf{0})\|_\infty - \|\mathbf{S}_0^{(2)}\|_\infty| \leq \|\mathbf{S}^{(2)}(\mathbf{0}) - \mathbf{S}_0^{(2)}\|_\infty$ , where

$$\begin{aligned} \mathbf{S}^{(2)}(\gamma) &= -2\{n(n-1)\}^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(2)} \sum_{j=1, j \neq i}^n \text{sign}\{(e_i^{(2)} - e_j^{(2)}) - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\} \\ &= -\{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) \text{sign}\{(e_i^{(2)} - e_j^{(2)}) - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\}, \end{aligned}$$

and

$$\begin{aligned} \mathbf{S}_0^{(2)} &= -2\{n(n-1)\}^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(2)} \sum_{j=1, j \neq i}^n \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)}) \\ &= -\{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)}). \end{aligned}$$

This can be achieved by upper bounding  $\|\mathbf{S}_1\|_\infty = \|\{n(n-1)\}^{-1} \sum_{i \neq j} \mathbf{S}_{1,i,j}\|_\infty$  and  $\|\mathbf{S}_2\|_\infty = \|\{n(n-1)\}^{-1} \sum_{i \neq j} \mathbf{S}_{2,i,j}\|_\infty$ , where  $\mathbf{S}_1 = (S_{1,1}, \dots, S_{1,p})^\top$ ,  $\mathbf{S}_{1,i,j} = (S_{1,i,j,1}, \dots, S_{1,i,j,p})^\top$ ,  $S_{1,i,j,k} = (\tilde{X}_{i,k}^{(2)} - \tilde{X}_{j,k}^{(2)})\{\text{sign}(e_i^{(2)} - e_j^{(2)}) - \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)})\}$ ,  $\mathbf{S}_2 = (S_{2,1}, \dots, S_{2,p})^\top$ ,  $\mathbf{S}_{2,i,j} = (S_{2,i,j,1}, \dots, S_{2,i,j,p})^\top$ , and  $S_{2,i,j,k} = \{\hat{\Delta}_k^{(1)}(Z_i^{(2)}) - \hat{\Delta}_k^{(1)}(Z_j^{(2)})\}\{\text{sign}(e_i^{(2)} - e_j^{(2)}) - \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)})\}$ .

Note that for any  $t > 0$ ,

$$\text{pr}_*(|S_{1,k}| \geq t) = \text{pr}_*(S_{1,k} \geq t) + \text{pr}_*(S_{1,k} \leq -t).$$

We denote  $M_n = \lfloor n/2 \rfloor$ , the largest integer smaller than or equal to  $n/2$ , and  $\pi$  the permutations of  $\{1, \dots, n\}$ . For  $\text{pr}_*(S_{1,k} \geq t)$ , we have for any  $\lambda > 0$ , that

$$\begin{aligned} \text{pr}_*(S_{1,k} \geq t) &\leq E_* \left( \exp \left[ \lambda \{n(n-1)\}^{-1} \sum_{i \neq j} S_{1,i,j,k} \right] \right) / \exp(\lambda t) \\ &= E_* \left\{ \exp \left( \lambda \frac{1}{n!} \sum_{\pi} M_n^{-1} \sum_{i=1}^{M_n} S_{1,\pi(i),\pi(M_n+i),k} \right) \right\} / \exp(\lambda t) \\ &\leq \frac{1}{n!} \sum_{\pi} E_* \left\{ \exp \left( \lambda M_n^{-1} \sum_{i=1}^{M_n} S_{1,\pi(i),\pi(M_n+i),k} \right) \right\} / \exp(\lambda t). \end{aligned}$$

The two inequalities follow from Markov's and Jensen's inequalities, respectively, and the equality is due to the property of  $U$ -statistic. It's thus sufficient to deal with the summand for any given  $\pi$ . We will frequently use the above arguments in our proofs. Because  $\tilde{X}_{\pi(i),k}^{(2)} - \tilde{X}_{\pi(M_n+i),k}^{(2)}$  is zero-mean sub-Gaussian with parameter  $\sqrt{2}\sigma$  under Condition (C2),  $E_* |\tilde{X}_{\pi(i),k}^{(2)} - \tilde{X}_{\pi(M_n+i),k}^{(2)}|^t \leq (4\sigma^2)^{t/2} t\Gamma(t/2)$  by using Proposition 2.5.2 (ii) in Vershynin (2018), where  $\Gamma(\cdot)$  represents the Gamma function. Therefore, on  $\mathcal{N}$ ,  $E_* |S_{1,\pi(i),\pi(M_n+i),k}|^2 \leq 32\sigma^2(2\kappa_u a_n)$ , and for  $t \geq 3$ ,

$$\begin{aligned} E_* |S_{1,\pi(i),\pi(M_n+i),k}|^t &\leq (4\sigma^2)^{t/2} t\Gamma(t/2) 2^t (4\kappa_u a_n) \\ &\leq \{t\Gamma(t/2)/2\} (64\sigma^2 \kappa_u a_n) (4\sigma)^{t-2}, \end{aligned}$$

so that  $S_{1,\pi(i),\pi(M_n+i),k}$  satisfies Bernstein's condition (see, for example, Wainwright, 2019, Chapter 2.1) with parameters  $(64\sigma^2 \kappa_u a_n, 4\sigma)$ . Therefore, following similar arguments to Wainwright (2019, Proposition 2.10),

$$E_* \left\{ \exp \left( \lambda M_n^{-1} \sum_{i=1}^{M_n} S_{1,\pi(i),\pi(M_n+i),k} \right) \right\} \leq \exp[32\lambda^2 \sigma^2 \kappa_u a_n / \{M_n(1 - 4\sigma|\lambda|M_n^{-1})\}],$$

for all  $|\lambda|M_n^{-1} < 1/(4\sigma)$ . By letting  $\lambda = M_n t / (4\sigma t + 64\sigma^2 \kappa_u a_n)$ , we have

$$\text{pr}_*(S_{1,k} \geq t) \leq \exp[-M_n t^2 / \{2(4\sigma t + 64\sigma^2 \kappa_u a_n)\}],$$

and hence

$$\begin{aligned} \text{pr}_*(|S_{1,k}| \geq t) &\leq 2 \exp[-M_n t^2 / \{2(4\sigma t + 64\sigma^2 \kappa_u a_n)\}] \\ &\leq 2 \exp[-M_n t^2 / \{4(4\sigma t \vee 64\sigma^2 \kappa_u a_n)\}], \end{aligned}$$

by processing  $\text{pr}_*(S_{1,k} \leq -t)$  in similar steps. By the union bound, we have that

$$\text{pr}_*(\|\mathbf{S}_1\|_{\infty} \geq t) \leq 2 \exp[-M_n t^2 / \{4(4\sigma t \vee 64\sigma^2 \kappa_u a_n)\}] + \log p.$$

Denote  $\mathcal{S}_1 = [\|\mathbf{S}_1\|_{\infty} \leq 16\sqrt{2}C\sigma\sqrt{\log p/n} \{(\kappa_u a_n)^{1/2} \vee (2\log p/n)^{1/4}\}]$ , and then  $\text{pr}_*(\mathcal{S}_1^c) \leq 2 \exp[-\{C^2 \wedge C(2\log p/n)^{-1/4}\} - 1] \log p$  on  $\mathcal{N}$ .

As for  $|S_{2,k}|$ , we can follow a similar paradigm to  $|S_{1,k}|$ . Specifically,

$$\begin{aligned} |E_*(S_{2,\pi(i),\pi(M_n+i),k})| &= \left| E_* \left( \{ \widehat{\Delta}_k^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_k^{(1)}(Z_j^{(2)}) \} \right. \right. \\ &\quad \left. \left. \times [2F^* \{ \widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)}) \} - 1] \right) \right| \leq 8\kappa_u r_n a_n, \end{aligned}$$

under Condition (C4) and on  $\mathcal{N}$ . Moreover, we note that

$$\begin{aligned} &E_* |S_{2,\pi(i),\pi(M_n+i),k} - E_*(S_{2,\pi(i),\pi(M_n+i),k})|^2 \\ &\leq E_* |S_{2,\pi(i),\pi(M_n+i),k}|^2 \leq 32r_n^2 \kappa_u a_n, \end{aligned}$$

and for  $t \geq 3$ ,

$$\begin{aligned} &E_* |S_{2,\pi(i),\pi(M_n+i),k} - E_*(S_{2,\pi(i),\pi(M_n+i),k})|^t \\ &\leq 2^t E_* |S_{2,\pi(i),\pi(M_n+i),k}|^t \leq (32r_n^2 \kappa_u a_n)(4r_n)^{t-2} 2^t \\ &= 2^{t-2} (128r_n^2 \kappa_u a_n)(4r_n)^{t-2} \leq t! / 2 (128r_n^2 \kappa_u a_n)(4r_n)^{t-2}, \end{aligned}$$

on  $\mathcal{N}$ , which implies  $S_{2,\pi(i),\pi(M_n+i),k}$  satisfies Bernstein's condition (see, for example, Wainwright, 2019, Chapter 2.1) with parameters  $(128r_n^2 \kappa_u a_n, 4r_n)$ . Following a similar argument to  $\mathcal{S}_1$ ,  $\text{pr}_*(\mathcal{S}_2^c) \leq 2 \exp(-[\{C^2 \wedge C(\log p/n)^{-1/4}\} - 1] \log p)$  when  $\mathcal{D}_1 \in \mathcal{N}$ , where  $\mathcal{S}_2 = [\|\mathbf{S}_2\|_\infty \leq 8\kappa_u r_n a_n + 32Cr_n \sqrt{\log p/n} \{(\kappa_u a_n)^{1/2} \vee (\log p/n)^{1/4}\}]$ .

To conclude,

$$\|\mathbf{S}_1\|_\infty + \|\mathbf{S}_2\|_\infty \leq 8\kappa_u r_n a_n + 32C(\sigma + r_n) \sqrt{\log p/n} \{(\kappa_u a_n)^{1/2} \vee (2 \log p/n)^{1/4}\},$$

on  $\mathcal{S}_1 \cap \mathcal{S}_2$ , and  $\text{pr}_*\{(\mathcal{S}_1 \cap \mathcal{S}_2)^c\} \leq 4 \exp(-[\{C^2 \wedge C(2 \log p/n)^{-1/4}\} - 1] \log p)$  on  $\mathcal{N}$ . This implies for any constant  $c_0 > 0$ , by letting

$$C = \left[ \{c_0 - 8\kappa_u r_n a_n (\log p/n)^{-1/2}\} / \{32(\sigma + r_n)\} \right] \{(\kappa_u a_n)^{1/2} \vee (2 \log p/n)^{1/4}\}^{-1},$$

we have  $\|\mathbf{S}_1\|_\infty + \|\mathbf{S}_2\|_\infty \leq c_0 \sqrt{\log p/n}$ . Therefore,  $\text{pr}_*\{(\mathcal{S}_1 \cap \mathcal{S}_2)^c\} \leq 4p^{-(C_n^2 - 1)}$  on  $\mathcal{N}$ , where

$$C_n = \left( \left[ \{c_0 - 8\kappa_u r_n a_n (\log p/n)^{-1/2}\} / \{32(\sigma + r_n)\} \right] \wedge 1 \right) \{(\kappa_u a_n)^{1/2} \vee (2 \log p/n)^{1/4}\}^{-1}.$$

Denote  $\widehat{\gamma}^{(1)}(\widehat{\lambda}^{(1)})$  by  $\widehat{\gamma}^{(1)}$  for simplicity. By the definition of  $\widehat{\beta}^{(1)}(\widehat{\lambda}^{(1)})$ , we have,

$$Q^{(2)}(\widehat{\gamma}^{(1)}) + \widehat{\lambda}^{(1)} \|\widehat{\beta}^{(1)}\|_1 \leq Q^{(2)}(\mathbf{0}) + \widehat{\lambda}^{(1)} \|\beta_0\|_1.$$

This implies

$$Q^{(2)}(\widehat{\gamma}^{(1)}) - Q^{(2)}(\mathbf{0}) \leq \widehat{\lambda}^{(1)} (\|\beta_0\|_1 - \|\widehat{\beta}^{(1)}\|_1) \leq \widehat{\lambda}^{(1)} (\|\widehat{\gamma}_B^{(1)}\|_1 - \|\widehat{\gamma}_{B^c}^{(1)}\|_1).$$

By simple calculation, the subgradient of  $Q^{(2)}(\gamma)$  is

$$\begin{aligned} &-\{n(n-1)\}^{-1} \sum_{i \neq j} (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)}) \text{sign}\{(e_i^{(2)} - e_j^{(2)}) - (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \gamma\} \\ &\quad - \{n(n-1)\}^{-1} \sum_{i \neq j} (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)}) v_{i,j}^{(2)} I\{e_i^{(2)} - e_j^{(2)} = (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \gamma\}, \end{aligned}$$

where  $v_{i,j}^{(2)} \in [-1, 1]$ . We note that with probability one, the subgradient of  $Q^{(2)}(\gamma)$  at  $\gamma = \mathbf{0}$  is exactly  $\mathbf{S}^{(2)}(\mathbf{0})$  because the distribution of  $\epsilon_i^{(2)}$  has no point masses. By the convexity of  $Q(\cdot)$ , the definition of subdifferential and the above discussion,

$$\begin{aligned} Q^{(2)}(\hat{\gamma}^{(1)}) - Q^{(2)}(\mathbf{0}) &\geq -\|\hat{\gamma}^{(1)}\|_1 \|\mathbf{S}^{(2)}(\mathbf{0})\|_\infty \\ &\geq -\|\hat{\gamma}^{(1)}\|_1 (\|\mathbf{S}_0^{(2)}\|_\infty + c_0 \sqrt{\log p/n}) \\ &\geq -(\|\hat{\gamma}_{\mathcal{B}}^{(1)}\|_1 + \|\hat{\gamma}_{\mathcal{B}^c}^{(1)}\|_1) (\hat{\lambda}^{(1)}/c), \end{aligned}$$

where the second inequality holds on  $\mathcal{S}_1 \cap \mathcal{S}_2$ , and the third inequality holds with probability at least  $1 - \alpha_0$  by the definition of  $\hat{\lambda}^{(1)}$ . This implies  $\text{pr}_* \left( \|\hat{\gamma}_{\mathcal{B}^c}^{(1)}\|_1 \leq \bar{c} \|\hat{\gamma}_{\mathcal{B}}^{(1)}\|_1 \right) = \text{pr}_* \left\{ \hat{\gamma}^{(1)}(\hat{\lambda}^{(1)}) \in \Gamma \right\} \geq 1 - \alpha_0 - 4p^{-(C_n^2-1)}$ . Here,  $C_n \rightarrow \infty$  as  $n \rightarrow \infty$  since  $a_n = o(1)$ ,  $a_n r_n = o(\sqrt{\log p/n})$  and  $\log p/n = o(1)$  by Condition (C1).

For the second part of Lemma 1, it suffices to upper bound  $\|\mathbf{S}_0^{(2)}\|_\infty = \|\{n(n-1)\}^{-1} \sum_{i \neq j} \mathbf{S}_{0,i,j}\|_\infty$ , conditioning on  $\{\hat{\mathbf{x}}_i^{(2)}\}_{i=1}^n$ , where  $\mathbf{S}_0^{(2)} = (S_{0,1}, \dots, S_{0,p})^\top$ ,  $\mathbf{S}_{0,i,j} = (S_{0,i,j,1}, \dots, S_{0,i,j,p})^\top$  and  $S_{0,i,j,k} = (\hat{X}_{i,k}^{(2)} - \hat{X}_{j,k}^{(2)}) \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)})$ .

We note that  $\mathcal{D}_1$  and  $\{\mathbf{x}_i^{(2)}, Z_i^{(2)}\}_{i=1}^n$  are independent of  $\{\epsilon_i^{(2)}\}_{i=1}^n$ , then conditioning on  $\mathcal{D}_1$  and  $\{\mathbf{x}_i^{(2)}, Z_i^{(2)}\}_{i=1}^n$ ,  $M_n^{-1} \sum_{i=1}^{M_n} S_{0,\pi(i),\pi(M_n+i),k}$  for any given  $\pi$  is zero-mean sub-Gaussian, with parameter at most

$$M_n^{-1} \left\{ \sum_{i=1}^{M_n} (\hat{X}_{\pi(i),k}^{(2)} - \hat{X}_{\pi(M_n+i),k}^{(2)})^2 \right\}^{1/2} \leq \left[ (16/n) \left\{ n^{-1} \sum_{i=1}^n (\tilde{X}_{i,k}^{(2)})^2 + r_n^2 \right\} \right]^{1/2},$$

on  $\mathcal{N}$ . Denote the event  $\mathcal{M}_1 = \{\max_k n^{-1} \sum_{i=1}^n (\tilde{X}_{i,k}^{(2)})^2 \leq 12\sigma^2\}$ , and then the above equation is upper bounded by  $\{(16/n)(12\sigma^2 + r_n^2)\}^{1/2}$  on  $\mathcal{M}_1 \cap \mathcal{N}$ . Using Honorio and Jaakkola (2014, Appendix B), since  $\tilde{X}_{i,k}^{(2)}$  is sub-Gaussian with parameter  $\sigma$ , its square is sub-exponential with parameter  $(4\sqrt{2}\sigma^2, 4\sigma^2)$ . Besides,  $E\{(\tilde{X}_{i,k}^{(2)})^2\} \leq 4\sigma^2$  by Vershynin (2018, Proposition 2.5.2), then

$$\begin{aligned} &\text{pr} \left\{ \max_k n^{-1} \sum_{i=1}^n (\tilde{X}_{i,k}^{(2)})^2 \geq t \right\} \leq \sum_{k=1}^p \text{pr} \left\{ n^{-1} \sum_{i=1}^n (\tilde{X}_{i,k}^{(2)})^2 \geq t \right\} \\ &\leq \sum_{k=1}^p \inf_{\lambda \in [0, (4\sigma^2/n)^{-1}]} E \left( \exp \left[ (\lambda/n) \sum_{i=1}^n [(\tilde{X}_{i,k}^{(2)})^2 - E\{(\tilde{X}_{i,k}^{(2)})^2\}] + E\{(\tilde{X}_{i,k}^{(2)})^2\} \right] \right) / \exp(\lambda t) \\ &\leq \sum_{k=1}^p \inf_{\lambda \in [0, (4\sigma^2/n)^{-1}]} \exp \{ 16\sigma^2 \lambda^2 / n - \lambda(t - 4\sigma^2) \} \\ &\leq \exp \left[ -n(t - 4\sigma^2)^2 / \max\{64\sigma^4, 8\sigma^2(t - 4\sigma^2)\} + \log p \right]. \end{aligned}$$

Consequently,  $\text{pr}(\mathcal{M}_1^c) \leq p \exp(-n)$ . Then, by properties of sub-Gaussian random variables and similar arguments to  $\|\mathbf{S}_1\|_\infty$  and  $\|\mathbf{S}_2\|_\infty$ , it follows that

$$\text{pr}[\|\mathbf{S}_0^{(2)}\|_\infty \geq t \mid \mathcal{D}_1, \{\mathbf{x}_i^{(2)}, Z_i^{(2)}\}_{i=1}^n] \leq 2 \exp\{-(n/32)t^2/(12\sigma^2 + r_n^2) + \log p\},$$

on  $\mathcal{N} \cap \mathcal{M}_1$ , where  $\boldsymbol{\epsilon}^{(2)} = (\epsilon_1^{(2)}, \dots, \epsilon_n^{(2)})^\top$ . Consequently, for any  $C > 1$ ,

$$\begin{aligned} & \Pr[\|\mathbf{S}_0^{(2)}\|_\infty \geq 4\sqrt{2}C(2\sqrt{3}\sigma + r_n)\sqrt{\log p/n} \mid \mathcal{D}_1, \{\mathbf{x}_i^{(2)}, Z_i^{(2)}\}_{i=1}^n] \\ & \leq 2 \exp\{-(C^2 - 1) \log p\}, \end{aligned}$$

on  $\mathcal{N} \cap \mathcal{M}_1$ , which implies, for  $\alpha_0 > 2p^{-(C^2-1)}$ ,  $\Pr[\|\mathbf{S}_0^{(2)}\|_\infty \geq 4\sqrt{2}C(2\sqrt{3}\sigma + r_n)\sqrt{\log p/n} \mid \mathcal{D}_1, \{\mathbf{x}_i^{(2)}, Z_i^{(2)}\}_{i=1}^n] < \alpha_0$ . By the definition of  $\hat{\lambda}^{(1)}$ ,  $\hat{\lambda}^{(1)}$  is the infimum of  $\lambda^{(1)}$  such that

$$\begin{aligned} & \Pr\left[c(\|\mathbf{S}_0^{(2)}\|_\infty + c_0\sqrt{\log p/n}) > \lambda^{(1)} \mid \mathcal{D}_1, \{\mathbf{x}_i^{(2)}, Z_i^{(2)}\}_{i=1}^n\right] \\ & = \Pr\left[c(\|\mathbf{S}_0^{(2)}\|_\infty + c_0\sqrt{\log p/n}) > \lambda^{(1)} \mid \{\hat{\mathbf{x}}_i^{(2)}\}_{i=1}^n\right] < \alpha_0. \end{aligned}$$

We can then conclude that  $\hat{\lambda}^{(1)} \leq c\{c_0 + 4\sqrt{2}C(2\sqrt{3}\sigma + r_n)\}\sqrt{\log p/n}$  on  $\mathcal{N} \cap \mathcal{M}_1$ . Thus, on  $\mathcal{N}$ , we have  $\hat{\lambda}^{(1)} \leq c\{c_0 + 4\sqrt{2}C(2\sqrt{3}\sigma + r_n)\}\sqrt{\log p/n}$  with probability at least  $1 - p \exp(-n)$ , which tends to one as  $n \rightarrow \infty$ . Then the proof is completed.  $\blacksquare$

## A.2 Proof of Theorem 2

**Proof** We focus on proving the properties of  $\hat{\beta}^{(1)}$ , because the properties of  $\hat{\beta}^{(2)}$  follow directly by swapping the roles of the two subsamples.

We write  $L^{(2)}(\boldsymbol{\gamma}) = Q^{(2)}(\boldsymbol{\gamma}) + \hat{\lambda}^{(1)}\|\boldsymbol{\beta}_0 + \boldsymbol{\gamma}\|_1$ , and then by definition,

$$\hat{\boldsymbol{\gamma}}^{(1)}(\hat{\lambda}^{(1)}) := \hat{\boldsymbol{\beta}}^{(1)}(\hat{\lambda}^{(1)}) - \boldsymbol{\beta}_0 = \arg \min_{\boldsymbol{\gamma} \in \mathbb{R}^p} L^{(2)}(\boldsymbol{\gamma}).$$

Recall that  $\Gamma = \{\boldsymbol{\gamma} \in \mathbb{R}^p : \|\boldsymbol{\gamma}_{\mathcal{B}^c}\|_1 \leq \bar{c}\|\boldsymbol{\gamma}_{\mathcal{B}}\|_1\}$  and  $\hat{\lambda}^{(1)} = c\{G_{(1)}^{-1}(1 - \alpha_0) + c_0\sqrt{\log p/n}\}$ . It follows from Lemma 1 that  $\Pr_*\{\hat{\boldsymbol{\gamma}}^{(1)}(\hat{\lambda}^{(1)}) \in \Gamma\} \geq 1 - \alpha_0 - \alpha_n$  and  $\hat{\lambda}^{(1)} \leq c_1\sqrt{\log p/n}$  on  $\mathcal{M}_1 \cap \mathcal{N}$ , where  $c_1 = c\{c_0 + 4\sqrt{2}C(2\sqrt{3}\sigma + r_n)\}$  and  $\alpha_n = 4p^{-(C_n^2-1)} \rightarrow 0$  as  $n \rightarrow \infty$ . Let  $h_n = \sqrt{q \log p/n}$ , and denote

$$\Gamma^* = \{\boldsymbol{\gamma} \in \mathbb{R}^p : \boldsymbol{\gamma} \in \Gamma, \|\boldsymbol{\gamma}\|_2 = \delta h_n\},$$

where  $\delta$  is a constant to be set large enough later. We note that

$$\begin{aligned} \Pr\left\{\|\hat{\boldsymbol{\gamma}}^{(1)}(\hat{\lambda}^{(1)})\|_2 \leq \delta h_n \mid \mathcal{N}\right\} & \geq 1 - \Pr\left\{\|\hat{\boldsymbol{\gamma}}^{(1)}(\hat{\lambda}^{(1)})\|_2 \geq \delta h_n, \hat{\boldsymbol{\gamma}}^{(1)}(\hat{\lambda}^{(1)}) \in \Gamma \mid \mathcal{N}\right\} \\ & \quad - \Pr\left\{\hat{\boldsymbol{\gamma}}^{(1)}(\hat{\lambda}^{(1)}) \notin \Gamma \mid \mathcal{N}\right\} \\ & \geq \Pr\left\{\inf_{\boldsymbol{\gamma} \in \Gamma, \|\boldsymbol{\gamma}\|_2 \geq \delta h_n} L^{(2)}(\boldsymbol{\gamma}) > L^{(2)}(\mathbf{0}) \mid \mathcal{N}\right\} - \alpha_0 - \alpha_n \\ & = \Pr\left\{\inf_{\boldsymbol{\gamma} \in \Gamma^*} L^{(2)}(\boldsymbol{\gamma}) > L^{(2)}(\mathbf{0}) \mid \mathcal{N}\right\} - \alpha_0 - \alpha_n, \end{aligned}$$

where the third inequality is due to the convexity of  $L^{(2)}(\cdot)$ . To see this, for an arbitrary point  $\mathbf{s}$  outside the  $\ell_2$  ball centered at  $\mathbf{0}$  with radius  $\delta h_n$ , we can write it as  $\mathbf{s} = l\mathbf{u}$ , where

$\mathbf{u}$  is a unit vector and  $l > \delta h_n$  is a positive constant. The convexity of  $L^{(2)}(\cdot)$  implies that  $(1 - l^{-1}\delta h_n)L^{(2)}(\mathbf{0}) + l^{-1}\delta h_n L^{(2)}(\mathbf{s}) \geq L^{(2)}(\delta h_n \mathbf{u})$ . Hence,

$$\begin{aligned} l^{-1}\delta h_n \{L^{(2)}(\mathbf{s}) - L^{(2)}(\mathbf{0})\} &\geq L^{(2)}(\delta h_n \mathbf{u}) - L^{(2)}(\mathbf{0}) \\ &\geq \inf_{\|\gamma\|_2 = \delta h_n} \{L^{(2)}(\gamma) - L^{(2)}(\mathbf{0})\}. \end{aligned}$$

Besides, on  $\mathcal{N} \cap \mathcal{M}_1$ ,

$$\begin{aligned} &\widehat{\lambda}^{(1)} \|\beta_0 + \gamma\|_1 - \|\beta_0\|_1 \\ &\leq \widehat{\lambda}^{(1)} \|\gamma\|_1 = \widehat{\lambda}^{(1)} (\|\gamma_B\|_1 + \|\gamma_{B^c}\|_1) \\ &\leq \widehat{\lambda}^{(1)} (1 + \bar{c}) \|\gamma_B\|_1 \leq \widehat{\lambda}^{(1)} (1 + \bar{c}) \sqrt{q} \|\gamma_B\|_2 \\ &\leq \widehat{\lambda}^{(1)} (1 + \bar{c}) \sqrt{q} \delta h_n \leq c_1 (1 + \bar{c}) \delta h_n^2, \end{aligned}$$

for any  $\gamma \in \Gamma^*$ . We denote  $Q^{(2)}(\gamma) = \{n(n-1)\}^{-1} \sum_{i \neq j} h_{i,j}(\gamma)$  with  $h_{i,j}(\gamma) = |(e_i^{(2)} - e_j^{(2)}) - (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \gamma|$ . We further define for simplicity that  $\zeta_{i,j}^{(2)} = e_i^{(2)} - e_j^{(2)}$ , and by the Knight's identity (Koenker, 2005), we have

$$\begin{aligned} Q^{(2)}(\gamma) - Q^{(2)}(\mathbf{0}) &= 2\{n(n-1)\}^{-1} \sum_{i \neq j} (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \gamma \{I(\zeta_{i,j}^{(2)} \leq 0) - 1/2\} \\ &\quad + 2\{n(n-1)\}^{-1} \sum_{i \neq j} \int_0^{(\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \gamma} \{I(\zeta_{i,j}^{(2)} \leq u) - I(\zeta_{i,j}^{(2)} \leq 0)\} du \\ &= 2Q_1^{(1)}(\gamma) + 2Q_2^{(1)}(\gamma). \end{aligned}$$

We further denote  $E_{\epsilon^{(2)}}(\cdot)$ ,  $\text{var}_{\epsilon^{(2)}}(\cdot)$  and  $\text{pr}_{\epsilon^{(2)}}(\cdot)$  as taking expectation, variance and probability with respect to  $\epsilon^{(2)} = (\epsilon_1^{(2)}, \dots, \epsilon_n^{(2)})^\top$  only.

We first note that for any  $\gamma \in \Gamma^*$ ,

$$\begin{aligned} E_{\epsilon^{(2)}} \{Q_1^{(1)}(\gamma)\} &= \{n(n-1)\}^{-1} \sum_{i \neq j} \left\{ (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \gamma \right. \\ &\quad \times [F^* \{\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})\} - 1/2] \Big\} \\ &\geq - \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} \left\{ [F^* \{\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})\} - 1/2] \right. \right. \\ &\quad \times (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)}) \Big\} \Big\|_\infty (1 + \bar{c}) \sqrt{q} \delta h_n. \end{aligned}$$

On the one hand, by Condition (C4) and the mean value theorem, there exists some  $\mu_{i,j}$  lying between 0 and  $\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})$ , such that,

$$\begin{aligned} &\left\| \{n(n-1)\}^{-1} \sum_{i \neq j} \left( [F^* \{\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})\} - 1/2] \right. \right. \\ &\quad \times \{\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})\} \Big) \Big\|_\infty \\ &= \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} \left[ f^*(\mu_{i,j}) \{\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})\} \right. \right. \\ &\quad \times \{\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})\} \Big] \Big\|_\infty \leq 4\kappa_u r_n a_n, \end{aligned}$$

under Condition (C1) and on  $\mathcal{N}$ , where  $\widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_i^{(2)}) = (\widehat{\Delta}_1^{(1)}(Z_i^{(2)}), \dots, \widehat{\Delta}_p^{(1)}(Z_i^{(2)}))^T$ .

On the other hand, denote  $H_{i,j,k} = [F^*\{\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})\} - 1/2](\widetilde{X}_{i,k}^{(2)} - \widetilde{X}_{j,k}^{(2)})$ . Then we claim that conditioning on  $\mathcal{D}_1$ ,  $H_{\pi(i),\pi(M_n+i),k}$  is zero-mean sub-Gaussian with parameter  $4\sqrt{2}\kappa_u a_n \sigma$  under Condition (C2) and on  $\mathcal{N}$ . To see this, recall that  $E_*|\widetilde{X}_{\pi(i),k}^{(2)} - \widetilde{X}_{\pi(M_n+i),k}^{(2)}|^t \leq (4\sigma^2)^{t/2} t\Gamma(t/2)$  by Proposition 2.5.2 (ii) in Vershynin (2018), and then

$$\begin{aligned} & E_* \left| [F^*\{\widehat{\Delta}_g^{(1)}(Z_{\pi(i)}^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_{\pi(M_n+i)}^{(2)})\} - 1/2](\widetilde{X}_{\pi(i),k}^{(2)} - \widetilde{X}_{\pi(M_n+i),k}^{(2)}) \right|^t \\ & \leq (4\kappa_u a_n \sigma)^t t\Gamma(t/2). \end{aligned}$$

Therefore,

$$\begin{aligned} & E_* \{\exp(\lambda^2 H_{\pi(i),\pi(M_n+i),k}^2)\} = E_* \left\{ 1 + \sum_{t=1}^{\infty} (\lambda^2 H_{\pi(i),\pi(M_n+i),k}^2)^t / t! \right\} \\ & = 1 + \sum_{t=1}^{\infty} \lambda^{2t} E_*(H_{\pi(i),\pi(M_n+i),k}^{2t}) / t! \leq 1 + 2 \sum_{t=1}^{\infty} (4\lambda\kappa_u a_n \sigma)^{2t} \\ & \leq \sum_{t=0}^{\infty} (4\sqrt{2}\lambda\kappa_u a_n \sigma)^{2t} = 1/(1 - 32\lambda^2 \kappa_u^2 a_n^2 \sigma^2). \end{aligned}$$

Similar to the proof of Vershynin (2018, Proposition 2.5.2),  $E_* \{\exp(\lambda^2 H_{\pi(i),\pi(M_n+i),k}^2)\} \leq \exp(64\lambda^2 \kappa_u^2 a_n^2 \sigma^2)$  for  $|\lambda| \leq (8\kappa_u a_n \sigma)^{-1}$ , since  $1/(1-x) \leq \exp(2x)$  for  $x \in [0, 1/2]$ . Furthermore, when  $\lambda \leq (8\kappa_u a_n \sigma)^{-1}$ ,

$$\begin{aligned} E_* \{\exp(\lambda H_{\pi(i),\pi(M_n+i),k})\} & \leq E_* \{\lambda H_{\pi(i),\pi(M_n+i),k} + \exp(\lambda^2 H_{\pi(i),\pi(M_n+i),k}^2)\} \\ & \leq \exp(64\lambda^2 \kappa_u^2 a_n^2 \sigma^2), \end{aligned}$$

as  $\exp(x) \leq x + \exp(x^2)$ . When  $\lambda > (8\kappa_u a_n \sigma)^{-1}$ ,

$$\begin{aligned} E_* \{\exp(\lambda H_{\pi(i),\pi(M_n+i),k})\} & \leq \exp(32\lambda^2 \kappa_u^2 a_n^2 \sigma^2) E_* [\exp\{H_{\pi(i),\pi(M_n+i),k}^2 / (128\kappa_u^2 a_n^2 \sigma^2)\}] \\ & \leq \exp(32\lambda^2 \kappa_u^2 a_n^2 \sigma^2) \exp(1/2) \leq \exp(64\lambda^2 \kappa_u^2 a_n^2 \sigma^2). \end{aligned}$$

That is,  $H_{\pi(i),\pi(M_n+i),k}$  is sub-Gaussian with parameter at most  $8\sqrt{2}\kappa_u a_n \sigma$ . Denote the event

$$\begin{aligned} \mathcal{M}_2 &= \left[ \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} \left\{ [F^*\{\widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})\} - 1/2] \right. \right. \right. \\ & \quad \left. \left. \left. \times (\widetilde{\mathbf{x}}_i^{(2)} - \widetilde{\mathbf{x}}_j^{(2)}) \right\} \right\|_{\infty} \leq 16\sqrt{2}C\kappa_u \sigma a_n \sqrt{\log p/n} \right], \end{aligned}$$

then  $\text{pr}_*(\mathcal{M}_2^c) \leq 2\exp\{-(C^2 - 1)\log p\}$  on  $\mathcal{N}$ .

To sum up,  $\inf_{\gamma \in \Gamma^*} E_{\epsilon^{(2)}} \{Q_1^{(1)}(\gamma)\} \geq -4(1 + \bar{c})\sqrt{q}\delta h_n \kappa_u a_n (r_n + 4\sqrt{2}C\sigma\sqrt{\log p/n})$  on  $\mathcal{M}_2 \cap \mathcal{N}$ .

As for  $Q_2^{(1)}(\gamma)$ , we have

$$E_{\epsilon^{(2)}} \{Q_2^{(1)}(\gamma)\} = \{n(n-1)\}^{-1} \sum_{i \neq j} \int_0^{(\widetilde{\mathbf{x}}_i^{(2)} - \widetilde{\mathbf{x}}_j^{(2)})^T \gamma} \{F^*(u + E_{i,j}^0) - F^*(E_{i,j}^0)\} du,$$

where  $E_{i,j}^0 = \widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})$ . Again by the mean value theorem, there exists some  $\mu_{i,j}$  lying between  $E_{i,j}^0$  and  $E_{i,j}^0 + (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \boldsymbol{\gamma}$ , such that

$$E_{\epsilon^{(2)}}\{Q_2^{(1)}(\boldsymbol{\gamma})\} = \{n(n-1)\}^{-1} \sum_{i \neq j} \int_0^{|(\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \boldsymbol{\gamma}|} f^*(\mu_{i,j}) u \, du,$$

which we claim is lower bounded by  $(\kappa_l/2)\{n(n-1)\}^{-1} \sum_{i \neq j} \{(\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \boldsymbol{\gamma}\}^2$  on  $\mathcal{N}$ , provided that  $2\{(2C\sigma L_n + r_n)(1 + \bar{c})\sqrt{q\delta}h_n + a_n\} \leq \delta_0$ . To see this, note that

$$\begin{aligned} \Pr\left(\max_i \|\widehat{\mathbf{x}}_i^{(2)}\|_\infty \geq t\right) &\leq \inf_{\lambda > 0} E\left\{\exp\left(\lambda \max_i \|\widehat{\mathbf{x}}_i^{(2)}\|_\infty\right)\right\} / \exp(\lambda t) \\ &= \inf_{\lambda > 0} E\left[\sum_{i=1}^n \sum_{k=1}^p \{\exp(\lambda \widetilde{X}_{i,k}^{(2)}) + \exp(-\lambda \widetilde{X}_{i,k}^{(2)})\}\right] / \exp(\lambda t) \\ &\leq 2 \inf_{\lambda > 0} \exp\left\{\lambda^2 \sigma^2 / 2 - \lambda t + \log(np)\right\} \leq 2 \exp\left\{-t^2 / (2\sigma^2) + 2 \log(p \vee n)\right\}. \end{aligned}$$

Therefore,

$$\max_{i \neq j} |(\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \boldsymbol{\gamma}| \leq 2(2C\sigma L_n + r_n)(1 + \bar{c})\sqrt{q\delta}h_n,$$

under Condition (C1), and on  $\mathcal{N}$  and  $\mathcal{M}_3 = \{\max_i \|\widetilde{\mathbf{x}}_i^{(2)}\|_\infty \leq 2C\sigma L_n\}$ , where  $\Pr(\mathcal{M}_3^c) \leq 2 \exp\{-2(C^2 - 1)L_n^2\}$  and  $L_n = \sqrt{\log(p \vee n)}$ . Consequently, we have  $f^*(\mu_{i,j}) \geq \kappa_l$  under Condition (C4), provided that  $2\{(2C\sigma L_n + r_n)(1 + \bar{c})\sqrt{q\delta}h_n + a_n\} \leq \delta_0$ .

Then, we observe that

$$\begin{aligned} \{(\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \boldsymbol{\gamma}\}^2 &= \{(\widetilde{\mathbf{x}}_i^{(2)} - \widetilde{\mathbf{x}}_j^{(2)})^\top \boldsymbol{\gamma}\}^2 + [\{\widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_j^{(2)})\}^\top \boldsymbol{\gamma}]^2 \\ &\quad - 2[\boldsymbol{\gamma}^\top (\widetilde{\mathbf{x}}_i^{(2)} - \widetilde{\mathbf{x}}_j^{(2)})\{\widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_j^{(2)})\}^\top \boldsymbol{\gamma}]. \end{aligned}$$

On the one hand, we have

$$\begin{aligned} &\left|2\{n(n-1)\}^{-1} \sum_{i \neq j} \boldsymbol{\gamma}^\top (\widetilde{\mathbf{x}}_i^{(2)} - \widetilde{\mathbf{x}}_j^{(2)})\{\widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_j^{(2)})\}^\top \boldsymbol{\gamma}\right| \\ &\leq 2 \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} (\widetilde{\mathbf{x}}_i^{(2)} - \widetilde{\mathbf{x}}_j^{(2)})\{\widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_j^{(2)})\}^\top \right\|_\infty \|\boldsymbol{\gamma}\|_1^2 \\ &\leq 64C(1 + \bar{c})^2 \sigma r_n q \delta^2 h_n^2 \sqrt{\log p / n}, \end{aligned}$$

on  $\mathcal{N}$  and the event

$$\begin{aligned} \mathcal{M}_4 &= \left[ \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} (\widetilde{\mathbf{x}}_i^{(2)} - \widetilde{\mathbf{x}}_j^{(2)})\{\widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_{\mathbf{x}}^{(1)}(Z_j^{(2)})\}^\top \right\|_\infty \right. \\ &\quad \left. \leq 32C\sigma r_n \sqrt{\log p / n} \right]. \end{aligned}$$

Following similar arguments to  $\mathcal{M}_2$ , we note that  $[(\widetilde{X}_{\pi(i),k}^{(2)} - \widetilde{X}_{\pi(M_n+i),k}^{(2)})\{\widehat{\Delta}_{k'}^{(1)}(Z_{\pi(i)}^{(2)}) - \widehat{\Delta}_{k'}^{(1)}(Z_{\pi(M_n+i)}^{(2)})\}]$ , conditioning on  $\mathcal{D}_1$ , is zero-mean sub-Gaussian with parameter  $8\sqrt{2}\sigma r_n$  under Condition (C2) and on  $\mathcal{N}$ . We then have  $\Pr_*(\mathcal{M}_4^c) \leq 2 \exp\{-2(C^2 - 1) \log p\}$  on  $\mathcal{N}$ .

On the other hand, we note that for any  $\gamma' \in \mathbb{R}^p$ ,

$$\begin{aligned} n^{-1} \sum_{i=1}^n (\gamma')^T \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T \gamma' &= (\gamma')^T \Sigma \gamma' + (\gamma')^T \left\{ n^{-1} \sum_{i=1}^n \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T - \Sigma \right\} \gamma' \\ &\geq \lambda_{\min}(\Sigma) \|\gamma'\|_2^2 - \left\| n^{-1} \sum_{i=1}^n \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T - \Sigma \right\|_{\infty} \|\gamma'\|_1^2, \end{aligned}$$

and similarly,

$$n^{-1} \sum_{i=1}^n (\gamma')^T \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T \gamma' \leq \lambda_{\max}(\Sigma) \|\gamma'\|_2^2 + \left\| n^{-1} \sum_{i=1}^n \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T - \Sigma \right\|_{\infty} \|\gamma'\|_1^2.$$

It is well-known that the product of sub-Gaussian is sub-exponential (Vershynin, 2018, Lemma 2.7.7). In particular, since for all  $k, k' = 1, \dots, p$  and  $t \geq 1$ ,

$$\begin{aligned} E|\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)} - E(\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)})|^t &= E|E(\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)} - \tilde{X}'_{i,k}{}^{(2)} \tilde{X}'_{i,k'}{}^{(2)} | \tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)})|^t \\ &\leq E(|\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)} - \tilde{X}'_{i,k}{}^{(2)} \tilde{X}'_{i,k'}{}^{(2)}|^t) = 2^t E\{(|\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)} - \tilde{X}'_{i,k}{}^{(2)} \tilde{X}'_{i,k'}{}^{(2)}|/2)^t\} \\ &\leq 2^t E(|\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)}|^t) \leq 2^t \{E(|\tilde{X}_{i,k}^{(2)}|^{2t}) E(|\tilde{X}_{i,k'}^{(2)}|^{2t})\}^{1/2}, \end{aligned}$$

where  $(\tilde{X}'_{i,k}{}^{(2)}, \tilde{X}'_{i,k'}{}^{(2)})$  is an independent copy of  $(\tilde{X}_{i,k}^{(2)}, \tilde{X}_{i,k'}^{(2)})$ , the first two inequalities follow from Jensen's inequality, and the last inequality follows from Cauchy-Schwarz inequality. Then following a similar argument to Honorio and Jaakkola (2014, Appendix B),  $\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)} - E(\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)})$  is sub-exponential with parameter  $(8\sqrt{2}\sigma^2, 8\sigma^2)$ , so that

$$\begin{aligned} &\text{pr}\left(\left\| (1/n) \sum_{i=1}^n \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T - \Sigma \right\|_{\infty} \geq t\right) \\ &\leq \sum_{k'=1}^p \sum_{k=1}^p \inf_{\lambda \in [0, (8\sigma^2/n)^{-1}]} E\left\{ \exp\left(\lambda \left| (1/n) \sum_{i=1}^n \tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)} - E(\tilde{X}_{i,k}^{(2)} \tilde{X}_{i,k'}^{(2)}) \right| \right) \right\} / \exp(\lambda t) \\ &\leq 2 \sum_{k'=1}^p \sum_{k=1}^p \exp\left\{ -nt^2 / \max(256\sigma^4, 16\sigma^2 t) \right\} \\ &= 2 \exp\left\{ -nt^2 / \max(256\sigma^4, 16\sigma^2 t) + 2 \log p \right\}. \end{aligned}$$

For notational simplicity, denote

$$\mathcal{L} := \mathcal{L}(t) = \left[ \left\| (1/n) \sum_{i=1}^n \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T - \Sigma \right\|_{\infty} \leq t \right].$$

This means when  $t = \lambda_{\min}(\Sigma) / \{4(1 + \bar{c})^2 q\}$ , we have

$$\begin{aligned} n^{-1} \sum_{i=1}^n (\gamma')^T \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T \gamma' &\geq \lambda_{\min}(\Sigma) \|\gamma'\|_2^2 - \lambda_{\min}(\Sigma) / \{4(1 + \bar{c})^2 q\} \|\gamma'\|_1^2, \\ n^{-1} \sum_{i=1}^n (\gamma')^T \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_i^{(2)})^T \gamma' &\leq \lambda_{\max}(\Sigma) \|\gamma'\|_2^2 + \lambda_{\min}(\Sigma) / \{4(1 + \bar{c})^2 q\} \|\gamma'\|_1^2, \end{aligned}$$

on  $\mathcal{L}$ , and

$$\begin{aligned} \text{pr}(\mathcal{L}^c) &\leq 2 \exp \left( -n[\lambda_{\min}(\mathbf{\Sigma})/\{4(1+\bar{c})^2q\}]^2 \right. \\ &\quad \left. / \max[256\sigma^4, 4\sigma^2\lambda_{\min}(\mathbf{\Sigma})/\{(1+\bar{c})^2q\}] + 2\log p \right). \end{aligned}$$

Because  $\gamma \in \Gamma^*$ , we have  $\|\gamma\|_1^2 \leq (1+\bar{c})^2q\|\gamma\|_2^2$ , and

$$\begin{aligned} n^{-1} \sum_{i=1}^n \{(\tilde{\mathbf{x}}_i^{(2)})^\top \gamma\}^2 &\geq (3/4)\lambda_{\min}(\mathbf{\Sigma})\|\gamma\|_2^2 = (3/4)\lambda_{\min}(\mathbf{\Sigma})\delta^2 h_n^2, \\ n^{-1} \sum_{i=1}^n \{(\tilde{\mathbf{x}}_i^{(2)})^\top \gamma\}^2 &\leq (5/4)\lambda_{\max}(\mathbf{\Sigma})\|\gamma\|_2^2 \leq 5\sigma^2\delta^2 h_n^2, \end{aligned}$$

where the last inequality holds because for any unit  $\alpha \in \mathbb{R}^p$ ,

$$\alpha^\top \mathbf{\Sigma} \alpha = E(\alpha^\top \tilde{\mathbf{x}} \tilde{\mathbf{x}}^\top \alpha) = E\{(\tilde{\mathbf{x}}^\top \alpha)^2\} \leq 4\sigma^2,$$

under Condition (C2) (see, for example, Vershynin, 2018, Proposition 2.5.2).

Besides, denote

$$\mathcal{M}_5 := \mathcal{M}_5(t) = \left[ \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_j^{(2)})^\top \right\|_\infty \leq t \right],$$

for simplicity. This means when  $t = \lambda_{\min}(\mathbf{\Sigma})/\{4(1+\bar{c})^2q\}$ , we have

$$\begin{aligned} &\left| 2\{n(n-1)\}^{-1} \sum_{i \neq j} \gamma^\top \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_j^{(2)})^\top \gamma \right| \\ &\leq 2 \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_j^{(2)})^\top \right\|_\infty \|\gamma\|_1^2 \leq (1/2)\lambda_{\min}(\mathbf{\Sigma})\delta^2 h_n^2, \end{aligned}$$

for any  $\gamma \in \Gamma^*$  and on  $\mathcal{M}_5$ . Because

$$\begin{aligned} &\text{pr} \left( \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} \tilde{\mathbf{x}}_i^{(2)} (\tilde{\mathbf{x}}_j^{(2)})^\top \right\|_\infty \geq t \right) \\ &\leq \sum_{k'=1}^p \sum_{k=1}^p \inf_{\lambda \in [0, (2\sqrt{2}\sigma^2/n)^{-1})} \frac{1}{n!} \\ &\quad \sum_{\pi} E \left\{ \exp \left( \lambda \left| M_n^{-1} \sum_{i=1}^{M_n} \tilde{X}_{\pi(i),k}^{(2)} \tilde{X}_{\pi(M_n+i),k'}^{(2)} \right| \right) \right\} / \exp(\lambda t) \\ &\leq 2 \sum_{k'=1}^p \sum_{k=1}^p \exp \left\{ -nt^2 / \max(8\sigma^4, 4\sqrt{2}\sigma^2 t) \right\} \\ &= 2 \exp \left\{ -nt^2 / \max(8\sigma^4, 4\sqrt{2}\sigma^2 t) + 2\log p \right\}, \end{aligned}$$

where the second inequality follows from the fact that the product of two independent sub-Gaussian random variables with parameter  $\sigma$  is sub-exponential with parameter  $(\sqrt{2}\sigma^2, \sqrt{2}\sigma^2)$  (Wainwright, 2019, Exercise 2.13), we then have

$$\begin{aligned} \text{pr}(\mathcal{M}_5^c) &\leq 2 \exp \left( -n[\lambda_{\min}(\mathbf{\Sigma})/\{4(1+\bar{c})^2q\}]^2 \right. \\ &\quad \left. / \max[8\sigma^4, \sqrt{2}\sigma^2\lambda_{\min}(\mathbf{\Sigma})/\{(1+\bar{c})^2q\}] + 2\log p \right). \end{aligned}$$

Putting the pieces together, we have  $\inf_{\gamma \in \Gamma^*} E_{\epsilon^{(2)}} \{Q_2^{(1)}(\gamma)\} \geq \kappa_l \delta^2 h_n^2 \{(1/2)\lambda_{\min}(\mathbf{\Sigma}) - 32C(1+\bar{c})^2 \sigma r_n q \sqrt{\log p/n}\}$  on  $\mathcal{L} \cap \mathcal{N} \cap \mathcal{M}_3 \cap \mathcal{M}_4 \cap \mathcal{M}_5$ .

Since we have known that

$$\begin{aligned} &\inf_{\gamma \in \Gamma^*} \{L^{(2)}(\gamma) - L^{(2)}(\mathbf{0})\} \\ &\geq \inf_{\gamma \in \Gamma^*} E_{\epsilon^{(2)}} \{Q^{(2)}(\gamma) - Q^{(2)}(\mathbf{0})\} \\ &\quad - \sup_{\gamma \in \Gamma^*} |Q^{(2)}(\gamma) - Q^{(2)}(\mathbf{0}) - E_{\epsilon^{(2)}} \{Q^{(2)}(\gamma) - Q^{(2)}(\mathbf{0})\}| \\ &\quad - \sup_{\gamma \in \Gamma^*} \hat{\lambda}^{(1)} \|\beta_0 + \gamma\|_1 - \|\beta_0\|_1, \end{aligned}$$

it remains to deal with the following term

$$W_n = \sup_{\gamma \in \Gamma^*} |Q^{(2)}(\gamma) - Q^{(2)}(\mathbf{0}) - E_{\epsilon^{(2)}} \{Q^{(2)}(\gamma) - Q^{(2)}(\mathbf{0})\}|.$$

By Markov's inequality, for any  $\lambda > 0$ ,

$$\begin{aligned} &\text{pr}_{\epsilon^{(2)}}(W_n \geq t) \\ &\leq E_{\epsilon^{(2)}} \left\{ \exp \left( \lambda \sup_{\gamma \in \Gamma^*} \left| \{n(n-1)\}^{-1} \sum_{i \neq j} \left[ h_{i,j}(\gamma) - h_{i,j}(\mathbf{0}) \right. \right. \right. \right. \\ &\quad \left. \left. \left. - E_{\epsilon^{(2)}} \{h_{i,j}(\gamma) - h_{i,j}(\mathbf{0})\} \right] \right| \right) \right\} / \exp(\lambda t) \\ &= E_{\epsilon^{(2)}} \left\{ \exp \left( \lambda \sup_{\gamma \in \Gamma^*} \left| \frac{1}{n!} \sum_{\pi} M_n^{-1} \sum_{i=1}^{M_n} \left[ h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \right. \right. \right. \\ &\quad \left. \left. \left. - E_{\epsilon^{(2)}} \{h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0})\} \right] \right| \right) \right\} / \exp(\lambda t) \\ &\leq \frac{1}{n!} \sum_{\pi} E_{\epsilon^{(2)}} \left\{ \exp \left( \lambda \sup_{\gamma \in \Gamma^*} \left| M_n^{-1} \sum_{i=1}^{M_n} \left[ h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \right. \right. \right. \\ &\quad \left. \left. \left. - E_{\epsilon^{(2)}} \{h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0})\} \right] \right| \right) \right\} / \exp(\lambda t). \end{aligned}$$

Write

$$\begin{aligned}\omega_+(\pi) &= \sup_{\gamma \in \Gamma^*} M_n^{-1} \sum_{i=1}^{M_n} \left[ h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \\ &\quad \left. - E_{\epsilon^{(2)}} \{ h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \} \right], \\ \omega_-(\pi) &= \sup_{\gamma \in \Gamma^*} -M_n^{-1} \sum_{i=1}^{M_n} \left[ h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \\ &\quad \left. - E_{\epsilon^{(2)}} \{ h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \} \right].\end{aligned}$$

Given the set  $\Gamma^*$ , we define its  $\varepsilon$ -covering  $\Gamma^*(\varepsilon)$  following Wainwright (2019, Chapter 5), and denote the covering number by  $|\Gamma^*(\varepsilon)| = N(\varepsilon)$ . We will show that the concentration inequalities for the supremum over  $\Gamma^*(\varepsilon)$  are independent of  $N(\varepsilon)$ . Therefore, by letting  $\varepsilon \rightarrow 0$ , we can extend the results to  $\Gamma^*$ . For simplicity, we will not distinguish between  $\Gamma^*$  and  $\Gamma^*(\varepsilon)$  in the following analysis. Using Massart (2000, Lemma 8), we have

$$\begin{aligned}& \lambda E_{\epsilon^{(2)}} [\omega_+(\pi) \exp\{\lambda \omega_+(\pi)\}] - E_{\epsilon^{(2)}} [\exp\{\lambda \omega_+(\pi)\}] \log E_{\epsilon^{(2)}} [\exp\{\lambda \omega_+(\pi)\}] \\ & \leq \sum_{i=1}^{M_n} E_{\epsilon^{(2)}} (\exp\{\lambda \omega_+(\pi)\} \phi[-\lambda \{\omega_+(\pi) - \omega_+^{\setminus i}(\pi)\}]),\end{aligned}$$

where  $\phi(u) = \exp(u) - u - 1$ , and  $\omega_+^{\setminus i}(\pi)$  is defined similarly to  $\omega_+(\pi)$ , but with  $h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) - E_{\epsilon^{(2)}} \{ h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \}$  replaced by  $-2|(\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^\top \gamma|$  for  $\omega_+^{\setminus i}(\pi)$  and  $2|(\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^\top \gamma|$  for  $\omega_-^{\setminus i}(\pi)$ , respectively. Since

$$|h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0})| \leq |(\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^\top \gamma|,$$

we have

$$\begin{aligned}& |h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) - E_{\epsilon^{(2)}} \{ h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \}| \\ & \leq 2|(\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^\top \gamma|.\end{aligned}$$

Therefore, for  $\gamma^*$  such that

$$\begin{aligned}\omega_+(\pi) &= M_n^{-1} \sum_{i=1}^{M_n} \left[ h_{\pi(i), \pi(M_n+i)}(\gamma^*) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \\ &\quad \left. - E_{\epsilon^{(2)}} \{ h_{\pi(i), \pi(M_n+i)}(\gamma^*) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \} \right],\end{aligned}$$

or

$$\begin{aligned}\omega_-(\pi) &= -M_n^{-1} \sum_{i=1}^{M_n} \left[ h_{\pi(i), \pi(M_n+i)}(\gamma^*) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \\ &\quad \left. - E_{\epsilon^{(2)}} \{ h_{\pi(i), \pi(M_n+i)}(\gamma^*) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \} \right],\end{aligned}$$

we have

$$0 \leq \omega.(\pi) - \omega.^{\setminus i}(\pi) \leq 4M_n^{-1} |(\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^T \boldsymbol{\gamma}^*|.$$

Note that  $\phi(u)/u^2$  is nondecreasing on  $\mathbb{R}$ ,  $\phi(u) \leq u^2/2$  for  $u \leq 0$ , which implies

$$\begin{aligned} & \sum_{i=1}^{M_n} E_{\epsilon^{(2)}}(\exp\{\lambda\omega.(\pi)\} \phi[-\lambda\{\omega.(\pi) - \omega.^{\setminus i}(\pi)\}]) \\ & \leq \sum_{i=1}^{M_n} 8M_n^{-2} \lambda^2 \{(\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^T \boldsymbol{\gamma}^*\}^2 E_{\epsilon^{(2)}}[\exp\{\lambda\omega.(\pi)\}], \end{aligned}$$

and

$$\begin{aligned} & \lambda E_{\epsilon^{(2)}}[\omega.(\pi) \exp\{\lambda\omega.(\pi)\}] - E_{\epsilon^{(2)}}[\exp\{\lambda\omega.(\pi)\}] \log E_{\epsilon^{(2)}}[\exp\{\lambda\omega.(\pi)\}] \\ & \leq \sup_{\boldsymbol{\gamma} \in \Gamma^*} M_n^{-2} \sum_{i=1}^{M_n} \{(\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^T \boldsymbol{\gamma}\}^2 8\lambda^2 E_{\epsilon^{(2)}}[\exp\{\lambda\omega.(\pi)\}]. \end{aligned}$$

On the one hand, on  $\mathcal{N}$ , we have that

$$M_n^{-1} \sum_{i=1}^{M_n} \{[\hat{\boldsymbol{\Delta}}_{\mathbf{x}}^{(1)}(Z_{\pi(i)}^{(2)}) - \hat{\boldsymbol{\Delta}}_{\mathbf{x}}^{(1)}(Z_{\pi(M_n+i)}^{(2)})]^T \boldsymbol{\gamma}\}^2 \leq 4(1 + \bar{c})^2 r_n^2 q \delta^2 h_n^2.$$

On the other hand, on  $\mathcal{L}$ ,

$$M_n^{-1} \sum_{i=1}^{M_n} \{(\tilde{\mathbf{x}}_{\pi(i)}^{(2)} - \tilde{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^T \boldsymbol{\gamma}\}^2 \leq (4/n) \sum_{i=1}^n \{(\tilde{\mathbf{x}}_i^{(2)})^T \boldsymbol{\gamma}\}^2 \leq 20\sigma^2 \delta^2 h_n^2.$$

Consequently, on  $\mathcal{L} \cap \mathcal{N}$ ,

$$\begin{aligned} & \lambda E_{\epsilon^{(2)}}[\omega.(\pi) \exp\{\lambda\omega.(\pi)\}] - E_{\epsilon^{(2)}}[\exp\{\lambda\omega.(\pi)\}] \log E_{\epsilon^{(2)}}[\exp\{\lambda\omega.(\pi)\}] \\ & \leq 128\{(1 + \bar{c})^2 r_n^2 q + 5\sigma^2\} (1/n) \delta^2 h_n^2 \lambda^2 E_{\epsilon^{(2)}}[\exp\{\lambda\omega.(\pi)\}]. \end{aligned}$$

This inequality holds for all positive  $\lambda$  and therefore yields the differential inequality

$$(1/\lambda)\{T'(\lambda)/T(\lambda)\} - (1/\lambda^2) \log T(\lambda) \leq 128\{(1 + \bar{c})^2 r_n^2 q + 5\sigma^2\} (1/n) \delta^2 h_n^2,$$

for  $T(\lambda) = E_{\epsilon^{(2)}}\{\exp[\lambda\{\omega.(\pi) - E_{\epsilon^{(2)}}\{\omega.(\pi)\}]\}\}$ . Setting  $H(\lambda) = (1/\lambda) \log T(\lambda)$ , the differential inequality becomes  $H'(\lambda) \leq 128\{(1 + \bar{c})^2 r_n^2 q + 5\sigma^2\} (1/n) \delta^2 h_n^2$ , which in turn implies since  $H(\lambda)$  tends to 0 as  $\lambda$  tends to 0,  $H(\lambda) \leq 128\{(1 + \bar{c})^2 r_n^2 q + 5\sigma^2\} (1/n) \delta^2 h_n^2 \lambda$ . Equivalently, on  $\mathcal{L} \cap \mathcal{N}$ ,

$$E_{\epsilon^{(2)}}\{\exp[\lambda\{\omega.(\pi) - E_{\epsilon^{(2)}}\{\omega.(\pi)\}]\}\} \leq \exp[128\{(1 + \bar{c})^2 r_n^2 q + 5\sigma^2\} (1/n) \delta^2 h_n^2 \lambda^2].$$

Putting the pieces together, on  $\mathcal{L} \cap \mathcal{N}$ ,

$$\begin{aligned} \text{pr}_{\epsilon^{(2)}}(W_n \geq t) & \leq \frac{2}{n!} \sum_{\pi} \exp\left(128\{(1 + \bar{c})^2 r_n^2 q + 5\sigma^2\} (1/n) \delta^2 h_n^2 \lambda^2 \right. \\ & \quad \left. - \lambda[t - E_{\epsilon^{(2)}}\{\omega.(\pi)\}]\right). \end{aligned}$$

It thus suffices to deal with  $E_{\epsilon^{(2)}}\{\omega.(\pi)\}$ :

$$\begin{aligned}
 E_{\epsilon^{(2)}}\{\omega.(\pi)\} &\leq 2E_{\epsilon^{(2)},\sigma}\left[\sup_{\gamma\in\Gamma^*}\left|M_n^{-1}\sum_{i=1}^{M_n}\sigma_i\{h_{\pi(i),\pi(M_n+i)}(\gamma)-h_{\pi(i),\pi(M_n+i)}(\mathbf{0})\}\right|\right] \\
 &\leq 4E_{\epsilon^{(2)},\sigma}\left\{\sup_{\gamma\in\Gamma^*}\left|M_n^{-1}\sum_{i=1}^{M_n}\sigma_i(\widehat{\mathbf{x}}_{\pi(i)}^{(2)}-\widehat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^\top\gamma\right|\right\} \\
 &\leq 4\sup_{\gamma\in\Gamma^*}\|\gamma\|_1E_{\epsilon^{(2)},\sigma}\left\{\max_k\left|M_n^{-1}\sum_{i=1}^{M_n}\sigma_i(\widehat{X}_{\pi(i),k}^{(2)}-\widehat{X}_{\pi(M_n+i),k}^{(2)})\right|\right\} \\
 &\leq 32\sup_{\gamma\in\Gamma^*}\|\gamma\|_1\left\{\max_k n^{-1}\sum_{i=1}^n(\widetilde{X}_{i,k}^{(2)})^2+r_n^2\right\}^{1/2}\sqrt{\log p/n},
 \end{aligned}$$

where  $\sigma = (\sigma_1, \dots, \sigma_{M_n})^\top$  are i.i.d. Rademacher random variables, the first inequality follows from the property of  $U$ -statistic, the symmetrization technique (Wainwright, 2019) and Ledoux and Talagrand (1991, Theorem 4.12), and the last equality follows from Wainwright (2019, Exercise 2.12). On  $\mathcal{M}_1 \cap \mathcal{N}$ ,  $E_{\epsilon^{(2)}}\{\omega.(\pi)\} \leq 32(1+\bar{c})(12\sigma^2+r_n^2)^{1/2}\delta h_n^2$ . As a result, on  $\mathcal{L} \cap \mathcal{N} \cap \mathcal{M}_1$ ,

$$\begin{aligned}
 \text{pr}_{\epsilon^{(2)}}(W_n \geq t) &\leq 2\inf_{\lambda>0}\exp\left[128\{(1+\bar{c})^2r_n^2q+5\sigma^2\}(1/n)\delta^2h_n^2\lambda^2\right. \\
 &\quad \left.-\lambda\{t-32(1+\bar{c})(12\sigma^2+r_n^2)^{1/2}\delta h_n^2\}\right] \\
 &\leq 2\exp\left\{-n\{t-32(1+\bar{c})(12\sigma^2+r_n^2)^{1/2}\delta h_n^2\}^2\right. \\
 &\quad \left./[512\{(1+\bar{c})^2r_n^2q+5\sigma^2\}\delta^2h_n^2]\right\}.
 \end{aligned}$$

That is, define

$$\begin{aligned}
 \mathcal{W} &= \left\{W_n \leq 16\sqrt{2}C\{(1+\bar{c})^2r_n^2q+5\sigma^2\}^{1/2}\delta h_n\sqrt{\log p/n}\right. \\
 &\quad \left.+32(1+\bar{c})(12\sigma^2+r_n^2)^{1/2}\delta h_n^2\right\},
 \end{aligned}$$

we have  $\text{pr}(\mathcal{W}^c \cap \mathcal{L} \cap \mathcal{N} \cap \mathcal{M}_1) \leq 2\exp(-C^2\log p)$ .

Putting the pieces together, we have

$$\begin{aligned}
 &\inf_{\gamma\in\Gamma^*}\{L^{(2)}(\gamma)-L^{(2)}(\mathbf{0})\} \\
 &> \kappa_l\delta^2h_n^2\{\lambda_{\min}(\Sigma)-64C(1+\bar{c})^2\sigma r_nq\sqrt{\log p/n}\} \\
 &\quad -8(1+\bar{c})\sqrt{q}\delta h_n\kappa_ua_n(r_n+4\sqrt{2}C\sigma\sqrt{\log p/n}) \\
 &\quad -64(1+\bar{c})C(3\sigma+r_n)\delta h_n^2-c_1(1+\bar{c})\delta h_n^2 \\
 &\geq \kappa_l\delta^2h_n^2\lambda_{\min}(\Sigma)-64C(1+\bar{c})\sigma\kappa_u\delta_0\delta h_n^2-c_0(1+\bar{c})\delta h_n^2 \\
 &\quad -64C(1+\bar{c})(3\sigma+r_n)\delta h_n^2-c_1(1+\bar{c})\delta h_n^2 \\
 &\geq \kappa_l\delta^2h_n^2\lambda_{\min}(\Sigma)-64C(1+\bar{c})\{4\sigma(\kappa_u\delta_0\vee 1)+r_n\}\delta h_n^2-2c_1(1+\bar{c})\delta h_n^2,
 \end{aligned}$$

on  $\mathcal{W} \cap \mathcal{N} \cap \mathcal{L} \cap \mathcal{M}_1 \cap \mathcal{M}_2 \cap \mathcal{M}_3 \cap \mathcal{M}_4 \cap \mathcal{M}_5$ , provided that  $2\{(2C\sigma L_n + r_n)(1 + \bar{c})\sqrt{q}\delta h_n + a_n\} \leq \delta_0$  and  $8\kappa_u r_n a_n \leq c_0 \sqrt{\log p/n}$ . By letting  $\delta = 2(1 + \bar{c})[c_1 + 32C\{4\sigma(\kappa_u \delta_0 \vee 1) + r_n\}]/\{\kappa_l \lambda_{\min}(\mathbf{\Sigma})\}$ ,  $\inf_{\gamma \in \Gamma^*} \{L^{(2)}(\gamma) - L^{(2)}(\mathbf{0})\} > 0$ . Therefore, for each  $l \in \{1, 2\}$ , we have

$$\begin{aligned} & \Pr[\|\hat{\gamma}^{(l)}(\hat{\lambda}^{(l)})\|_2 \leq \delta \sqrt{q \log p/n}] \\ & \geq 1 - \Pr\{(\mathcal{W} \cap \mathcal{L} \cap \mathcal{M}_1 \cap \mathcal{M}_2 \cap \mathcal{M}_3 \cap \mathcal{M}_4 \cap \mathcal{M}_5)^c \cap \mathcal{N}\} - \Pr(\mathcal{N}^c) - \alpha_0 - \alpha_n \\ & \geq 1 - 8 \exp\{-(C^2 - 1) \log p\} \\ & \quad - c_2 \exp(-c_3 n[\{\lambda_{\min}(\mathbf{\Sigma})/q\}^2/\sigma^4 \wedge 1] + \log p) - \rho_n - \alpha_0 - \alpha_n, \end{aligned}$$

for some universal constants  $c_2, c_3 > 0$ . Furthermore,

$$\|\hat{\gamma}^{(l)}(\hat{\lambda}^{(l)})\|_1 \leq (1 + \bar{c})\sqrt{q}\|\hat{\gamma}^{(l)}(\hat{\lambda}^{(l)})\|_2 \leq (1 + \bar{c})\delta q \sqrt{\log p/n},$$

with probability at least  $1 - 8 \exp\{-(C^2 - 1) \log p\} - c_2 \exp(-c_3 n[\{\lambda_{\min}(\mathbf{\Sigma})/q\}^2/\sigma^4 \wedge 1] + \log p) - \rho_n - \alpha_0 - \alpha_n$  as well. Then the proof is completed.  $\blacksquare$

**Remark 8** *The proof of Theorem 2 can be adapted to obtain similar results for the ANOPE (Fan et al., 2020b) subsampled version of DS Rank Lasso. Specifically, the loss function now becomes*

$$Q^{(2)}(\gamma) = (M_n |\mathcal{S}|)^{-1} \sum_{\pi \in \mathcal{S}} \sum_{i=1}^{M_n} |(e_{\pi(i)}^{(2)} - e_{\pi(M_n+i)}^{(2)}) - (\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^T \gamma|,$$

where  $\mathcal{S}$  is sampled with replacement from permutations of  $\{1, \dots, n\}$ , and  $|\mathcal{S}|$  denotes its cardinality. Recall that

$$W_n = \sup_{\gamma \in \Gamma^*} |Q^{(2)}(\gamma) - Q^{(2)}(\mathbf{0}) - E_{\epsilon^{(2)}}\{Q^{(2)}(\gamma) - Q^{(2)}(\mathbf{0})\}|,$$

we can similarly derive that

$$\begin{aligned} & \Pr_{\epsilon^{(2)}}(W_n \geq t) \\ & \leq E_{\epsilon^{(2)}} \left\{ \exp \left( \lambda \sup_{\gamma \in \Gamma^*} \left| (M_n |\mathcal{S}|)^{-1} \sum_{\pi \in \mathcal{S}} \sum_{i=1}^{M_n} [h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \right. \right. \\ & \quad \left. \left. \left. - E_{\epsilon^{(2)}}\{h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0})\}] \right| \right) \right\} / \exp(\lambda t) \\ & = E_{\epsilon^{(2)}} \left\{ \exp \left( \lambda \sup_{\gamma \in \Gamma^*} \left| \frac{1}{|\mathcal{S}|} \sum_{\pi \in \mathcal{S}} M_n^{-1} \sum_{i=1}^{M_n} [h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \right. \right. \\ & \quad \left. \left. \left. - E_{\epsilon^{(2)}}\{h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0})\}] \right| \right) \right\} / \exp(\lambda t) \\ & \leq \frac{1}{|\mathcal{S}|} \sum_{\pi \in \mathcal{S}} E_{\epsilon^{(2)}} \left\{ \exp \left( \lambda \sup_{\gamma \in \Gamma^*} \left| M_n^{-1} \sum_{i=1}^{M_n} [h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0}) \right. \right. \right. \\ & \quad \left. \left. \left. - E_{\epsilon^{(2)}}\{h_{\pi(i), \pi(M_n+i)}(\gamma) - h_{\pi(i), \pi(M_n+i)}(\mathbf{0})\}] \right| \right) \right\} / \exp(\lambda t), \end{aligned}$$

with  $h_{i,j}(\gamma) = |(e_i^{(2)} - e_j^{(2)}) - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma|$ . Then it is also sufficient to deal with the summand for any given  $\pi$ , and further details are thus omitted here.

### A.3 Proof of Lemma 4

**Proof** For consistency in notation, we next deal with  $\hat{\beta}^{d,(2)}$ . Now recall that

$$\mathbf{S}^{(2)}(\gamma) = -\{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) \text{sign}\{(e_i^{(2)} - e_j^{(2)}) - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\}.$$

Given any  $C' > 0$ , we note that uniformly for  $\gamma \in \mathbb{R}^p$  such that  $\|\gamma\|_0 \leq C'q$ , the number of  $(i, j)$  pairs such that  $e_i^{(2)} - e_j^{(2)} = (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma$  is at most of order  $O(q^2)$  with probability one, because the distribution of  $\epsilon_i^{(2)}$  has no point masses (see, for example, Koenker, 2005, Section 2.2). Consequently, with probability one,

$$\begin{aligned} & \sup_{\gamma \in \mathbb{R}^p: \|\gamma\|_0 \leq C'q} \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) I\{e_i^{(2)} - e_j^{(2)} = (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\} \right\|_\infty \\ &= O\{(q/n)^2 L_n\}, \end{aligned}$$

on  $\mathcal{M}_3$  defined in the proof of Theorem 2. We shall see this quantity is negligible. Thus, it is sufficient to consider the term  $2\{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) [I\{(e_i^{(2)} - e_j^{(2)}) \leq (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\} - 1/2]$ , which we also denote by  $\mathbf{S}^{(2)}(\gamma)$  for notational simplicity. Then the debiased estimator can be represented as

$$\begin{aligned} \hat{\beta}^{d,(2)} - \beta_0 &= -\{n(n-1)\}^{-1} \sum_{i \neq j} \hat{\Theta}^\top (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) \{I(\zeta_{i,j}^{(2)} < 0) - 1/2\} / \{2\hat{f}^*(0)\} \\ &\quad - \frac{\hat{f}^*(0)}{\hat{f}^*(0)} \hat{\Theta}^\top \mathbf{G}^{(2)}(\hat{\gamma}^{(2)}) - \frac{\hat{f}^*(0)}{\hat{f}^*(0)} \hat{\Theta}^\top E_{\epsilon^{(2)}} \{\boldsymbol{\nu}^{(2)}(\gamma)\} \Big|_{\gamma=\hat{\gamma}^{(2)}} + \hat{\gamma}^{(2)} \\ &= -\hat{\Theta}^\top \mathbf{S}^{(2)}(0) / \{4\hat{f}^*(0)\} - \frac{\hat{f}^*(0)}{\hat{f}^*(0)} \hat{\Theta}^\top \mathbf{G}^{(2)}(\hat{\gamma}^{(2)}) \\ &\quad - \frac{\hat{f}^*(0)}{\hat{f}^*(0)} \hat{\Theta}^\top E_{\epsilon^{(2)}} \{\boldsymbol{\nu}^{(2)}(\hat{\gamma}^{(2)})\} + \hat{\gamma}^{(2)}, \end{aligned}$$

where  $\mathbf{G}^{(2)}(\gamma)$  denotes the empirical process

$$\mathbf{G}^{(2)}(\gamma) = \boldsymbol{\nu}^{(2)}(\gamma) - E_{\epsilon^{(2)}} \{\boldsymbol{\nu}^{(2)}(\gamma)\},$$

and for  $\gamma \in \mathbb{R}^p$ ,

$$\boldsymbol{\nu}^{(2)}(\gamma) = \{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) / \{2\hat{f}^*(0)\} [I\{\zeta_{i,j}^{(2)} \leq (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma\} - I(\zeta_{i,j}^{(2)} \leq 0)].$$

Denote

$$\hat{P}_{i,j}(\gamma) = I\{\zeta_{0,i,j}^{(2)} \leq (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma + \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)})\},$$

and

$$P_{i,j}(\gamma) = F^*\{(\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma + \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)})\},$$

where  $\zeta_{0,i,j}^{(2)} = \epsilon_i^{(2)} - \epsilon_j^{(2)}$ . Then

$$\nu^{(2)}(\gamma) = \{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) / (2f^*(0)) \{ \hat{P}_{i,j}(\gamma) - \hat{P}_{i,j}(\mathbf{0}) \}.$$

We let  $\{\tilde{\gamma}_s\}_{s \in [N_\gamma]}$  be centers of the balls of radius  $\delta_n \xi_n$  that cover the set  $\mathcal{C}(\delta_n, q') = \{\gamma \in \mathbb{R}^p : \|\gamma\|_2 \leq \delta_n, \|\gamma\|_0 \leq q'\}$ . Such a cover can be constructed with  $N_\gamma \leq C_p^{q'} (3/\xi_n)^{q'}$  (see, for example, Wainwright, 2019, Example 5.8). Furthermore, let

$$\mathcal{B}(\tilde{\gamma}_s, r) = \{\gamma \in \mathbb{R}^p : \|\gamma - \tilde{\gamma}_s\|_2 \leq r \wedge \text{supp}(\gamma) \subseteq \text{supp}(\tilde{\gamma}_s)\}$$

be a ball of radius  $r$  centered at  $\tilde{\gamma}_s$  with points that have the same support as  $\tilde{\gamma}_s$  for  $s = 1, \dots, N_\gamma$ . We define  $G_k^{(2)}(\gamma)$  as the  $k$ -th component of  $\mathbf{G}^{(2)}(\gamma)$ , and in what follows we shall bound  $\sup_{\gamma \in \mathcal{C}(\delta_n, q')} \|\mathbf{G}^{(2)}(\gamma)\|_\infty = \max_k \sup_{\gamma \in \mathcal{C}(\delta_n, q')} |G_k^{(2)}(\gamma)|$  using a standard discretization argument. In particular, using the notation introduced above, we have the following decomposition

$$\begin{aligned} & \sup_{\gamma \in \mathcal{C}(\delta_n, q')} |G_k^{(2)}(\gamma)| \\ &= \max_{s \in [N_\gamma]} \sup_{\gamma \in \mathcal{B}(\tilde{\gamma}_s, \delta_n \xi_n)} |G_k^{(2)}(\gamma)| \\ &\leq \max_{s \in [N_\gamma]} |G_k^{(2)}(\tilde{\gamma}_s)| + \max_{s \in [N_\gamma]} \sup_{\gamma \in \mathcal{B}(\tilde{\gamma}_s, \delta_n \xi_n)} |G_k^{(2)}(\gamma) - G_k^{(2)}(\tilde{\gamma}_s)| = T_{1,k} + T_{2,k}. \end{aligned}$$

Notice that the term  $T_{1,k}$  arises from discretization of the set  $\mathcal{C}(\delta_n, q')$ . To control it, we will apply the union bound for  $s \in [N_\gamma]$ . The term  $T_{2,k}$  captures the deviation of the empirical process in a small neighborhood around the fixed center  $\tilde{\gamma}_s$ . We first bound  $T_{1,k}$ . Let  $a_{i,j,k} = (\hat{X}_{i,k}^{(2)} - \hat{X}_{j,k}^{(2)}) / (2f^*(0))$ , and

$$D_{i,j,k,s} = a_{i,j,k} [\{\hat{P}_{i,j}(\tilde{\gamma}_s) - P_{i,j}(\tilde{\gamma}_s)\} - \{\hat{P}_{i,j}(\mathbf{0}) - P_{i,j}(\mathbf{0})\}].$$

That is,

$$T_{1,k} = \max_{s \in [N_\gamma]} \left| \{n(n-1)\}^{-1} \sum_{i \neq j} D_{i,j,k,s} \right|.$$

We note that  $E_{\epsilon^{(2)}}(D_{i,j,k,s}) = 0$  and for some  $\eta_{i,j} \in [0, 1]$ ,

$$\begin{aligned} \text{var}_{\epsilon^{(2)}}(D_{i,j,k,s}) &= a_{i,j,k}^2 \left[ P_{i,j}(\tilde{\gamma}_s) - P_{i,j}^2(\tilde{\gamma}_s) + P_{i,j}(\mathbf{0}) - P_{i,j}^2(\mathbf{0}) \right. \\ &\quad \left. - 2\{P_{i,j}(\mathbf{0}) \wedge P_{i,j}(\tilde{\gamma}_s)\} + 2P_{i,j}(\tilde{\gamma}_s)P_{i,j}(\mathbf{0}) \right] \\ &\leq a_{i,j,k}^2 [P_{i,j}(\tilde{\gamma}_s) + P_{i,j}(\mathbf{0}) - 2\{P_{i,j}(\mathbf{0}) \wedge P_{i,j}(\tilde{\gamma}_s)\}] \\ &\leq a_{i,j,k}^2 |(\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \tilde{\gamma}_s| f^* \left\{ \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)}) \right. \\ &\quad \left. + \eta_{i,j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \tilde{\gamma}_s \right\} \\ &\leq \kappa_u a_{i,j,k}^2 |(\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \tilde{\gamma}_s| := \sigma_{i,j,k,s}^2, \end{aligned}$$

where the first inequality follows by dropping a negative term, the second inequality follows by the mean value theorem, and the last inequality is from Condition (C4). Furthermore,  $b_{k,s} := \max_{i,j} |D_{i,j,k,s}| \leq 2 \max_{i,j,k} |a_{i,j,k}|$ . Therefore, for a fixed  $s$ , we have that

$$\begin{aligned}
 & E_{\epsilon^{(2)}} \left[ \exp \{ \lambda (D_{\pi(i), \pi(M_n+i), k, s}) \} \right] \\
 &= 1 + \lambda^2 \sigma_{\pi(i), \pi(M_n+i), k, s}^2 / 2 + \sum_{t=3}^{\infty} \lambda^t E_{\epsilon^{(2)}} (D_{\pi(i), \pi(M_n+i), k, s}^t) / t! \\
 &\leq 1 + \lambda^2 \sigma_{\pi(i), \pi(M_n+i), k, s}^2 / 2 + \lambda^2 \sigma_{\pi(i), \pi(M_n+i), k, s}^2 / 2 \sum_{t=3}^{\infty} (|\lambda| b_{k,s})^{t-2} \\
 &= 1 + \frac{\lambda^2 \sigma_{\pi(i), \pi(M_n+i), k, s}^2 / 2}{1 - b_{k,s} |\lambda|} \leq \exp \left( \frac{\lambda^2 \sigma_{\pi(i), \pi(M_n+i), k, s}^2 / 2}{1 - b_{k,s} |\lambda|} \right),
 \end{aligned}$$

for  $|\lambda| \leq 1/b_{k,s}$ . Let  $C_0 > 0$  be a constant that may vary from line to line in the proof. On  $\mathcal{M}_1 \cap \mathcal{M}_3 \cap \mathcal{L} \cap \mathcal{N}$  and under Condition (C1), we have that

$$b_{k,s} \leq 2 \max_{i,j,k} |a_{i,j,k}| \leq (2/f^*(0))(2C\sigma L_n + r_n) \leq C_0 L_n,$$

and by the Cauchy-Schwarz inequality,

$$M_n^{-1} \sum_{i=1}^{M_n} \sigma_{\pi(i), \pi(M_n+i), k, s}^2 = M_n^{-1} \sum_{i=1}^{M_n} a_{\pi(i), \pi(M_n+i), k}^2 |(\hat{\mathbf{x}}_{\pi(i)}^{(2)} - \hat{\mathbf{x}}_{\pi(M_n+i)}^{(2)})^\top \tilde{\gamma}_s| \leq C_0 \delta_n L_n.$$

Then for arbitrary  $\varepsilon > 0$ , following similar arguments to Theorem 2 and Bernstein's inequality (Wainwright, 2019, Proposition 2.10), we have

$$\left| \{n(n-1)\}^{-1} \sum_{i \neq j} D_{i,j,k,s} \right| \leq C_0 [\sqrt{\{\log(2/\varepsilon)/n\} \delta_n L_n} \vee (L_n/n) \log(2/\varepsilon)],$$

with probability  $1 - \varepsilon$ . Therefore, using the union bound over  $s \in [N_\gamma]$ , with probability  $1 - \varepsilon$ , we have

$$T_{1,k} \leq C_0 [\{\log(2N_\gamma/\varepsilon) \delta_n L_n / n\}^{1/2} \vee (L_n/n) \log(2N_\gamma/\varepsilon)].$$

Let us now focus on bounding  $T_{2,k}$  term. For a fixed  $s$ , we denote

$$T_{2,k,s} = \sup_{\gamma \in \mathcal{B}(\tilde{\gamma}_s, \delta_n \xi_n)} |G_k^{(2)}(\gamma) - G_k^{(2)}(\tilde{\gamma}_s)|.$$

We will use the fact that  $I(\cdot < x)$  and  $\text{pr}(\cdot < x)$  are monotone function in  $x$ . Note that

$$\begin{aligned}
 & \max_{s \in [N_\gamma]} \max_{i,j} \sup_{\gamma \in \mathcal{B}(\tilde{\gamma}_s, \delta_n \xi_n)} |(\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top (\gamma - \tilde{\gamma}_s)| \\
 &\leq \delta_n \xi_n \sqrt{q'} \max_{i,j} \|\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}\|_\infty \leq C_0 \delta_n \xi_n \sqrt{q' L_n^2} := \tilde{L}_n,
 \end{aligned}$$

on  $\mathcal{M}_3 \cap \mathcal{N}$ , because

$$\begin{aligned} & |(\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top (\gamma - \widetilde{\gamma}_s)| \\ & \leq \|\gamma - \widetilde{\gamma}_s\|_2 \sqrt{\|\gamma - \widetilde{\gamma}_s\|_0} \max_{i,j} \|\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)}\|_\infty. \end{aligned}$$

We further denote  $E_{i,j}^s = (\widehat{\mathbf{x}}_i^{(2)} - \widehat{\mathbf{x}}_j^{(2)})^\top \widetilde{\gamma}_s + \widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})$  and  $E_{i,j}^0 = \widehat{\Delta}_g^{(1)}(Z_i^{(2)}) - \widehat{\Delta}_g^{(1)}(Z_j^{(2)})$ , and then

$$\begin{aligned} T_{2,k,s} & \leq \{n(n-1)\}^{-1} \sum_{i \neq j} \left[ |a_{i,j,k}| \left\{ I(\zeta_{0,i,j}^{(2)} \leq E_{i,j}^s + \widetilde{L}_n) - I(\zeta_{0,i,j}^{(2)} \leq E_{i,j}^0 + \widetilde{L}_n) \right. \right. \\ & \quad \left. \left. - I(\zeta_{0,i,j}^{(2)} \leq E_{i,j}^s - \widetilde{L}_n) + I(\zeta_{0,i,j}^{(2)} \leq E_{i,j}^0 - \widetilde{L}_n) \right\} \right. \\ & \quad \left. + |a_{i,j,k}| \left\{ -F^*(E_{i,j}^s - \widetilde{L}_n) + F^*(E_{i,j}^0 - \widetilde{L}_n) \right. \right. \\ & \quad \left. \left. + F^*(E_{i,j}^s + \widetilde{L}_n) - F^*(E_{i,j}^0 + \widetilde{L}_n) \right\} \right]. \end{aligned}$$

Furthermore,

$$\begin{aligned} T_{2,k,s} & \leq \{n(n-1)\}^{-1} \sum_{i \neq j} \left[ |a_{i,j,k}| \left\{ I(\zeta_{0,i,j}^{(2)} \leq E_{i,j}^s + \widetilde{L}_n) - I(\zeta_{0,i,j}^{(2)} \leq E_{i,j}^0 + \widetilde{L}_n) \right. \right. \\ & \quad \left. \left. - I(\zeta_{0,i,j}^{(2)} \leq E_{i,j}^s - \widetilde{L}_n) + I(\zeta_{0,i,j}^{(2)} \leq E_{i,j}^0 - \widetilde{L}_n) \right\} \right. \\ & \quad \left. + |a_{i,j,k}| \left\{ -F^*(E_{i,j}^s + \widetilde{L}_n) + F^*(E_{i,j}^0 + \widetilde{L}_n) \right. \right. \\ & \quad \left. \left. + F^*(E_{i,j}^s - \widetilde{L}_n) - F^*(E_{i,j}^0 - \widetilde{L}_n) \right\} \right] \\ & \quad + 2\{n(n-1)\}^{-1} \sum_{i \neq j} \left[ |a_{i,j,k}| \left\{ F^*(E_{i,j}^s + \widetilde{L}_n) - F^*(E_{i,j}^0 + \widetilde{L}_n) \right. \right. \\ & \quad \left. \left. - F^*(E_{i,j}^s - \widetilde{L}_n) + F^*(E_{i,j}^0 - \widetilde{L}_n) \right\} \right]. \end{aligned}$$

The first term in this equation can be bounded in a similar way to  $T_{1,k}$  by applying Bernstein's inequality and hence the details are omitted. For the second term, we have a bound  $C_0 L_n \widetilde{L}_n$ , since  $|a_{i,j,k}| \leq C_0 L_n$  and  $F^*(x + \widetilde{L}_n) - F^*(x - \widetilde{L}_n) \leq 2\kappa_u \widetilde{L}_n$  for all  $x \in \mathbb{R}$ . Therefore, with probability  $1 - \varepsilon$ ,

$$T_{2,k,s} \leq C_0 [\{L_n \widetilde{L}_n \log(2/\varepsilon)/n\}^{1/2} \vee L_n \{\log(2/\varepsilon)/n\} \vee L_n \widetilde{L}_n].$$

A bound on  $T_{2,k}$  now follows using a union bound over  $s \in [N_\gamma]$ . We choose  $\xi_n = n^{-1}$ , which gives  $N_\gamma \leq (3np)^{q'}$ . Then we let  $q' = C'q$ , and  $\delta_n = \delta h_n$  with  $\delta, h_n$  given in the proof of Theorem 2. Using a union bound over  $k \in \{1, \dots, p\}$ , we have, on  $\mathcal{M}_3 \cap \mathcal{L} \cap \mathcal{N}$ , that

$$\max_k \sup_{\gamma \in \mathcal{C}(\delta_n, q')} |G_k^{(2)}(\gamma)| \leq C_0 \left[ \sqrt{\{\log(np/\varepsilon)/n\} \delta_n q' L_n} \vee (L_n q'/n) \log(np/\varepsilon) \right],$$

with probability  $1 - 2\varepsilon$ . That is,

$$\sup_{\gamma \in \mathcal{C}(\delta_n, q')} \|\mathbf{G}^{(2)}(\gamma)\|_\infty = O_p[\lambda_{\min}^{-1/2}(\boldsymbol{\Sigma}) L_n^{3/2} \{\log(p)\}^{1/4} (q/n)^{3/4} \vee L_n^3 q/n].$$

Thus, we have

$$\sup_{\gamma \in \mathcal{C}(\delta_n, q')} \|\hat{\Theta}^T \mathbf{G}^{(2)}(\gamma)\|_\infty = O_p[\|\Theta\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{1/4} (q/n)^{3/4}],$$

since  $\|\Theta\|_1 L_n q \sqrt{\log p/n} = o(1)$  and  $\|\hat{\Theta} - \Theta\|_1 = o_p(1)$  by assumption, and since  $1/(4\sigma^2) \leq \lambda_{\min}^{-1}(\Sigma) = \lambda_{\max}(\Theta) = \|\Theta\|_2 \leq \|\Theta\|_1$ , where  $\|\Theta\|_1 = \max_k \|\theta_k\|_1$  and  $\|\Theta\|_2 = \lambda_{\max}^{1/2}(\Theta^T \Theta)$  are the  $\ell_1$  norm and spectral norm of  $\Theta$ , respectively.

Next, denote  $\tilde{\Gamma}^* = \{\gamma \in \mathbb{R}^p : \gamma \in \Gamma, \|\gamma\|_2 \leq \delta h_n\}$ . Using the mean value theorem, we note that there exists some  $\mu_{i,j}$  lying between  $\hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)})$  and  $(\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma + \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)})$ , such that

$$\begin{aligned} E_{\epsilon^{(2)}}\{\nu^{(2)}(\gamma)\} &= \{n(n-1)\}^{-1} \sum_{i \neq j} \left\{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) / (2f^*(0)) \right. \\ &\quad \times \left[ F^* \left\{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma + \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)}) \right\} \right. \\ &\quad \left. \left. - F^* \{ \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)}) \} \right] \right\} \\ &= \{n(n-1)\}^{-1} \sum_{i \neq j} \left[ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) / (2f^*(0)) \right. \\ &\quad \times \{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma \} f^* \{ \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)}) \} \Big] \\ &\quad + \{n(n-1)\}^{-1} \sum_{i \neq j} \left[ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) / (2f^*(0)) \right. \\ &\quad \times [f^*(\mu_{i,j}) - f^* \{ \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)}) \}] \{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma \} \Big] \\ &:= \mathbf{T}_3(\gamma) + \mathbf{T}_4(\gamma). \end{aligned}$$

For  $\mathbf{T}_4(\gamma)$ , by Condition (C1), we have known that  $\|\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}\|_\infty \leq C_0 L_n$  on  $\mathcal{M}_3 \cap \mathcal{N}$ . Furthermore, we have

$$\begin{aligned} &\sup_{\gamma \in \tilde{\Gamma}^*} \{n(n-1)\}^{-1} \sum_{i \neq j} [f^*(\mu_{i,j}) - f^* \{ \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)}) \}]^2 \\ &\leq L_1^2 \sup_{\gamma \in \tilde{\Gamma}^*} \{n(n-1)\}^{-1} \sum_{i \neq j} \{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma \}^2 \leq C_0 \{ \lambda_{\min}^{-2}(\Sigma) q \log(p) / n \}, \end{aligned}$$

as well as

$$\sup_{\gamma \in \tilde{\Gamma}^*} \{n(n-1)\}^{-1} \sum_{i \neq j} \{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma \}^2 \leq C_0 \{ \lambda_{\min}^{-2}(\Sigma) q \log(p) / n \},$$

on  $\mathcal{L} \cap \mathcal{N}$ , then

$$\sup_{\gamma \in \tilde{\Gamma}^*} \|\hat{\Theta}^T \mathbf{T}_4(\gamma)\|_\infty = O_p[\{\|\Theta\|_1 + o_p(1)\} \lambda_{\min}^{-2}(\Sigma) L_n q \log(p) / n] = O_p\{\|\Theta\|_1^3 L_n q \log(p) / n\}.$$

For  $\mathbf{T}_3(\gamma)$ , we further decompose it into

$$\begin{aligned}\mathbf{T}_{3,1}(\gamma) &= \{n(n-1)\}^{-1} \sum_{i \neq j} \left\{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) / (2f^*(0)) \{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma \} \right. \\ &\quad \left. \times [f^* \{ \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)}) \} - f^*(0)] \right\}, \\ \mathbf{T}_{3,2}(\gamma) &= \{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) \{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma \} / 2.\end{aligned}$$

Because on  $\mathcal{M}_1 \cap \mathcal{L} \cap \mathcal{N}$  and under Conditions (C1) and (C4),

$$\begin{aligned}\max_k \{n(n-1)\}^{-1} \sum_{i \neq j} \left\{ (\hat{X}_{i,k}^{(2)} - \hat{X}_{j,k}^{(2)})^2 \right. \\ \left. \times [f^* \{ \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)}) \} - f^*(0)]^2 \right\} \leq C_0 a_n^2,\end{aligned}$$

and

$$\sup_{\gamma \in \tilde{\Gamma}^*} \{n(n-1)\}^{-1} \sum_{i \neq j} \{ (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma \}^2 \leq C_0 \{ \lambda_{\min}^{-2}(\Sigma) q \log(p)/n \},$$

we have, by the Cauchy-Schwarz inequality, that

$$\sup_{\gamma \in \tilde{\Gamma}^*} \|\hat{\Theta}^\top \mathbf{T}_{3,1}(\gamma)\|_\infty = O_p[a_n \lambda_{\min}^{-1}(\Sigma) \|\Theta\|_1 \{q \log(p)/n\}^{1/2}] = O_p[a_n \|\Theta\|_1^2 \{q \log(p)/n\}^{1/2}].$$

On the other hand, recall  $\mathcal{L}(t)$  and  $\mathcal{M}_5(t)$  defined in the proof of Theorem 2, we note let  $t_1 = 16\sqrt{2}C\sigma^2\sqrt{\log p/n}(\sqrt{2}C\sqrt{\log p/n} \vee 1)$  and  $t_2 = 4C\sigma^2\sqrt{\log p/n}(2\sqrt{2}C\sqrt{\log p/n} \vee 1)$ , so that

$$\begin{aligned}& \sup_{\gamma \in \tilde{\Gamma}^*} \|\mathbf{T}_{3,2}(\gamma) - \Sigma\gamma\|_\infty \\ & \leq \sup_{\gamma \in \tilde{\Gamma}^*} \left\| \left[ \{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top - 2\Sigma \right] \gamma \right\|_\infty / 2 \\ & \leq C_0 \lambda_{\min}^{-1}(\Sigma) q \sqrt{\log p/n} \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)}) (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top - 2\Sigma \right\|_\infty \\ & \leq C_0 \lambda_{\min}^{-1}(\Sigma) \{q \log(p)/n\},\end{aligned}$$

on  $\mathcal{L}(t_1) \cap \mathcal{M}_4 \cap \mathcal{M}_5(t_2) \cap \mathcal{N}$ . This implies that

$$\begin{aligned}\sup_{\gamma \in \tilde{\Gamma}^*} \|\hat{\Theta}^\top \mathbf{T}_{3,2}(\gamma) - \hat{\Theta}^\top \Sigma\gamma\|_\infty &\leq \|\hat{\Theta}\|_1 \sup_{\gamma \in \tilde{\Gamma}^*} \|\mathbf{T}_{3,2}(\gamma) - \Sigma\gamma\|_\infty \\ &\leq O_p[\|\Theta\|_1 \lambda_{\min}^{-1}(\Sigma) \{q \log(p)/n\}].\end{aligned}$$

Moreover,

$$\begin{aligned}\sup_{\gamma \in \tilde{\Gamma}^*} \|\hat{\Theta}^\top \Sigma\gamma - \gamma\|_\infty &\leq \|\hat{\Theta} - \Theta\|_1 \|\Sigma\|_\infty \sup_{\gamma \in \tilde{\Gamma}^*} \|\gamma\|_1 \\ &= O_p[\|\hat{\Theta} - \Theta\|_1 \lambda_{\min}^{-1}(\Sigma) q \{\log(p)/n\}^{1/2}],\end{aligned}$$

and

$$\sup_{\gamma \in \tilde{\Gamma}^*} \left\| \frac{\hat{f}^*(0)}{f^*(0)} \gamma - \gamma \right\|_\infty = O_p \left\{ \left| \frac{\hat{f}^*(0)}{f^*(0)} - 1 \right| \lambda_{\min}^{-1}(\Sigma) \sqrt{q \log p/n} \right\}.$$

Thus,

$$\sup_{\gamma \in \tilde{\Gamma}^*} \|\hat{\Theta}^T \mathbf{T}_{3,2}(\gamma) - \gamma\|_\infty = O_p \left\{ \|\Theta\|_1 q \sqrt{\log p/n} (\|\hat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1 \sqrt{\log p/n}) \right\},$$

and

$$\begin{aligned} & \sup_{\gamma \in \tilde{\Gamma}^*} \left\| \hat{\Theta}^T E_{\epsilon^{(2)}} \{\nu^{(2)}(\gamma)\} - \frac{\hat{f}^*(0)}{f^*(0)} \gamma \right\|_\infty \\ &= O_p \left[ \|\Theta\|_1 \sqrt{q \log p/n} \left\{ \left| \frac{\hat{f}^*(0)}{f^*(0)} - 1 \right| \vee \sqrt{q} \|\hat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 L_n \sqrt{q \log p/n} \vee \|\Theta\|_1 a_n \right\} \right]. \end{aligned}$$

Next we derive the convergence rate of  $\hat{f}^*(0) = (\hat{f}^{*,(1)}(0) + \hat{f}^{*,(2)}(0))/2$  to  $f^*(0)$ . Denote

$$\begin{aligned} T_{5,1}(\gamma) &= \{n(n-1)\}^{-1} \sum_{i \neq j} \left( K[\{e_i^{(2)} - e_j^{(2)} - (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma\}/h] \right. \\ &\quad \left. - K\{(e_i^{(2)} - e_j^{(2)})/h\} \right) / h, \end{aligned}$$

and

$$T_{5,2} = \{n(n-1)\}^{-1} \sum_{i \neq j} K\{(e_i^{(2)} - e_j^{(2)})/h\}/h.$$

We first upper bound  $\sup_{\gamma \in \tilde{\Gamma}^*} |T_{5,1}(\gamma)|$ . On the one hand, following an argument analogous to that used for  $W_n$  in the proof of Theorem 2, we have that

$$\sup_{\gamma \in \tilde{\Gamma}^*} |T_{5,1}(\gamma) - E_{\epsilon^{(2)}}\{T_{5,1}(\gamma)\}| = O_p \{\lambda_{\min}^{-1}(\Sigma) q \log p / (nh^2)\},$$

under Condition (C5). On the other hand,

$$\begin{aligned} & \sup_{\gamma \in \tilde{\Gamma}^*} |E_{\epsilon^{(2)}}\{T_{5,1}(\gamma)\}| \\ & \leq \sup_{\gamma \in \tilde{\Gamma}^*} \left| \{n(n-1)\}^{-1} \sum_{i \neq j} h^{-1} \int_{-\infty}^{\infty} \left( K[\{u - E_{i,j}(\gamma)\}/h] \right. \right. \\ & \quad \left. \left. - K\{(u - E_{i,j}^0)/h\} \right) f^*(u) du \right| \\ & = \sup_{\gamma \in \tilde{\Gamma}^*} \left| \{n(n-1)\}^{-1} \sum_{i \neq j} \int_{-\infty}^{\infty} K(u) [f^*\{hu + E_{i,j}(\gamma)\} - f^*(hu + E_{i,j}^0)] du \right| \\ & \leq \sup_{\gamma \in \tilde{\Gamma}^*} L_1 \{n(n-1)\}^{-1} \sum_{i \neq j} |(\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^T \gamma| \\ & = O_p [\lambda_{\min}^{-1}(\Sigma) \{q \log(p)/n\}^{1/2}], \end{aligned}$$

on  $\mathcal{L} \cap \mathcal{N}$  and under Condition (C1), where  $E_{i,j}^0 = \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)})$  and  $E_{i,j}(\gamma) = (\hat{\mathbf{x}}_i^{(2)} - \hat{\mathbf{x}}_j^{(2)})^\top \gamma + \hat{\Delta}_g^{(1)}(Z_i^{(2)}) - \hat{\Delta}_g^{(1)}(Z_j^{(2)})$ . Thus, we have that

$$\sup_{\gamma \in \tilde{\Gamma}^*} |T_{5,1}(\gamma)| = O_p\{\lambda_{\min}^{-1}(\Sigma) \sqrt{q \log p/n} (\sqrt{q \log p/n}/h^2 \vee 1)\}.$$

Then we upper bound  $|T_{5,2} - f^*(0)|$ , which is further decomposed into  $|T_{5,2} - E_{\epsilon^{(2)}}(T_{5,2})| + |E_{\epsilon^{(2)}}(T_{5,2}) - f^*(0)|$ . By standard arguments for kernel density estimation, we have that

$$\begin{aligned} & |E_{\epsilon^{(2)}}(T_{5,2}) - f^*(0)| \\ & \leq |f^*(E_{i,j}^0) - f^*(0)| + \left| \int_{-\infty}^{+\infty} f^*(E_{i,j}^0 + hu) K(u) du - f^*(E_{i,j}^0) \right| \\ & = O_p(a_n) + \left| \int_{-\infty}^{+\infty} hu \{ (f^*)'(\mu_{i,j}(u)) - (f^*)'(E_{i,j}^0) \} K(u) du \right| \\ & \leq O_p(a_n) + L_3 h^2 \int_{-\infty}^{+\infty} u^2 K(u) du = O_p(a_n \vee h^2), \end{aligned}$$

under Conditions (C1), (C4) and (C5), where  $\mu_{i,j}(u)$  lies between  $E_{i,j}^0$  and  $E_{i,j}^0 + hu$ . As for the first term, we note that

$$\begin{aligned} \text{var}_{\epsilon^{(2)}}(T_{5,2}) & \leq \frac{1}{n!} \sum_{\pi} M_n^{-2} \sum_{i=1}^{M_n} E_{\epsilon^{(2)}}[K^2\{(\epsilon_{\pi(i)}^{(2)} - \epsilon_{\pi(M_n+i)}^{(2)} - E_{i,j}^0)/h\}]/h^2 \\ & \leq M_n^{-1} \kappa_u \int_{-\infty}^{\infty} K^2(u) du/h = O\{1/(nh)\}, \end{aligned}$$

by Jensen's inequality, so that  $|T_{5,2} - E_{\epsilon^{(2)}}(T_{5,2})| = O_p\{1/(nh)^{1/2}\}$  by Chebyshev's inequality.

Therefore, we have

$$\begin{aligned} |\hat{f}^*(0) - f^*(0)| & = O_p[\{\|\Theta\|_1 \sqrt{q \log p/n} (\sqrt{q \log p/n}/h^2 \vee 1) \vee 1/(nh)^{1/2} \vee h^2\} \vee a_n] \\ & = O_p(\|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n), \end{aligned}$$

if we let  $h \asymp \{\|\Theta\|_1 q \log(p)/n\}^{1/4} \vee n^{-1/5}$ , and

$$\left| \frac{\hat{f}^*(0)}{f^*(0)} - 1 \right| = O_p(\|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n),$$

because  $f^*(0) \geq \kappa_l$  under Condition (C4). Putting the pieces together, we have

$$\begin{aligned} & \sup_{\gamma \in \tilde{\Gamma}^*} \left\| \hat{\Theta}^\top E_{\epsilon^{(2)}}\{\nu^{(2)}(\gamma)\} - \frac{\hat{f}^*(0)}{f^*(0)} \gamma \right\|_{\infty} \\ & = O_p\left\{ \|\Theta\|_1 \sqrt{q \log p/n} \left( \|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee \sqrt{q} \|\hat{\Theta} - \Theta\|_1 \right. \right. \\ & \quad \left. \left. \vee \|\Theta\|_1^2 L_n \sqrt{q \log p/n} \vee \|\Theta\|_1 a_n \right) \right\} \\ & = O_p\left\{ \|\Theta\|_1 \sqrt{q \log p/n} \left( n^{-2/5} \vee \sqrt{q} \|\hat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 L_n \sqrt{q \log p/n} \vee \|\Theta\|_1 a_n \right) \right\}, \end{aligned}$$

and

$$\begin{aligned}
 & \|\widehat{\beta}^{d,(2)} - \beta_0 + \widehat{\Theta}^T \mathbf{S}^{(2)}(\mathbf{0}) / (4\widehat{f}^*(0))\|_\infty \\
 = & O_p \left\{ \|\Theta\|_1 \sqrt{q \log p/n} \left( n^{-2/5} \vee \sqrt{q} \|\widehat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 L_n \sqrt{q \log p/n} \vee \|\Theta\|_1 a_n \right) \right. \\
 & \left. \vee \|\Theta\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{1/4} (q/n)^{3/4} \right\} \\
 = & O_p \left\{ \|\Theta\|_1 \sqrt{q \log p/n} \left( \sqrt{q} \|\widehat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 L_n \sqrt{q \log p/n} \vee \|\Theta\|_1 a_n \right) \right. \\
 & \left. \vee \|\Theta\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{1/4} (q/n)^{3/4} \right\}.
 \end{aligned}$$

This is because, by Theorem 2 and the assumption that  $\|\widehat{\beta}^{(2)}\|_0 = O_p(q)$ , for any  $\varepsilon > 0$ , there exist sufficiently large constants  $C > 0$  and  $C' > 0$  (associated with  $\delta$  and  $q'$ ) such that  $\text{pr}\{\widehat{\gamma}^{(2)} \notin \mathcal{C}(\delta_n, q') \cap \bar{\Gamma}^*\} < \varepsilon$ .

Because  $\|\mathbf{S}^{(2)}(\mathbf{0})\|_\infty = O_p(\sqrt{\log p/n})$  by the proof of Lemma 1, and  $|1/\widehat{f}^*(0)| = O_p(1)$ , we have

$$\|(\widehat{\Theta} - \Theta)^T \mathbf{S}^{(2)}(\mathbf{0}) / (4\widehat{f}^*(0))\|_\infty = O_p(\|\widehat{\Theta} - \Theta\|_1 \sqrt{\log p/n}).$$

Moreover, because

$$\left| \frac{\widehat{f}^*(0)}{f^*(0)} - 1 \right| = O_p \left( \|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n \right),$$

we have

$$\begin{aligned}
 & \|\Theta^T \mathbf{S}^{(2)}(\mathbf{0}) [\{1/(4\widehat{f}^*(0))\} - \{1/(4f^*(0))\}]\|_\infty \\
 = & O_p \left\{ \left( \|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n \right) \|\Theta\|_1 \sqrt{\log p/n} \right\}.
 \end{aligned}$$

Thus,

$$\begin{aligned}
 & \|\widehat{\beta}^{d,(2)} - \beta_0 + \Theta^T \mathbf{S}^{(2)}(\mathbf{0}) / (4f^*(0))\|_\infty \\
 = & O_p \left\{ \|\Theta\|_1 \sqrt{q \log p/n} \left( \sqrt{q} \|\widehat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 L_n \sqrt{q \log p/n} \vee \|\Theta\|_1 a_n \right) \right. \\
 & \left. \vee \|\Theta\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{1/4} (q/n)^{3/4} \right\}.
 \end{aligned}$$

Then we prove the second part of Lemma 4. Denote  $\mathbf{T}_6 = \Theta^T (\mathbf{S}^{(2)}(\mathbf{0}) - \mathbf{S}_0^{(2)})$ . We note that

$$\|\mathbf{T}_6\|_\infty \leq \|\Theta\|_1 \|\mathbf{S}^{(2)}(\mathbf{0}) - \mathbf{S}_0^{(2)}\|_\infty = O_p \left\{ \|\Theta\|_1 \left( r_n a_n \vee a_n^{1/2} \sqrt{\log p/n} \vee \log p/n \right) \right\},$$

by using the proof of Lemma 1. It thus remains to deal with  $\Theta^T \mathbf{S}_0^{(2)}$ , and we claim it is asymptotically equivalent to

$$\mathbf{T}_7 = -\{n(n-1)\}^{-1} \sum_{i \neq j} \Theta^T (\widetilde{\mathbf{x}}_i^{(2)} - \widetilde{\mathbf{x}}_j^{(2)}) \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)}),$$

in the sense that

$$\|\Theta^T \mathbf{S}_0^{(2)} - \mathbf{T}_7\|_\infty = O_p(\|\Theta\|_1 r_n \sqrt{\log p/n}).$$

To see this, we note that

$$\begin{aligned} \|\Theta^T \mathbf{S}_0^{(2)} - \mathbf{T}_7\|_\infty &\leq \|\Theta\|_1 \left\| \{n(n-1)\}^{-1} \sum_{i \neq j} \{\hat{\Delta}_{\mathbf{x}}^{(1)}(Z_i^{(2)}) - \hat{\Delta}_{\mathbf{x}}^{(1)}(Z_j^{(2)})\} \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)}) \right\|_\infty \\ &= O_p(\|\Theta\|_1 r_n \sqrt{\log p/n}), \end{aligned}$$

following similar arguments to the proof of Lemma 1 (ii). Finally, let's consider

$$\mathbf{T}_0 = -(4/n) \sum_{i=1}^n \Theta^T \tilde{\mathbf{x}}_i^{(2)} \{F(\epsilon_i^{(2)}) - 1/2\}.$$

We first define a truncated version of  $\tilde{X}_{i,k}^{(2)}$ ,  $\tilde{X}_{i,k,T_n}^{(2)} = \tilde{X}_{i,k}^{(2)} I(|\tilde{X}_{i,k}^{(2)}| \leq T_n)$ , for some  $T_n > 0$ . Then, we define in addition

$$\begin{aligned} \mathbf{T}_{7,T_n} &= -\{n(n-1)\}^{-1} \sum_{i \neq j} \Theta^T (\tilde{\mathbf{x}}_{i,T_n}^{(2)} - \tilde{\mathbf{x}}_{j,T_n}^{(2)}) \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)}), \\ \mathbf{T}_{0,T_n} &= -(4/n) \sum_{i=1}^n \Theta^T \{\tilde{\mathbf{x}}_{i,T_n}^{(2)} - E(\tilde{\mathbf{x}}_{i,T_n}^{(2)})\} \{F(\epsilon_i^{(2)}) - 1/2\}, \end{aligned}$$

where  $\tilde{\mathbf{x}}_{i,T_n}^{(2)} = (\tilde{X}_{i,1,T_n}^{(2)}, \dots, \tilde{X}_{i,p,T_n}^{(2)})$ . Therefore,

$$\|\mathbf{T}_7 - \mathbf{T}_0\|_\infty \leq \|\mathbf{T}_7 - \mathbf{T}_{7,T_n}\|_\infty + \|\mathbf{T}_{7,T_n} - \mathbf{T}_{0,T_n}\|_\infty + \|\mathbf{T}_{0,T_n} - \mathbf{T}_0\|_\infty.$$

For  $\|\mathbf{T}_{7,T_n} - \mathbf{T}_{0,T_n}\|_\infty$ , we note that it is upper bounded by  $\|\Theta\|_1 \|\Delta_{T_n}\|_\infty$ , where

$$\begin{aligned} \Delta_{T_n} &= -\{n(n-1)\}^{-1} \\ &\quad \sum_{i \neq j} \left[ (\tilde{\mathbf{x}}_{i,T_n}^{(2)} - \tilde{\mathbf{x}}_{j,T_n}^{(2)}) \text{sign}(\epsilon_i^{(2)} - \epsilon_j^{(2)}) - 2\{\tilde{\mathbf{x}}_{i,T_n}^{(2)} - E(\tilde{\mathbf{x}}_{i,T_n}^{(2)})\} \{F(\epsilon_i^{(2)}) - 1/2\} \right. \\ &\quad \left. - 2\{\tilde{\mathbf{x}}_{j,T_n}^{(2)} - E(\tilde{\mathbf{x}}_{j,T_n}^{(2)})\} \{F(\epsilon_j^{(2)}) - 1/2\} \right]. \end{aligned}$$

Thus, it suffices to upper bound  $\Delta_{T_n}$  componentwise. For the  $k$ -th component,  $\Delta_{k,T_n}$ , we observe that it is a degenerate  $U$ -statistic. Therefore, according to Proposition 2.3 (d) in Arcones and Gine (1993),

$$\text{pr}\{(n-1)|\Delta_{k,T_n}| \geq t\} \leq c_4 \exp\{-c_5 t/(6T_n)\},$$

for some universal constants  $c_4, c_5$ . By the union bound, we have that

$$\text{pr}\{(n-1)\|\Delta_{T_n}\|_\infty \geq C_0 T_n \log p\} \leq c_4 \exp\{-(c_5 C_0/6 - 1) \log p\}.$$

By letting  $T_n = 2C\sigma L_n$ ,  $\|\Delta_{T_n}\|_\infty = O_p(\|\Theta\|_1 L_n \log p/n)$ . Moreover,  $\|\mathbf{T}_7 - \mathbf{T}_{7,T_n}\|_\infty = 0$ , and

$$\|\mathbf{T}_{0,T_n} - \mathbf{T}_0\|_\infty = 4 \left\| \Theta^T E(\tilde{\mathbf{x}}_i^{(2)} - \tilde{\mathbf{x}}_{i,T_n}^{(2)}) (1/n) \sum_{i=1}^n \{F(\epsilon_i^{(2)}) - 1/2\} \right\|_\infty,$$

on  $\mathcal{M}_3$ . For  $\|E(\tilde{\mathbf{x}}_i^{(2)} - \tilde{\mathbf{x}}_{i,T_n}^{(2)})\|_\infty$ , we observe that

$$\begin{aligned} |E(\tilde{X}_{i,k}^{(2)} - \tilde{X}_{i,k,T_n}^{(2)})| &\leq |E\{\tilde{X}_{i,k}^{(2)} I(|\tilde{X}_{i,k}^{(2)}| \geq T_n)\}| \\ &\leq 2\sqrt{2}\sigma \exp\{-T_n^2/(4\sigma^2)\} = 2\sqrt{2}\sigma \exp(-C^2 L_n^2), \end{aligned}$$

by Cauchy-Schwarz inequality,  $E\{(\tilde{X}_{i,k}^{(2)})^2\} \leq 4\sigma^2$ , and  $\text{pr}(|\tilde{X}_{i,k}^{(2)}| \geq T_n) \leq 2 \exp\{-T_n^2/(2\sigma^2)\}$ . Therefore,

$$\|\mathbf{T}_{0,T_n} - \mathbf{T}_0\|_\infty = O_p\{\|\boldsymbol{\Theta}\|_1 n^{-1/2} \|E(\tilde{\mathbf{x}}_i^{(2)} - \tilde{\mathbf{x}}_{i,T_n}^{(2)})\|_\infty\} = o_p(\|\boldsymbol{\Theta}\|_1 n^{-1}),$$

and

$$\|\mathbf{T}_7 - \mathbf{T}_0\|_\infty = O_p\left(\|\boldsymbol{\Theta}\|_1 L_n \log p/n\right).$$

To sum up,

$$\begin{aligned} &\left\| \boldsymbol{\Theta}^\top \mathbf{S}^{(2)}(\mathbf{0}) + (4/n) \sum_{i=1}^n \boldsymbol{\Theta}^\top \tilde{\mathbf{x}}_i^{(2)} \{F(\epsilon_i^{(2)}) - 1/2\} \right\|_\infty \\ &= O_p\left\{ \|\boldsymbol{\Theta}\|_1 \left( r_n a_n \vee a_n^{1/2} \sqrt{\log p/n} \vee L_n \log p/n \right) \right\}. \end{aligned}$$

Then the proof is completed. ■

#### A.4 Proof of Theorem 5

**Proof** Note that the assumptions  $\|\boldsymbol{\Theta}\|_1^3 L_n^3 \{\log(p)\}^{1/2} q^{3/2} n^{-1/2} = o(1)$  and  $\|\boldsymbol{\Theta}\|_1^2 (\sqrt{q \log p} \vee \log p) a_n = o(1)$  imply that

$$\begin{aligned} \|\boldsymbol{\Theta}\|_1^3 L_n q \log(p) n^{-1/2} &= o(1), \\ \|\boldsymbol{\Theta}\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{1/4} q^{3/4} n^{-1/4} &= o(1), \\ \|\boldsymbol{\Theta}\|_1 L_n \log(p) n^{-1/2} &= o(1), \\ \|\boldsymbol{\Theta}\|_1^2 \sqrt{q \log p} a_n &= o(1), \\ \|\boldsymbol{\Theta}\|_1 \sqrt{\log p} a_n^{1/2} &= o(1). \end{aligned}$$

Therefore, we have  $\|\mathbf{R}_{n,1}^{(l)}\|_\infty = o_p(n^{-1/2})$  and  $\|\mathbf{R}_{n,2}^{(l)}\|_\infty = o_p(n^{-1/2})$  by assumptions. Recall that  $\hat{\beta}_k^d = (\hat{\beta}_k^{d,(1)} + \hat{\beta}_k^{d,(2)})/2$ . We then have

$$\sqrt{N} \left[ (\hat{\beta}_k^d - \beta_{0,k}) - N^{-1} \sum_{i=1}^N \boldsymbol{\theta}_k^\top \tilde{\mathbf{x}}_i (F(\epsilon_i) - 1/2) / f^*(0) \right] = o_p(1),$$

by Lemma 4, where  $N = 2n$ . For  $\hat{\boldsymbol{\theta}}_k^\top \hat{\boldsymbol{\Sigma}}^{(l)} \hat{\boldsymbol{\theta}}_k$  with  $\hat{\boldsymbol{\Sigma}}^{(l)} = n^{-1} \sum_{i=1}^n (\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)})(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)})^\top$  and  $\bar{\mathbf{x}}^{(l)} = n^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(l)}$ , we note that

$$\|\hat{\boldsymbol{\theta}}^\top \hat{\boldsymbol{\Sigma}}^{(l)} \hat{\boldsymbol{\theta}} - \boldsymbol{\Theta}^\top \boldsymbol{\Sigma} \boldsymbol{\Theta}\|_\infty \leq \|\hat{\boldsymbol{\Sigma}}^{(l)} - \boldsymbol{\Sigma}\|_\infty \|\hat{\boldsymbol{\theta}}\|_1^2 + \|\hat{\boldsymbol{\theta}} \boldsymbol{\Sigma} \hat{\boldsymbol{\theta}} - \boldsymbol{\Theta}^\top \boldsymbol{\Sigma} \boldsymbol{\Theta}\|_\infty.$$

On the one hand, because  $\widehat{\Sigma}^{(l)} = n^{-2} \sum_{i,j=1}^n (\widehat{\mathbf{x}}_i^{(l)} - \widehat{\mathbf{x}}_j^{(l)})(\widehat{\mathbf{x}}_i^{(l)} - \widehat{\mathbf{x}}_j^{(l)})^\top / 2 = n^{-2} \sum_{i \neq j} (\widehat{\mathbf{x}}_i^{(l)} - \widehat{\mathbf{x}}_j^{(l)})(\widehat{\mathbf{x}}_i^{(l)} - \widehat{\mathbf{x}}_j^{(l)})^\top / 2$ , it follows from the proof of Lemma 4 that  $\|\widehat{\Sigma}^{(l)} - \Sigma\|_\infty = O_p\{(r_n \vee 1)\sqrt{\log p/n} + r_n^2\} = O_p(\sqrt{\log p/n})$ . Besides,  $\|\widehat{\Theta}\|_1^2 = O(\|\Theta\|_1^2 \vee \|\widehat{\Theta} - \Theta\|_1^2) = O_p(\|\Theta\|_1^2)$  by assumption. Therefore,

$$\|\widehat{\Sigma}^{(l)} - \Sigma\|_\infty \|\widehat{\Theta}\|_1^2 = O_p(\|\Theta\|_1^2 \sqrt{\log p/n}).$$

On the other hand, we have that

$$\begin{aligned} \|\widehat{\Theta}^\top \Sigma \widehat{\Theta} - \Theta^\top \Sigma \Theta\|_\infty &\leq \|(\widehat{\Theta} - \Theta)^\top \Sigma (\widehat{\Theta} - \Theta)\|_\infty + 2\|\Theta^\top \Sigma (\widehat{\Theta} - \Theta)\|_\infty \\ &= O(\|\widehat{\Theta} - \Theta\|_1^2) + O(\|\widehat{\Theta} - \Theta\|_1). \end{aligned}$$

Therefore,

$$\|\widehat{\Theta}^\top \widehat{\Sigma}^{(l)} \widehat{\Theta} - \Theta^\top \Sigma \Theta\|_\infty = O_p(\|\widehat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 \sqrt{\log p/n}).$$

Furthermore, by the proof of Lemma 4 and the assumptions, we have

$$\begin{aligned} |\widehat{\omega}_k^2 - \omega_k^2| &= O_p\left\{\|\widehat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 \sqrt{\log p/n} \right. \\ &\quad \left. \vee \|\Theta\|_1 \left(\|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n\right)\right\} \\ &= O_p\left\{\|\widehat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1 \left(\|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n\right)\right\} = o_p(1), \\ \left|\frac{\widehat{\omega}_k}{\omega_k} - 1\right| &= \left|\frac{\widehat{\omega}_k^2 - \omega_k^2}{\omega_k \left(\sqrt{\widehat{\omega}_k^2 - \omega_k^2} + \omega_k\right)}\right| = O_p\left(\left|\frac{\widehat{\omega}_k^2 - \omega_k^2}{\omega_k^2}\right|\right) \\ &= O_p\left(\|\widehat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 \sqrt{\log p/n} \vee \|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n\right) = o_p(1). \end{aligned}$$

Hence we have  $\widehat{\omega}_k/\omega_k \xrightarrow{p} 1$ , where  $\omega_k^2 = \Theta_{k,k}/[12\{f^*(0)\}^2]$  and

$$\widehat{\omega}_k^2 = \widehat{\theta}_k^\top \left( \frac{\widehat{\Sigma}^{(1)} + \widehat{\Sigma}^{(2)}}{2} \right) \widehat{\theta}_k / [12\{\widehat{f}^*(0)\}^2].$$

By Lindeberg-Feller CLT for triangular arrays (Resnick, 2019, Exercise 9.9.1),

$$N^{-1/2} \sum_{i=1}^N \theta_k^\top \widetilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\} / \{\omega_k f^*(0)\} \xrightarrow{d} \mathcal{N}(0, 1),$$

since  $E[\theta_k^\top \widetilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\} / \{N^{1/2} \omega_k f^*(0)\}] = 0$ ,

$$\sum_{i=1}^N E[\theta_k^\top \widetilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\} / \{N^{1/2} \omega_k f^*(0)\}]^2 = 1,$$

and for any  $\varepsilon > 0$ ,

$$\begin{aligned}
 & \sum_{i=1}^N E \left( [\boldsymbol{\theta}_k^T \tilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\} / \{N^{1/2} \omega_k f^*(0)\}]^2 \right. \\
 & \quad \left. \times I[|\boldsymbol{\theta}_k^T \tilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\} / \{N^{1/2} \omega_k f^*(0)\}| > \varepsilon] \right) \\
 &= NE \left( [\boldsymbol{\theta}_k^T \tilde{\mathbf{x}}_1 \{F(\epsilon_1) - 1/2\} / \{N^{1/2} \omega_k f^*(0)\}]^2 \right. \\
 & \quad \left. \times I[|\boldsymbol{\theta}_k^T \tilde{\mathbf{x}}_1 \{F(\epsilon_1) - 1/2\} / \{N^{1/2} \omega_k f^*(0)\}| > \varepsilon] \right) \\
 &\leq E[|\boldsymbol{\theta}_k^T \tilde{\mathbf{x}}_1 \{F(\epsilon_1) - 1/2\} / \{\omega_k f^*(0)\}|^3] / (\varepsilon N^{1/2}) \rightarrow 0.
 \end{aligned}$$

This is because  $E|\boldsymbol{\theta}_k^T / \|\boldsymbol{\theta}_k\|_1 \tilde{\mathbf{x}}_1|^3 = O(1)$  and  $\Theta_{k,k} \geq 1/E(\tilde{X}_{i,k}^2) \geq 1/(4\sigma^2)$  by sub-Gaussianity of  $\tilde{\mathbf{x}}_1$  (see, for exmaple, Vershynin, 2018, Proposition 2.5.2),  $f^*(0) \geq \kappa_l$  and  $\|\boldsymbol{\Theta}\|_1^3 N^{-1/2} = o(1)$ . It thus implies by Slutsky's theorem that

$$\sqrt{N}(\hat{\beta}_k^d - \beta_{0,k})/\hat{\omega}_k \xrightarrow{d} \mathcal{N}(0, 1),$$

since  $f^*(0)/\hat{f}^*(0) \xrightarrow{P} 1$ , and then the proof is finished.  $\blacksquare$

### A.5 Proof of Theorem 7

**Proof** The proof will follow similar arguments to Chernozhukov et al. (2013) and Zhang and Cheng (2017). To begin, we apply Chernozhukov et al. (2013, Corollary 2.1) to derive the following result under the assumption that  $\|\boldsymbol{\Theta}\|_1^6 L_n^{14}/n \leq Mn^{-\iota}$  for some  $M, \iota > 0$ :

$$\rho := \sup_{t \in \mathbb{R}} \left| \Pr \left( \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \xi_{i,k} \right| / \sqrt{N} \leq t \right) - \Pr \left( \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \tilde{\xi}_{i,k} \right| / \sqrt{N} \leq t \right) \right| \leq C'' N^{-C'}, \quad (7)$$

for any  $\mathcal{G} \subseteq \{1, \dots, p\}$  and some constants  $C', C'' > 0$ , where  $\xi_{i,k} = \boldsymbol{\theta}_k^T \tilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\} / f^*(0)$  and  $\tilde{\boldsymbol{\xi}}_i = (\tilde{\xi}_{i,1}, \dots, \tilde{\xi}_{i,p})^T$ ,  $i = 1, \dots, N$  is a sequence of zero-mean independent Gaussian vectors with  $E(\tilde{\boldsymbol{\xi}}_i \tilde{\boldsymbol{\xi}}_i^T) = \boldsymbol{\Theta} / [12\{f^*(0)\}^2]$ . Because

$$\begin{aligned}
 \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \xi_{i,k} \right| / \sqrt{N} &= \max_{k \in \mathcal{G}} \max \left( \sum_{i=1}^N \xi_{i,k}, -\sum_{i=1}^N \xi_{i,k} \right) / \sqrt{N}, \\
 \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \tilde{\xi}_{i,k} \right| / \sqrt{N} &= \max_{k \in \mathcal{G}} \max \left( \sum_{i=1}^N \tilde{\xi}_{i,k}, -\sum_{i=1}^N \tilde{\xi}_{i,k} \right) / \sqrt{N},
 \end{aligned}$$

it suffices to verify Condition (E.1) of Chernozhukov et al. (2013), that is, there exist some constants  $C_1, C_2 > 0$  and a sequence of constants  $B_n$  such that

$$C_1 \leq E(\xi_{i,k}^2) \leq C_2, \quad \max_{t=1,2} E(|\xi_{i,k}|^{2+t}/B_n^t) + E\{\exp(|\xi_{i,k}|/B_n)\} \leq 4,$$

uniformly over  $k \in \{1, \dots, p\}$ . We have that  $E(\xi_{i,k}^2) = \Theta_{k,k} / [12\{f^*(0)\}^2]$  is bounded away from zero and infinity, uniformly over  $k \in \{1, \dots, p\}$ , because  $1/(4\sigma^2) \leq \min_{k \in \{1, \dots, p\}} \Theta_{k,k} \leq$

$\max_{k \in \{1, \dots, p\}} \Theta_{k,k} < +\infty$  and  $\kappa_l \leq f^*(0) \leq \kappa_u$ . By the independence between  $\tilde{\mathbf{x}}_i$  and  $\epsilon_i$ , we have, for  $B_n = B\|\Theta\|_1^3$  with sufficiently large  $B > 0$ , that

$$\begin{aligned} \max_{t=1,2} E(|\xi_{i,k}|^{2+t}/B_n^t) &= \max_{t=1,2} \frac{E|\Theta_k^T \tilde{\mathbf{x}}_i|^{2+t} E[|\{F(\epsilon_i) - 1/2\}/f^*(0)|^{2+t}]}{B_n^t} \\ &\leq \max_{t=1,2} (t+2)\Gamma(t/2+1) \left( \frac{\|\Theta\|_1 \sigma}{\sqrt{2}\kappa_l} \right)^{t+2} / B_n^t < 1, \end{aligned}$$

and

$$\begin{aligned} E\{\exp(|\xi_{i,k}|/B_n)\} &= 1 + \sum_{t=1}^{\infty} \frac{E(|\xi_{i,k}|^t)}{B_n^t t!} = 1 + \sum_{t=1}^{\infty} \frac{E|\Theta_k^T \tilde{\mathbf{x}}_i|^t E[|\{F(\epsilon_i) - 1/2\}/f^*(0)|^t]}{B_n^t t!} \\ &\leq 1 + 2 \sum_{t=1}^{\infty} \left( \frac{\|\Theta\|_1 \sigma}{\sqrt{2} B_n \kappa_l} \right)^t \leq \frac{2}{1 - \frac{\|\Theta\|_1 \sigma}{\sqrt{2} B_n \kappa_l}} < 3, \end{aligned}$$

uniformly over  $k \in \{1, \dots, p\}$ , where we use the fact  $f^*(0) \geq \kappa_l$ ,  $E|\Theta_k^T \tilde{\mathbf{x}}_i|^t \leq (2\sigma^2)^{t/2} t\Gamma(t/2)$  under Condition (C2) (see, for example, Vershynin, 2018, Proposition 2.5.2) and  $t\Gamma(t/2) \leq 2t!$  for all  $t \geq 1$ . Therefore, for  $B_n = B\|\Theta\|_1^3$  with sufficiently large  $B > 0$ , we have  $\max_{t=1,2} E(|\xi_{i,k}|^{2+t}/B_n^t) + E\{\exp(|\xi_{i,k}|/B_n)\} \leq 4$ . Thus, (7) is proved, with constants  $C', C'' > 0$  depending only on  $C_1, C_2, M$ , and  $\nu$ .

Then, without loss of generality, we set  $\mathcal{G} = \{1, \dots, p\}$ . Define

$$T_{\mathcal{G}}^* = \max_{k \in \mathcal{G}} \sqrt{N} |\hat{\beta}_k^d - \beta_{0,k}|, \quad T_{1,\mathcal{G}} = \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \xi_{i,k} \right| / \sqrt{N}, \quad T_{0,\mathcal{G}} = \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \tilde{\xi}_{i,k} \right| / \sqrt{N}.$$

By Lemma 4, we have

$$\begin{aligned} &\left\| \sqrt{N}(\hat{\beta}^d - \beta_0) - \sum_{i=1}^N \xi_i / \sqrt{N} \right\|_{\infty} \\ &= O_p \left[ \|\Theta\|_1 \sqrt{q \log p} \left( \sqrt{q} \|\hat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 L_n \sqrt{q \log p / n} \vee \|\Theta\|_1 a_n \right) \right. \\ &\quad \left. \vee \|\Theta\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{1/4} q^{3/4} n^{-1/4} \vee \|\Theta\|_1 \left( \sqrt{n} r_n a_n \vee a_n^{1/2} \sqrt{\log p} \right) \right], \end{aligned}$$

where  $\xi_i = (\xi_{i,1}, \dots, \xi_{i,p})^T$ .

Note that the assumptions  $\|\Theta\|_1^3 L_n^3 \{\log(p)\}^{3/2} q^{3/2} n^{-1/2} = o(1)$ ,  $\|\Theta\|_1^2 a_n [\sqrt{q} \log p \vee \{\log(p)\}^2] = o(1)$  and the existence of  $\zeta_n \rightarrow +\infty$  such that  $\zeta_n [\|\Theta\|_1 q \log p \vee \{\log(p)\}^2] \|\hat{\Theta} - \Theta\|_1 = o_p(1)$ , imply that

$$\begin{aligned} \|\Theta\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{3/4} q^{3/4} n^{-1/4} &= o(1), \\ \|\Theta\|_1^3 L_n q \{\log(p)\}^{3/2} n^{-1/2} &= o(1), \\ \|\Theta\|_1 \sqrt{a_n} \log p &= o(1), \\ \|\Theta\|_1^2 a_n \sqrt{q} \log p &= o(1), \\ \zeta_n \|\Theta\|_1 q \log p \|\hat{\Theta} - \Theta\|_1 &= o_p(1). \end{aligned}$$

Then we can always choose  $\zeta_1, \zeta_2 > 0$  such that  $\zeta_1 \sqrt{1 \vee \log(p/\zeta_1)} = o(1)$ ,  $\zeta_2 = o(1)$  and

$$\Pr(|T_{\mathcal{G}}^* - T_{1,\mathcal{G}}| > \zeta_1) < \zeta_2.$$

For example, we can choose  $\zeta_1$  such that  $\zeta_1 = o(p^{-1})$  and

$$\begin{aligned} \zeta_1^2 &= \|\Theta\|_1 \sqrt{q} \left( \|\Theta\|_1^2 L_n \sqrt{q \log p/n} \vee \|\Theta\|_1 a_n \right) \vee (\zeta_n \sqrt{\log p})^{-2} \\ &\quad \vee \|\Theta\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{-1/4} q^{3/4} n^{-1/4} \vee \|\Theta\|_1 \left( \sqrt{n/\log p} r_n a_n \vee a_n^{1/2} \right), \end{aligned}$$

so that  $\zeta_1 \sqrt{1 \vee \log(p/\zeta_1)} = o(1)$  and  $|T_{\mathcal{G}}^* - T_{1,\mathcal{G}}| \leq \|\sqrt{N}(\hat{\beta}^d - \beta_0) - \sum_{i=1}^N \xi_i / \sqrt{N}\|_\infty = o_p(\zeta_1)$ .

Recall that

$$c_{\mathcal{G}}(\alpha) = \inf \left( t \in \mathbb{R} : \Pr \left[ W_{\mathcal{G}} \leq t \mid \{\mathbf{x}_i, Y_i, Z_i\}_{i=1}^N \right] \geq 1 - \alpha \right),$$

where

$$W_{\mathcal{G}} = \max_{k \in \mathcal{G}} \left| \sum_{l=1}^2 \sum_{i=1}^n \hat{\theta}_k^{\text{T}}(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)}) u_i^{(l)} \right| / [\sqrt{12N} \{\hat{f}^*(0)\}],$$

$\bar{\mathbf{x}}^{(l)} = n^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(l)}$ , and  $u_1^{(1)}, \dots, u_n^{(1)}, u_1^{(2)}, \dots, u_n^{(2)}$  are i.i.d.  $\mathcal{N}(0, 1)$  random variables independent of  $\{\mathbf{x}_i, Y_i, Z_i\}_{i=1}^N$ . Let  $c_{0,\mathcal{G}}(\alpha) = \inf \{t \in \mathbb{R} : \Pr(T_{0,\mathcal{G}} \leq t) \geq 1 - \alpha\}$ . Because  $C_1 \leq E(\tilde{\xi}_{i,k}^2) = E(\xi_{i,k}^2) \leq C_2$ , and

$$\begin{aligned} &\max_{k \in \mathcal{G}} \left| \sum_{l=1}^2 \sum_{i=1}^n \hat{\theta}_k^{\text{T}}(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)}) u_i^{(l)} \right| \\ &= \max_{k \in \mathcal{G}} \max \left\{ \sum_{l=1}^2 \sum_{i=1}^n \hat{\theta}_k^{\text{T}}(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)}) u_i^{(l)}, - \sum_{l=1}^2 \sum_{i=1}^n \hat{\theta}_k^{\text{T}}(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)}) u_i^{(l)} \right\}, \end{aligned}$$

we have, for some constant  $C_3 > 0$  depending only on  $C_1$  and  $C_2$ , that

$$\begin{aligned} &\sup_{t \in \mathbb{R}} \left| \Pr \left[ W_{\mathcal{G}} \leq t \mid \{\mathbf{x}_i, Y_i, Z_i\}_{i=1}^N \right] - \Pr(T_{0,\mathcal{G}} \leq t) \right| \\ &\leq C_3 \Delta^{1/3} \{1 \vee \log(p/\Delta)\}^{2/3}, \end{aligned}$$

by Chernozhukov et al. (2013, Lemma 3.1) or Chernozhukov et al. (2015, Theorem 2), where

$$\Delta = \left\| \hat{\Theta}^{\text{T}} \left( \frac{\hat{\Sigma}^{(1)} + \hat{\Sigma}^{(2)}}{2} \right) \hat{\Theta} / [12 \{\hat{f}^*(0)\}^2] - \Theta / [12 \{f^*(0)\}^2] \right\|_\infty.$$

By the proofs of Lemma 4 and Theorem 5, we have

$$\Delta = O_p \left( \|\hat{\Theta} - \Theta\|_1 \vee \|\Theta\|_1^2 \sqrt{\log p/n} \vee \|\Theta\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n \right),$$

under the assumption that  $\max_{k \in \{1, \dots, p\}} \Theta_{k,k} < +\infty$ .

Moreover, on the event  $\{\Delta \leq \zeta_3\}$ , we have

$$\sup_{t \in \mathbb{R}} \left| \Pr \left[ W_{\mathcal{G}} \leq t \mid \{\mathbf{x}_i, Y_i, Z_i\}_{i=1}^N \right] - \Pr(T_{0,\mathcal{G}} \leq t) \right| \leq \pi(\zeta_3),$$

where  $\pi(\zeta_3) = C_3 \zeta_3^{1/3} \{1 \vee \log(p/\zeta_3)\}^{2/3}$ . Then, for any  $\alpha \in (0, 1)$ , we have

$$\begin{aligned} & \Pr \left[ W_{\mathcal{G}} \leq c_{0,\mathcal{G}}(\alpha - \pi(\zeta_3)) \mid \{\mathbf{x}_i, Y_i, Z_i\}_{i=1}^N \right] \\ & \geq \Pr \left\{ T_{0,\mathcal{G}} \leq c_{0,\mathcal{G}}(\alpha - \pi(\zeta_3)) \right\} - \pi(\zeta_3) \\ & = 1 - \alpha + \pi(\zeta_3) - \pi(\zeta_3) = 1 - \alpha, \end{aligned}$$

because the distribution of  $T_{0,\mathcal{G}}$  has no point masses, and hence

$$\Pr \{c_{\mathcal{G}}(\alpha) \leq c_{0,\mathcal{G}}(\alpha - \pi(\zeta_3))\} \geq 1 - \Pr(\Delta > \zeta_3),$$

by the definition of  $c_{\mathcal{G}}(\alpha)$  and because the distribution of  $W_{\mathcal{G}}$  has no point masses, either. Similarly, we can derive that

$$\Pr \{c_{0,\mathcal{G}}(\alpha + \pi(\zeta_3)) \leq c_{\mathcal{G}}(\alpha)\} \geq 1 - \Pr(\Delta > \zeta_3).$$

Now, we have that

$$\begin{aligned} & |\Pr(T_{\mathcal{G}}^* > c_{\mathcal{G}}(\alpha)) - \Pr(T_{1,\mathcal{G}} > c_{0,\mathcal{G}}(\alpha))| \\ & \leq \Pr(\{T_{\mathcal{G}}^* > c_{\mathcal{G}}(\alpha)\} \ominus \{T_{1,\mathcal{G}} > c_{0,\mathcal{G}}(\alpha)\}, |T_{\mathcal{G}}^* - T_{1,\mathcal{G}}| \leq \zeta_1) + \zeta_2 \\ & \leq \Pr \left[ \{T_{\mathcal{G}}^* > c_{\mathcal{G}}(\alpha)\} \ominus \{T_{1,\mathcal{G}} > c_{0,\mathcal{G}}(\alpha)\}, |T_{\mathcal{G}}^* - T_{1,\mathcal{G}}| \leq \zeta_1, \right. \\ & \quad \left. c_{0,\mathcal{G}}(\alpha + \pi(\zeta_3)) \leq c_{\mathcal{G}}(\alpha) \leq c_{0,\mathcal{G}}(\alpha - \pi(\zeta_3)) \right] + \zeta_2 + 2\Pr(\Delta > \zeta_3) \\ & \leq \Pr \{c_{0,\mathcal{G}}(\alpha + \pi(\zeta_3)) - \zeta_1 \leq T_{1,\mathcal{G}} \leq c_{0,\mathcal{G}}(\alpha - \pi(\zeta_3)) + \zeta_1\} + \zeta_2 + 2\Pr(\Delta > \zeta_3) \\ & \leq \Pr \{c_{0,\mathcal{G}}(\alpha + \pi(\zeta_3)) - \zeta_1 \leq T_{0,\mathcal{G}} \leq c_{0,\mathcal{G}}(\alpha - \pi(\zeta_3)) + \zeta_1\} + \zeta_2 + 2\Pr(\Delta > \zeta_3) + 2\rho \\ & \leq 2\pi(\zeta_3) + \zeta_2 + 2\Pr(\Delta > \zeta_3) + 4C_4\zeta_1\sqrt{1 \vee \log(p/\zeta_1)} + 2\rho, \end{aligned}$$

for some constant  $C_4 > 0$  depending only on  $C_1$  and  $C_2$ , where  $A \ominus B = (A \cap B^c) \cup (B \cap A^c)$ , the first two inequalities hold because  $\Pr(|T_{\mathcal{G}}^* - T_{1,\mathcal{G}}| > \zeta_1) < \zeta_2$  and  $\Pr \{c_{0,\mathcal{G}}(\alpha + \pi(\zeta_3)) \leq c_{\mathcal{G}}(\alpha) \leq c_{0,\mathcal{G}}(\alpha - \pi(\zeta_3))\} \geq 1 - 2\Pr(\Delta > \zeta_3)$ , respectively, and the second-to-last inequality follows from the definition of  $\rho$ . The last inequality is because the anti-concentration inequality (see, for exmaple, Chernozhukov et al., 2013, Lemma 2.1 or Chernozhukov et al., 2015, Corollary 1), that is,

$$\begin{aligned} & \sup_{z \in \mathbb{R}} \Pr \left( \left| \max_k \sum_{i=1}^N \tilde{\xi}_{i,k} / \sqrt{N} - z \right| \leq \zeta_1 \right) \\ & = \sup_{z \in \mathbb{R}} \left| \Pr \left( \max_k \sum_{i=1}^N \tilde{\xi}_{i,k} / \sqrt{N} \leq z + \zeta_1 \right) - \Pr \left( \max_k \sum_{i=1}^N \tilde{\xi}_{i,k} / \sqrt{N} \leq z - \zeta_1 \right) \right| \\ & \leq C_4 \zeta_1 \sqrt{1 \vee \log(p/\zeta_1)}. \end{aligned}$$

To obtain the last inequality, we also use the fact that the distribution of  $T_{0,\mathcal{G}}$  has no point masses, and

$$\max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \tilde{\xi}_{i,k} \right| / \sqrt{N} = \max_{k \in \mathcal{G}} \max \left( \sum_{i=1}^N \tilde{\xi}_{i,k}, -\sum_{i=1}^N \tilde{\xi}_{i,k} \right) / \sqrt{N}.$$

Finally, we note that

$$\begin{aligned} \sup_{\alpha \in (0,1)} |\text{pr}(T_{\mathcal{G}}^* > c_{\mathcal{G}}(\alpha)) - \alpha| &\leq \sup_{\alpha \in (0,1)} |\text{pr}(T_{\mathcal{G}}^* > c_{\mathcal{G}}(\alpha)) - \text{pr}(T_{1,\mathcal{G}} > c_{0,\mathcal{G}}(\alpha))| \\ &\quad + \sup_{\alpha \in (0,1)} |\text{pr}(T_{1,\mathcal{G}} > c_{0,\mathcal{G}}(\alpha)) - \text{pr}(T_{0,\mathcal{G}} > c_{0,\mathcal{G}}(\alpha))| \\ &\leq 2\pi(\zeta_3) + \zeta_2 + 2\text{pr}(\Delta > \zeta_3) + C_4 \zeta_1 \sqrt{1 \vee \log(p/\zeta_1)} + 3\rho, \end{aligned}$$

because the distribution of  $T_{0,\mathcal{G}}$  has no point masses and  $\text{pr}(T_{0,\mathcal{G}} > c_{0,\mathcal{G}}(\alpha)) = \alpha$ . By the previous discussion,  $\zeta_1 \sqrt{1 \vee \log(p/\zeta_1)} = o(1)$ ,  $\zeta_2 = o(1)$  and  $\rho \leq C'' N^{-C'} = o(1)$ , we can conclude that

$$\sup_{\alpha \in (0,1)} |\text{pr}(T_{\mathcal{G}}^* > c_{\mathcal{G}}(\alpha)) - \alpha| = o(1),$$

if we choose  $\zeta_3$  such that

$$\begin{aligned} \zeta_3^2 &= \zeta_n^{-2} \{\log(p)\}^{-4} \vee a_n \{\log(p)\}^{-2} \vee \|\Theta\|_1 q^{1/2} \{\log(p)\}^{-3/2} n^{-1/2} \\ &\quad \vee \{\log(p)\}^{-2} n^{-2/5} \vee \|\Theta\|_1^2 \{\log(p)\}^{-3/2} n^{-1/2}. \end{aligned}$$

This is because under the assumptions  $\|\Theta\|_1^3 q L_n^3 \{\log(p)\}^{3/2} q^{1/2} n^{-1/2} = o(1)$ ,  $\|\Theta\|_1^2 a_n [\sqrt{q} \log p \vee \{\log(p)\}^2] = o(1)$ , and the existence of  $\zeta_n \rightarrow +\infty$  such that  $\zeta_n [\|\Theta\|_1 q \log p \vee \{\log(p)\}^2] \|\hat{\Theta} - \Theta\|_1 = o_p(1)$ , we have

$$\begin{aligned} \|\Theta\|_1^2 \{\log(p)\}^{5/2} n^{-1/2} &= o(1), \\ \|\Theta\|_1 \{\log(p)\}^{5/2} q^{1/2} n^{-1/2} &= o(1), \\ \{\log(p)\}^2 n^{-2/5} &= o(1), \\ a_n \{\log(p)\}^2 &= o(1), \\ \zeta_n \{\log(p)\}^2 \|\hat{\Theta} - \Theta\|_1 &= o_p(1), \end{aligned}$$

and then  $\pi(\zeta_3) = o(1)$  and  $\Delta = o_p(\zeta_3)$ .

Next, we can follow a similar argument to prove

$$\sup_{\alpha \in (0,1)} |\text{pr}(\bar{T}_{\mathcal{G}}^* > \bar{c}_{\mathcal{G}}(\alpha)) - \alpha| = o(1),$$

where  $\bar{T}_{\mathcal{G}}^* = \max_{k \in \mathcal{G}} \sqrt{N} |\hat{\beta}_k^d - \beta_{0,k}| / \hat{\omega}_k$ ,

$$\hat{\omega}_k^2 = \hat{\theta}_k^T \left( \frac{\hat{\Sigma}^{(1)} + \hat{\Sigma}^{(2)}}{2} \right) \hat{\theta}_k / [12 \{\hat{f}^*(0)\}^2],$$

and  $\widehat{\boldsymbol{\Sigma}}^{(l)} = n^{-1} \sum_{i=1}^n (\widehat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)}) (\widehat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)})^\top$ . We define accordingly

$$\bar{T}_{1,\mathcal{G}} = \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \xi_{i,k} / \omega_k \right| / \sqrt{N}, \quad \bar{T}_{0,\mathcal{G}} = \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \tilde{\xi}_{i,k} / \omega_k \right| / \sqrt{N},$$

with  $\omega_k^2 = \Theta_{k,k} / [12 \{f^*(0)\}^2]$ , and  $\bar{c}_{0,\mathcal{G}}(\alpha) = \inf \{t \in \mathbb{R} : \text{pr}(\bar{T}_{0,\mathcal{G}} \leq t) \geq 1 - \alpha\}$ . Then we can similarly derive: (i) Under the assumption that  $\|\boldsymbol{\Theta}\|_1^6 L_n^{14} / n \leq M n^{-\iota}$  for some  $M, \iota > 0$ , we have

$$\bar{\rho} := \sup_{t \in \mathbb{R}} \left| \text{pr} \left( \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \xi_{i,k} / \omega_k \right| / \sqrt{N} \leq t \right) - \text{pr} \left( \max_{k \in \mathcal{G}} \left| \sum_{i=1}^N \tilde{\xi}_{i,k} / \omega_k \right| / \sqrt{N} \leq t \right) \right| \leq \bar{C}'' N^{-\bar{C}'},$$

where  $\bar{C}', \bar{C}'' > 0$  depend only on  $\bar{C}_1 = \bar{C}_2 = 1$ ,  $M$ , and  $\iota$ . This is because  $E\{(\xi_{i,k}/\omega_k)^2\} = 1$ , and for  $\bar{B}_n = \bar{B} \|\boldsymbol{\Theta}\|_1^3$  with sufficiently large  $\bar{B} > 0$ ,

$$\begin{aligned} \max_{t=1,2} E(|\xi_{i,k}/\omega_k|^{2+t} / \bar{B}_n^t) &= \max_{t=1,2} \frac{E|\boldsymbol{\theta}_k^\top \tilde{\mathbf{x}}_i|^{2+t} E[|\sqrt{12/\Theta_{k,k}} \{F(\epsilon_i) - 1/2\}|^{2+t}]}{\bar{B}_n^t} \\ &\leq \max_{t=1,2} (t+2) \Gamma(t/2 + 1) (2\sqrt{6} \|\boldsymbol{\Theta}\|_1 \sigma^2)^{t+2} / \bar{B}_n^t < 1, \end{aligned}$$

and

$$\begin{aligned} E\{\exp(|\xi_{i,k}/\omega_k| / \bar{B}_n)\} &= 1 + \sum_{t=1}^{\infty} \frac{E(|\xi_{i,k}/\omega_k|^t)}{\bar{B}_n^t t!} \\ &= 1 + \sum_{t=1}^{\infty} \frac{E|\boldsymbol{\theta}_k^\top \tilde{\mathbf{x}}_i|^t E[|\sqrt{12/\Theta_{k,k}} \{F(\epsilon_i) - 1/2\}|^t]}{\bar{B}_n^t t!} \\ &\leq 1 + 2 \sum_{t=1}^{\infty} \left( \frac{2\sqrt{6} \|\boldsymbol{\Theta}\|_1 \sigma^2}{\bar{B}_n} \right)^t \leq \frac{2}{1 - \frac{2\sqrt{6} \|\boldsymbol{\Theta}\|_1 \sigma^2}{\bar{B}_n}} < 3, \end{aligned}$$

uniformly over  $k \in \{1, \dots, p\}$ .

(ii) Under the assumptions, there exist some  $\bar{\zeta}_1, \bar{\zeta}_2 > 0$  such that  $\bar{\zeta}_1 \{1 \vee \log(p/\bar{\zeta}_1)\}^{1/2} = o(1)$ ,  $\bar{\zeta}_2 = o(1)$  and

$$\text{pr}(|\bar{T}_{\mathcal{G}}^* - \bar{T}_{1,\mathcal{G}}| > \bar{\zeta}_1) < \bar{\zeta}_2.$$

This is because

$$\begin{aligned} &\max_{k \in \{1, \dots, p\}} \left| \sqrt{N} (\hat{\beta}_k^d - \beta_{0,k}) / \hat{\omega}_k - \sum_{i=1}^N \xi_{i,k} / (\sqrt{N} \hat{\omega}_k) \right| \\ &= O_p \left[ \|\boldsymbol{\Theta}\|_1 \sqrt{q \log p} \left( \sqrt{q} \|\hat{\boldsymbol{\Theta}} - \boldsymbol{\Theta}\|_1 \vee \|\boldsymbol{\Theta}\|_1^2 L_n \sqrt{q \log p / n} \vee \|\boldsymbol{\Theta}\|_1 a_n \right) \right. \\ &\quad \left. \vee \|\boldsymbol{\Theta}\|_1^{3/2} L_n^{3/2} \{\log(p)\}^{1/4} q^{3/4} n^{-1/4} \vee \|\boldsymbol{\Theta}\|_1 \left( \sqrt{n} r_n a_n \vee a_n^{1/2} \sqrt{\log p} \right) \right], \end{aligned}$$

and

$$\begin{aligned} & \max_{k \in \{1, \dots, p\}} \left| \sum_{i=1}^N \xi_{i,k} / \sqrt{N} (1/\hat{\omega}_k - 1/\omega_k) \right| \\ &= O_p \left\{ \|\boldsymbol{\Theta}\|_1 \sqrt{\log p} \left( \|\hat{\boldsymbol{\Theta}} - \boldsymbol{\Theta}\|_1 \vee \|\boldsymbol{\Theta}\|_1^2 \sqrt{\log p/n} \vee \|\boldsymbol{\Theta}\|_1 \sqrt{q \log p/n} \vee n^{-2/5} \vee a_n \right) \right\}, \end{aligned}$$

by the proofs of Lemma 4 and Theorem 5, the sub-Gaussianity of  $\boldsymbol{\theta}_k^T / \|\boldsymbol{\theta}_k\|_1 \tilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\} / f^*(0)$ , and the boundedness of  $\max_{k \in \{1, \dots, p\}} 1/\omega_k$ .

(iii)

$$\text{pr}\{\bar{c}_{0,\mathcal{G}}(\alpha + \bar{\pi}(\bar{\zeta}_3)) \leq \bar{c}_{\mathcal{G}}(\alpha) \leq \bar{c}_{0,\mathcal{G}}(\alpha - \bar{\pi}(\bar{\zeta}_3))\} \geq 1 - 2\text{pr}(\bar{\Delta} > \bar{\zeta}_3),$$

where  $\bar{\pi}(\bar{\zeta}_3) = \bar{C}_3 \bar{\zeta}_3^{1/3} \{1 \vee \log(p/\bar{\zeta}_3)\}^{2/3}$  for some  $\bar{C}_3 > 0$  depending only on  $\bar{C}_1 = \bar{C}_2 = 1$ ,

$$\bar{\Delta} = \max_{k,k'} \left| \frac{\hat{\boldsymbol{\theta}}_k^T \hat{\boldsymbol{\Sigma}} \hat{\boldsymbol{\theta}}_{k'}}{\sqrt{\hat{\boldsymbol{\theta}}_k^T \hat{\boldsymbol{\Sigma}} \hat{\boldsymbol{\theta}}_k} \sqrt{\hat{\boldsymbol{\theta}}_{k'}^T \hat{\boldsymbol{\Sigma}} \hat{\boldsymbol{\theta}}_{k'}}} - \frac{\Theta_{k,k'}}{\sqrt{\Theta_{k,k}} \sqrt{\Theta_{k',k'}}} \right|,$$

and  $\hat{\boldsymbol{\Sigma}} := (\hat{\boldsymbol{\Sigma}}^{(1)} + \hat{\boldsymbol{\Sigma}}^{(2)})/2$ . By the proof of Theorem 5, we have

$$\bar{\Delta} = O_p(\|\hat{\boldsymbol{\Theta}} - \boldsymbol{\Theta}\|_1 \vee \|\boldsymbol{\Theta}\|_1^2 \sqrt{\log p/n}).$$

Under the assumptions, there exist some  $\bar{\zeta}_3 > 0$  such that  $\bar{\pi}(\bar{\zeta}_3) = o(1)$  and  $\bar{\Delta} = o_p(\bar{\zeta}_3)$ .

(iv)

$$\sup_{z \in \mathbb{R}} \text{pr} \left( \left| \max_k \sum_{i=1}^N \tilde{\xi}_{i,k} / \omega_k / \sqrt{N} - z \right| \leq \bar{\zeta}_1 \right) \leq \bar{C}_4 \bar{\zeta}_1 \{1 \vee \log(p/\bar{\zeta}_1)\}^{1/2},$$

for some  $\bar{C}_4 > 0$  depending only on  $\bar{C}_1 = \bar{C}_2 = 1$ .

Finally, following similar arguments to  $\sup_{\alpha \in (0,1)} |\text{pr}(\bar{T}_{\mathcal{G}}^* > \bar{c}_{\mathcal{G}}(\alpha)) - \alpha|$ , we can derive that

$$\begin{aligned} & \sup_{\alpha \in (0,1)} |\text{pr}(\bar{T}_{\mathcal{G}}^* > \bar{c}_{\mathcal{G}}(\alpha)) - \alpha| \\ & \leq 2\bar{\pi}(\bar{\zeta}_3) + \bar{\zeta}_2 + 2\text{pr}(\bar{\Delta} > \bar{\zeta}_3) + \bar{C}_4 \bar{\zeta}_1 \{1 \vee \log(p/\bar{\zeta}_1)\}^{1/2} + 3\bar{\rho} = o(1). \end{aligned}$$

Note that the condition  $\max_{k \in \{1, \dots, p\}} \Theta_{k,k} < +\infty$  is not required for this part of the proof. This completes the proof.  $\blacksquare$

## A.6 The Assumptions on Random Errors

In general, the assumptions on random errors in Condition (C4) is not restrictive and can be satisfied by a wide range of density functions. Below, we provide two examples for the density of  $\epsilon_i - \epsilon_j$ .

- (i) Suppose  $\epsilon_i$  follows a normal distribution with mean 0 and variance  $\sigma^2$ . Then, the density function of  $\epsilon_i - \epsilon_j$  is  $f^*(t) = (4\pi\sigma^2)^{-1/2} \exp\{-t^2/(4\sigma^2)\}$ . We have  $\inf_{t \in [-\delta_0, \delta_0]} f^*(t) = (4\pi\sigma^2)^{-1/2} \exp\{-\delta_0^2/(4\sigma^2)\} =: \kappa_l$  and  $\sup_{t \in \mathbb{R}} f^*(t) \leq (4\pi\sigma^2)^{-1/2} =: \kappa_u$ . The first derivative of  $f^*(t)$  satisfies  $\sup_{t \in \mathbb{R}} |f^{*'}(t)| = (2\sqrt{2\pi}\sigma^2)^{-1} \exp(-1/2)$ . Hence,  $|f^*(t_1) - f^*(t_2)| \leq L_1 |t_1 - t_2|$  holds with  $L_1 = (2\sqrt{2\pi}\sigma^2)^{-1} \exp(-1/2)$ .
- (ii) Now, suppose  $\epsilon_i$  follows a  $t$  distribution with  $\nu$  degrees of freedom, location parameter  $\mu$ , and scale parameter  $\sigma$ . The density function of  $\epsilon_i$  is

$$f(t) := f(t; \mu, \sigma, \nu) = \frac{\Gamma(\frac{\nu+1}{2})}{\Gamma(\frac{\nu}{2}) \sigma \sqrt{\pi\nu}} \left(1 + \frac{1}{\nu} \left(\frac{t - \mu}{\sigma}\right)^2\right)^{-(\nu+1)/2},$$

and the density function of  $\epsilon_i - \epsilon_j$  is  $f^*(t) = \int_{-\infty}^{\infty} f(x - t)f(x) dx$ . By the symmetry and unimodality of the  $t$  density around  $\mu$ , for any  $0 \leq t \leq \delta_0$ ,

$$\begin{aligned} f^*(t) &= \int_{-\infty}^{\infty} f(x - t)f(x) dx \\ &\geq \int_{-\infty}^{\mu} f(x - t)f(x) dx + \int_{\mu+t}^{\infty} f(x - t)f(x) dx \\ &\geq \int_{-\infty}^{\mu} f^2(x - t) dx + \int_{\mu+t}^{\infty} f^2(x) dx \\ &\geq \int_{-\infty}^{\mu} f^2(x - \delta_0) dx + \int_{\mu+\delta_0}^{\infty} f^2(x) dx \\ &\geq 2 \int_{-\infty}^{\mu - \delta_0} f^2(x) dx. \end{aligned}$$

Since  $f^*(t)$  is symmetric about zero, it follows that

$$\begin{aligned} \inf_{t \in [-\delta_0, \delta_0]} f^*(t) &\geq 2 \int_{-\infty}^{\mu - \delta_0} f^2(x) dx \\ &= 2 \frac{\Gamma^2(\frac{\nu+1}{2}) \Gamma(\nu + \frac{1}{2})}{\Gamma^2(\frac{\nu}{2}) \Gamma(\nu + 1) \sigma \sqrt{\pi\nu}} \int_{-\infty}^{\mu - \delta_0} f\left(x; \mu, \sqrt{\nu/(2\nu + 1)}\sigma, 2\nu + 1\right) dx \\ &=: \kappa_l > 0. \end{aligned}$$

On the other hand, we have

$$\begin{aligned} \sup_{t \in \mathbb{R}} f^*(t) &\leq \frac{\Gamma(\frac{\nu+1}{2})}{\Gamma(\frac{\nu}{2}) \sigma \sqrt{\pi\nu}} \int_{-\infty}^{\infty} f(x) dx \\ &= \frac{\Gamma(\frac{\nu+1}{2})}{\Gamma(\frac{\nu}{2}) \sigma \sqrt{\pi\nu}} =: \kappa_u. \end{aligned}$$

Moreover, the first derivative of  $f^*(t)$  satisfies

$$\begin{aligned} \sup_{t \in \mathbb{R}} |f^{*'}(t)| &\leq \sup_{t \in \mathbb{R}} |f'(t)| \int_{-\infty}^{\infty} f(x) dx \\ &= \frac{\Gamma(\frac{\nu+1}{2}) (\nu + 1)}{\Gamma(\frac{\nu}{2}) \sigma \sqrt{2\pi\nu(\nu + 2)}} \left(1 + \frac{1}{\nu + 2}\right)^{-(\nu+3)/2} =: L_1. \end{aligned}$$

Thus,  $|f^*(t_1) - f^*(t_2)| \leq L_1|t_1 - t_2|$ .

### A.7 The Necessity of Data Splitting

**Illustration of Remark 3.** Consider the following toy example: Assume that  $Z_i$  is a continuous random variable, and both  $\{\mathbf{x}_i - \mathbf{m}_\mathbf{x}(Z_i)\}^\top \boldsymbol{\beta}_0$  and  $\epsilon_i$  are bounded. Define the nonparametric estimators as  $\hat{\mathbf{m}}_\mathbf{x}(z) = \mathbf{m}_\mathbf{x}(z)$ , and

$$\begin{aligned} \hat{m}_0(z) &= m_0(z) + r_N \sum_{i=1}^N \{Y_i - m_0(Z_i)\} I(Z_i = z) \\ &= m_0(z) + r_N \sum_{i=1}^N [\{\mathbf{x}_i - \mathbf{m}_\mathbf{x}(Z_i)\}^\top \boldsymbol{\beta}_0 + \epsilon_i] I(Z_i = z), \end{aligned}$$

so that  $\hat{m}_0(\cdot)$  converges uniformly to  $m_0(\cdot)$  at rate  $r_N$ . Direct calculation yields

$$\mathbf{S}(\mathbf{0}) = -\{N(N-1)\}^{-1} \sum_{i \neq j} (\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j) \text{sign}\{(1-r_N)(\epsilon_i - \epsilon_j) - r_N(\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)^\top \boldsymbol{\beta}_0\}.$$

Consequently,

$$\begin{aligned} E\{\mathbf{S}(\mathbf{0})\} &= -2E\{(\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)[F^*\{r_N/(1-r_N)(\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)^\top \boldsymbol{\beta}_0\} - 1/2]\} \\ &\asymp -2r_N/(1-r_N)E\{(\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)(\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)^\top \boldsymbol{\beta}_0\} = -4r_N/(1-r_N)\boldsymbol{\Sigma}\boldsymbol{\beta}_0. \end{aligned}$$

Therefore, if  $\sqrt{\log p/N} = o(r_N)$  and  $\|\boldsymbol{\Sigma}\boldsymbol{\beta}_0\|_\infty$  is bounded away from zero, it follows that  $\|\mathbf{S}(\mathbf{0})\|_\infty \asymp_p r_N$ . Consequently, using arguments similar to those in the proof of Theorem 2, selecting  $\lambda \asymp r_N$  yields an  $\ell_2$  convergence rate of  $O_p(\sqrt{q}r_N)$  for the rank Lasso estimator. This demonstrates that the near-oracle rate  $O_p(\sqrt{q \log p/N})$  cannot be guaranteed. In contrast, data-splitting removes the bias induced by overfitting and ensures  $E\{\mathbf{S}^{(l)}(\mathbf{0})\} = \mathbf{0}$ . Thus, the proposed DS rank Lasso achieves near-oracle rates under mild conditions on nonparametric estimators, even when they are overfitted.

**Illustration of Remark 6.** Consider the same toy example and the remainder term  $\mathbf{R}_{n,2}^{(l)}$  introduced in Lemma 4. The full-sample version of the  $k$ -th component of the remainder term is

$$\begin{aligned} R_{N,2,k} &= \boldsymbol{\theta}_k^\top \left[ -\{4N(N-1)\}^{-1} \sum_{i \neq j} (\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j) \text{sign}\{(1-r_N)(\epsilon_i - \epsilon_j) - r_N(\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)^\top \boldsymbol{\beta}_0\} \right. \\ &\quad \left. - N^{-1} \sum_{i=1}^N \tilde{\mathbf{x}}_i \{F(\epsilon_i) - 1/2\} \right] / f^*(0). \end{aligned}$$

Through direct calculations, we obtain

$$\begin{aligned} E(R_{N,2,k}) &= -E\{\boldsymbol{\theta}_k^\top (\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)[F^*\{r_N/(1-r_N)(\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)^\top \boldsymbol{\beta}_0\} - 1/2]\} / \{2f^*(0)\} \\ &\asymp -r_N/(1-r_N)E\{\boldsymbol{\theta}_k^\top (\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)(\tilde{\mathbf{x}}_i - \tilde{\mathbf{x}}_j)^\top \boldsymbol{\beta}_0\} / \{2f^*(0)\} \\ &= -r_N/(1-r_N)\boldsymbol{\beta}_{0,k}/f^*(0). \end{aligned}$$

| Error                 | Method                   | CP <sub>1</sub> | CP <sub>500</sub> | ACP <sub>B</sub> | ACP <sub>B<sup>c</sup></sub> | ARP <sub>B</sub> | RCP <sub>1</sub> | AIL <sub>B</sub> | AIL <sub>B<sup>c</sup></sub> |
|-----------------------|--------------------------|-----------------|-------------------|------------------|------------------------------|------------------|------------------|------------------|------------------------------|
| $\mathcal{N}(0, 2^2)$ | DSRLasso                 | 0.955           | 0.955             | 0.942            | 0.949                        | 1.000            | 0.955            | 0.483            | 0.476                        |
|                       | DSRLasso <sup>sub</sup>  | 0.960           | 0.950             | 0.943            | 0.949                        | 1.000            | 0.960            | 0.483            | 0.477                        |
|                       | FSORLasso                | 0.965           | 0.955             | 0.940            | 0.950                        | 1.000            | 0.965            | 0.462            | 0.468                        |
|                       | FSORLasso <sup>sub</sup> | 0.970           | 0.955             | 0.943            | 0.950                        | 1.000            | 0.970            | 0.462            | 0.468                        |
| MN                    | DSRLasso                 | 0.945           | 0.950             | 0.937            | 0.948                        | 1.000            | 0.945            | 0.329            | 0.326                        |
|                       | DSRLasso <sup>sub</sup>  | 0.955           | 0.950             | 0.942            | 0.948                        | 1.000            | 0.955            | 0.330            | 0.326                        |
|                       | FSORLasso                | 0.930           | 0.965             | 0.937            | 0.950                        | 1.000            | 0.930            | 0.314            | 0.318                        |
|                       | FSORLasso <sup>sub</sup> | 0.935           | 0.965             | 0.940            | 0.950                        | 1.000            | 0.935            | 0.314            | 0.318                        |

Table 5: A Comparison of the Proposed Methods and Full-sample Oracle Methods. Empirical Coverage Probabilities, Power and Lengths of Confidence Intervals Based on 200 Simulations.

Given that the optimal rate of convergence for nonparametric regression is  $r_N = N^{-1/2+\varepsilon}$  for some  $\varepsilon > 0$  (Stone, 1982), it follows that  $\sqrt{N}R_{N,2,k} \asymp_p N^\varepsilon \rightarrow \infty$  when  $\beta_{0,k} \neq 0$ . On the other hand, data splitting eliminates the bias induced by overfitting and ensures  $E(\mathbf{R}_{n,2}^{(l)}) = \mathbf{0}$ . Consequently, the remainder terms in  $\sqrt{N}(\hat{\beta}_k^d - \beta_{0,k})$  can converge to zero even when overfitting occurs in the nonparametric estimators.

## Appendix B. Additional Numerical Results

In this section, we provide additional simulation results to support our proposed method and complement the simulations presented in the main article. The data generation process is the same as that in Example 2 of Section 4.1. Unless otherwise stated, we fix  $N = 450$  and consider two error distributions: a light-tailed normal distribution  $\epsilon_i \sim \mathcal{N}(0, 2^2)$  and a heavy-tailed mixture normal distribution  $\epsilon_i \sim 0.95\mathcal{N}(0, 1) + 0.05\mathcal{N}(0, 10^2)$ .

### B.1 Comparison with Full-sample Oracle Methods

As established in Theorem 5 and the subsequent discussions in Section 3.1, the proposed debiased data-splitting (DS) rank Lasso estimator does not suffer from asymptotic efficiency loss due to data splitting. Specifically, it achieves the same asymptotic variance as the debiased full-sample (FS) oracle rank Lasso estimator introduced in Section 3.1. Table 5 compares the finite-sample performance of the inference procedure based on the proposed debiased DS rank Lasso estimator with that based on the debiased FS oracle rank Lasso estimator. The evaluation metrics are the same as those in Example 2, and the results are based on 200 simulations. We denote the debiased full-sample oracle rank Lasso estimator and its ANOPE (Fan et al., 2020b) subsampled version as FSORLasso and FSORLasso<sup>sub</sup>, respectively. Overall, the finite-sample performance of the two methods is roughly comparable. In particular, the proposed debiased DS rank Lasso yields slightly wider but still comparable confidence intervals relative to the debiased FS oracle rank Lasso. This indicates that data splitting introduces only minor efficiency loss in finite samples when the sample size is moderately large. Such efficiency loss diminishes asymptotically, as supported by our theoretical analysis.

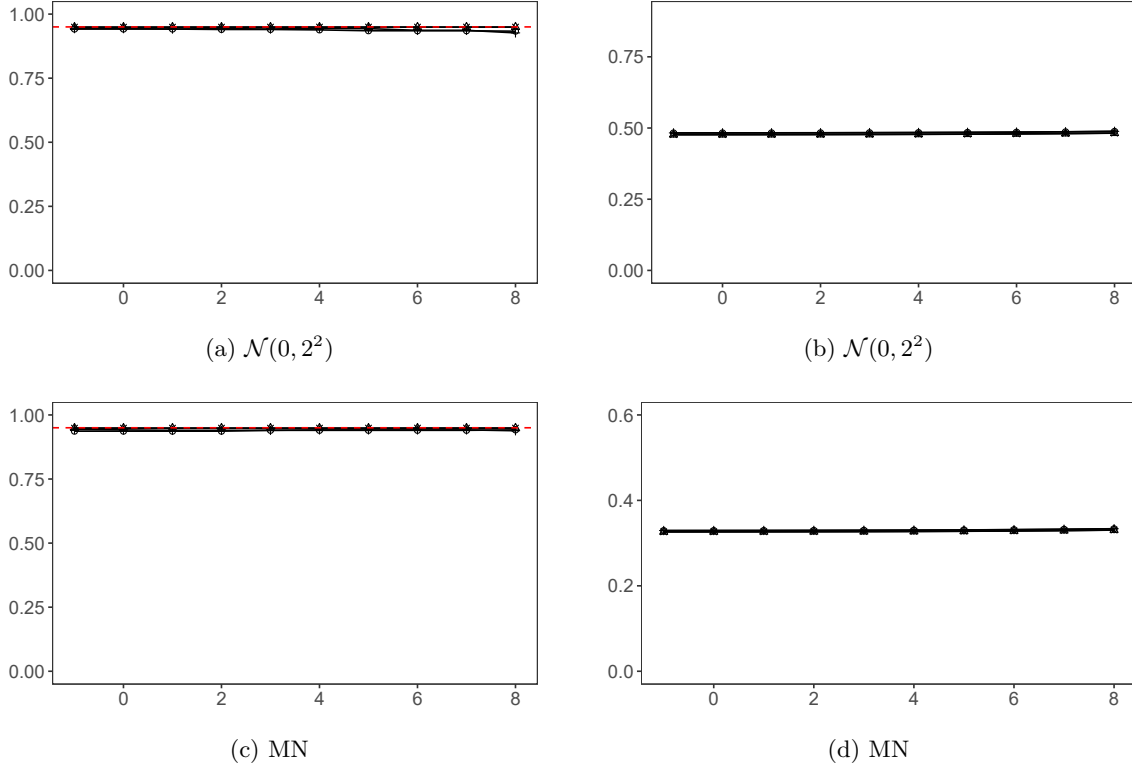


Figure 2: Sensitivity Analysis for  $c$  with  $\alpha_0 = 0.10$  and  $c_0 = 0.01$  Fixed. Empirical coverage probabilities (a, c) and interval lengths (b, d) of confidence intervals are shown on the vertical axis against  $i \in \{-1, 0, \dots, 8\}$  (for  $c = 1 + 10^{i/8-2}$ ) on the horizontal axis, based on 200 simulations. Results are displayed for DSRLasso ( $\circ$  for  $\mathcal{B}$ ,  $\triangle$  for  $\mathcal{B}^c$ ) and DSRLasso<sup>sub</sup> ( $+$  for  $\mathcal{B}$ ,  $\times$  for  $\mathcal{B}^c$ ). A red dashed line indicating the nominal 95% level is overlaid in (a) and (c).

## B.2 Sensitivity Analysis for Constants Used in Tuning Parameter Selection

In this section, we conduct a sensitivity analysis of the constants  $c$ ,  $\alpha_0$  and  $c_0$  used in the selection of the tuning parameter given in (3). We evaluate the empirical coverage probabilities and lengths of confidence intervals based on 200 simulations, as shown in Figures 2–4. Specifically, in Figure 2, we fix  $\alpha_0 = 0.10$  and  $c_0 = 0.01$ , and vary  $c = 1 + 10^{i/8-2}$  for  $i \in \{-1, 0, \dots, 8\}$ . In Figure 3, we fix  $c = 1.01$  and  $c_0 = 0.01$ , and vary  $\alpha_0 \in \{0.025, 0.050, \dots, 0.250\}$ . In Figure 4, we fix  $\alpha_0 = 0.10$  and  $c = 1.01$ , and vary  $c_0 = 20^{i/8-1}/5$  for  $i \in \{-1, 0, \dots, 8\}$ .

Overall, the performance of the proposed methods, in terms of both coverage probability and interval length, is quite robust to the choice of these constants. The results suggest that it is preferable to choose  $\alpha_0$  not too small, and  $c$  and  $c_0$  not too large.

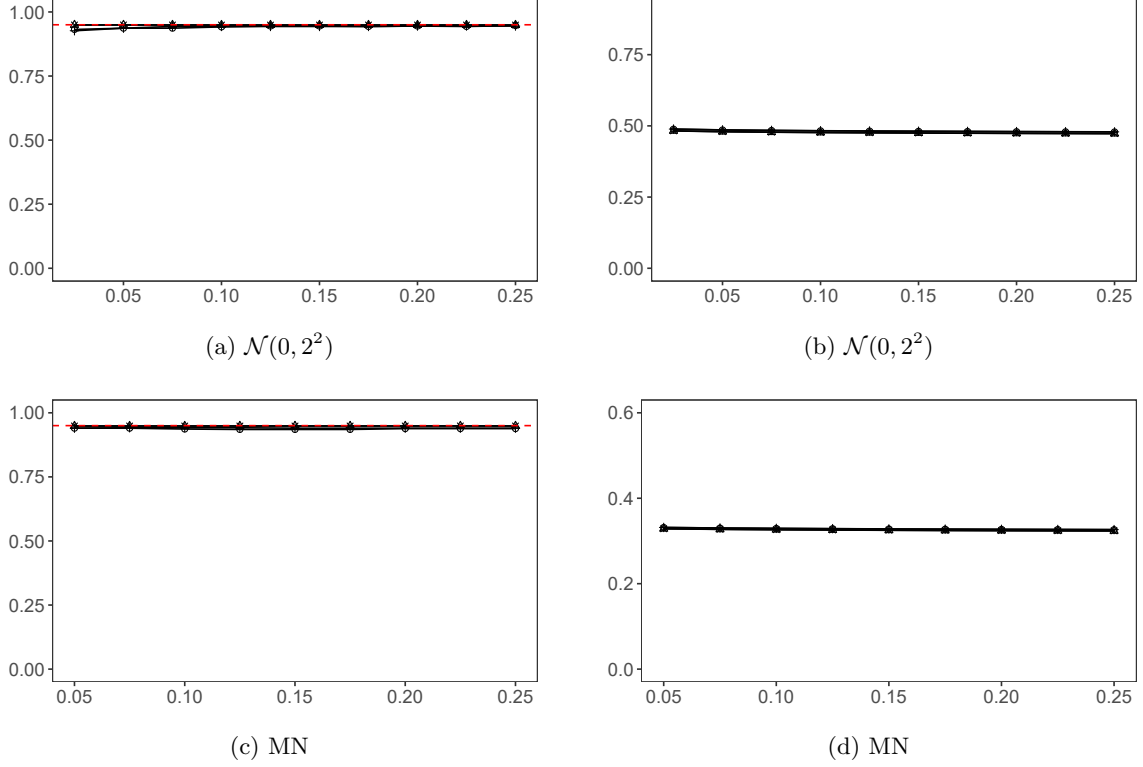


Figure 3: Sensitivity Analysis for  $\alpha_0$  with  $c = 1.01$  and  $c_0 = 0.01$  Fixed. Empirical coverage probabilities (a, c) and interval lengths (b, d) of confidence intervals are shown on the vertical axis against  $\alpha_0 \in \{0.025, 0.050, \dots, 0.250\}$  on the horizontal axis, based on 200 simulations. Results are displayed for DSRLasso ( $\circ$  for  $\mathcal{B}$ ,  $\triangle$  for  $\mathcal{B}^c$ ) and DSRLasso<sup>sub</sup> ( $+$  for  $\mathcal{B}$ ,  $\times$  for  $\mathcal{B}^c$ ). A red dashed line indicating the nominal 95% level is overlaid in (a) and (c).

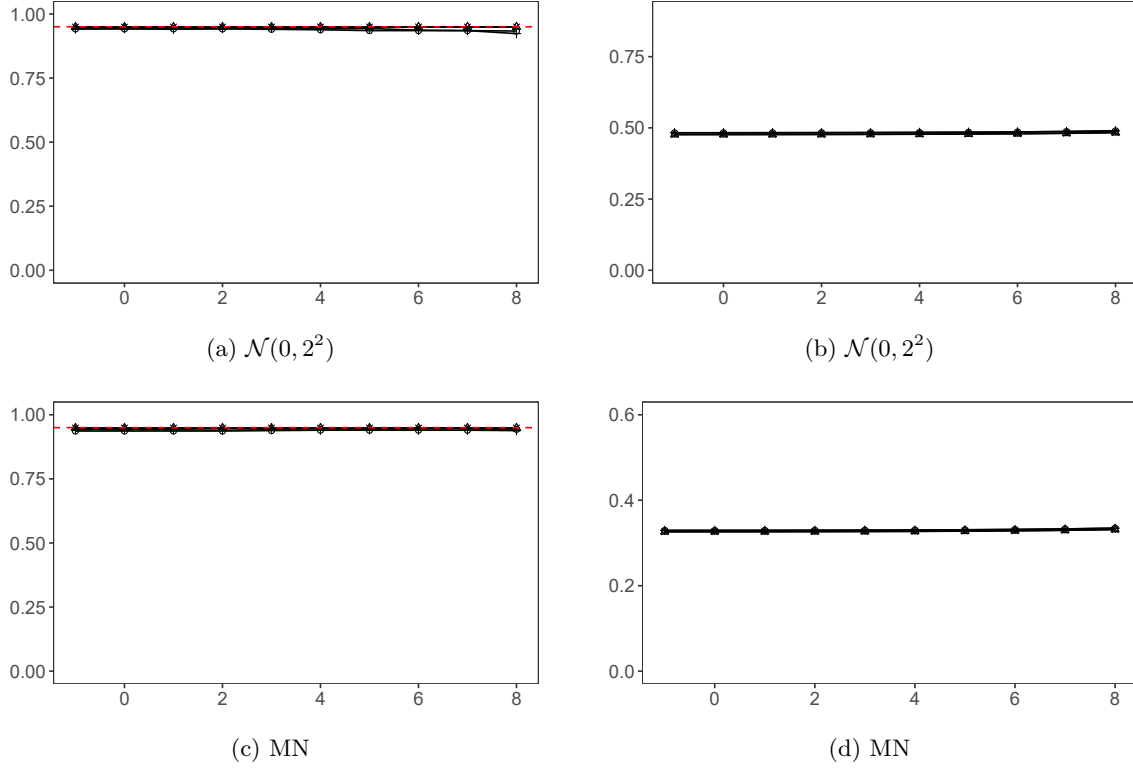


Figure 4: Sensitivity Analysis for  $c_0$  with  $c = 1.01$  and  $\alpha_0 = 0.10$  Fixed. Empirical coverage probabilities (a, c) and interval lengths (b, d) of confidence intervals are shown on the vertical axis against  $i \in \{-1, 0, \dots, 8\}$  (for  $c_0 = 20^{i/8-1}/5$ ) on the horizontal axis, based on 200 simulations. Results are displayed for DSRLasso ( $\circ$  for  $\mathcal{B}$ ,  $\triangle$  for  $\mathcal{B}^c$ ) and DSRLasso<sup>sub</sup> ( $+$  for  $\mathcal{B}$ ,  $\times$  for  $\mathcal{B}^c$ ). A red dashed line indicating the nominal 95% level is overlaid in (a) and (c).

| Error                 | $g_0(z)$ | Method                  | CP <sub>1</sub> | CP <sub>500</sub> | ACP <sub>B</sub> | ACP <sub>B<sup>c</sup></sub> | ARP <sub>B</sub> | RCP <sub>1</sub> | AIL <sub>B</sub> | AIL <sub>B<sup>c</sup></sub> |
|-----------------------|----------|-------------------------|-----------------|-------------------|------------------|------------------------------|------------------|------------------|------------------|------------------------------|
| $\mathcal{N}(0, 2^2)$ | (i)      | TSLasso                 | 0.925           | 0.945             | 0.923            | 0.948                        | 1.000            | 0.925            | 0.386            | 0.390                        |
|                       |          | DSRLasso                | 0.960           | 0.965             | 0.942            | 0.950                        | 1.000            | 0.960            | 0.482            | 0.476                        |
|                       |          | DSRLasso <sup>sub</sup> | 0.955           | 0.965             | 0.938            | 0.950                        | 1.000            | 0.955            | 0.483            | 0.477                        |
|                       | (ii)     | TSLasso                 | 0.920           | 0.935             | 0.925            | 0.949                        | 1.000            | 0.920            | 0.385            | 0.389                        |
|                       |          | DSRLasso                | 0.940           | 0.940             | 0.938            | 0.949                        | 1.000            | 0.940            | 0.483            | 0.477                        |
|                       |          | DSRLasso <sup>sub</sup> | 0.940           | 0.935             | 0.940            | 0.949                        | 1.000            | 0.940            | 0.483            | 0.477                        |
|                       | (iii)    | TSLasso                 | 0.920           | 0.945             | 0.918            | 0.949                        | 1.000            | 0.920            | 0.384            | 0.388                        |
|                       |          | DSRLasso                | 0.960           | 0.950             | 0.940            | 0.949                        | 1.000            | 0.960            | 0.483            | 0.477                        |
|                       |          | DSRLasso <sup>sub</sup> | 0.965           | 0.960             | 0.942            | 0.949                        | 1.000            | 0.965            | 0.483            | 0.477                        |
| MN                    | (i)      | TSLasso                 | 0.910           | 0.950             | 0.872            | 0.949                        | 1.000            | 0.910            | 0.466            | 0.472                        |
|                       |          | DSRLasso                | 0.955           | 0.955             | 0.945            | 0.949                        | 1.000            | 0.955            | 0.327            | 0.324                        |
|                       |          | DSRLasso <sup>sub</sup> | 0.955           | 0.950             | 0.943            | 0.949                        | 1.000            | 0.955            | 0.327            | 0.324                        |
|                       | (ii)     | TSLasso                 | 0.905           | 0.955             | 0.880            | 0.949                        | 1.000            | 0.905            | 0.464            | 0.470                        |
|                       |          | DSRLasso                | 0.965           | 0.950             | 0.948            | 0.949                        | 1.000            | 0.965            | 0.326            | 0.323                        |
|                       |          | DSRLasso <sup>sub</sup> | 0.955           | 0.950             | 0.943            | 0.949                        | 1.000            | 0.955            | 0.326            | 0.323                        |
|                       | (iii)    | TSLasso                 | 0.905           | 0.955             | 0.872            | 0.949                        | 1.000            | 0.905            | 0.465            | 0.471                        |
|                       |          | DSRLasso                | 0.950           | 0.960             | 0.943            | 0.950                        | 1.000            | 0.950            | 0.330            | 0.327                        |
|                       |          | DSRLasso <sup>sub</sup> | 0.950           | 0.945             | 0.942            | 0.950                        | 1.000            | 0.950            | 0.330            | 0.327                        |

Table 6: Simulation Results for Example 2 with Different Forms of Nuisance Functions  $g_0(z)$ . Empirical Coverage Probabilities, Power and Lengths of Confidence Intervals Based on 200 Simulations.

### B.3 Different Forms of Nuisance Functions

In this section, we examine the sensitivity of the proposed inference methods to different forms of the nuisance function  $g_0(z)$ . Specifically, we consider the following additional forms:

(i)  $g_0(z) = \frac{(z/2)^{10}(1-z/2)^4 + (z/2)^4(1-z/2)^{10}}{B(11,5)}$ , where  $B(a, b) = \int_0^1 x^{a-1}(1-x)^{b-1} dx$  is the beta function;

(ii)  $g_0(z) = \exp\{\sin(2\pi z)\} + (z-1)^2$ ; and

(iii)  $g_0(z) = z\{1 + \sin(2\pi z)\}$ .

The results based on 200 simulations are summarized in Table 6. The evaluation metrics remain the same as those used in Example 2. Overall, the performance of the proposed inference methods is robust across the different functional forms of  $g_0(z)$ .

### B.4 Different Nonparametric Estimators

To investigate the sensitivity of the proposed methods to the choice of nonparametric estimators, we first compare the simulation results for Example 2 using kernel-based and B-spline-based estimators for the nonparametric functions  $m_0(\cdot)$  and  $m_{\mathbf{x}}(\cdot)$ , as summarized in Table 7. Here, DSRLasso(Kernel) and DSRLasso<sup>sub</sup>(Kernel) denote the methods using kernel estimation, while DSRLasso(Spline) and DSRLasso<sup>sub</sup>(Spline) denote those using

| Error                 | Method                           | CP <sub>1</sub> | CP <sub>500</sub> | ACP <sub>B</sub> | ACP <sub>B<sup>c</sup></sub> | ARP <sub>B</sub> | RCP <sub>1</sub> | AIL <sub>B</sub> | AIL <sub>B<sup>c</sup></sub> |
|-----------------------|----------------------------------|-----------------|-------------------|------------------|------------------------------|------------------|------------------|------------------|------------------------------|
| $\mathcal{N}(0, 2^2)$ | DSRLasso(Kernel)                 | 0.955           | 0.955             | 0.942            | 0.949                        | 1.000            | 0.955            | 0.483            | 0.476                        |
|                       | DSRLasso <sup>sub</sup> (Kernel) | 0.960           | 0.950             | 0.943            | 0.949                        | 1.000            | 0.960            | 0.483            | 0.477                        |
|                       | DSRLasso(Spline)                 | 0.945           | 0.955             | 0.938            | 0.952                        | 1.000            | 0.945            | 0.531            | 0.533                        |
|                       | DSRLasso <sup>sub</sup> (Spline) | 0.950           | 0.955             | 0.940            | 0.952                        | 1.000            | 0.950            | 0.532            | 0.534                        |
| MN                    | DSRLasso(Kernel)                 | 0.945           | 0.950             | 0.937            | 0.948                        | 1.000            | 0.945            | 0.329            | 0.326                        |
|                       | DSRLasso <sup>sub</sup> (Kernel) | 0.955           | 0.950             | 0.942            | 0.948                        | 1.000            | 0.955            | 0.330            | 0.326                        |
|                       | DSRLasso(Spline)                 | 0.960           | 0.930             | 0.943            | 0.954                        | 1.000            | 0.960            | 0.386            | 0.387                        |
|                       | DSRLasso <sup>sub</sup> (Spline) | 0.960           | 0.935             | 0.943            | 0.954                        | 1.000            | 0.960            | 0.387            | 0.387                        |

Table 7: A Comparison of Kernel-Based and Spline-Based Estimators for Nonparametric Functions. Empirical Coverage Probabilities, Power and Lengths of Confidence Intervals Based on 200 Simulations.

B-splines. We note that, under certain regularity conditions (Newey, 1997), uniform convergence rates have been established for B-splines and, more generally, for series estimators. Hence, Condition (C1) can also be satisfied by such nonparametric estimators under appropriate regularity assumptions. The results indicate that the spline-based methods yield coverage probabilities comparable to those of the kernel-based methods, though with slightly wider interval lengths.

In addition, as discussed in Section 4.1, even when the moment condition for  $\tilde{Y}_i$  is not satisfied, such as when the error follows a  $t(2)$  or Cauchy distribution, the performance of the proposed methods remains reasonable. This indicates that even mild violations of Condition (C1) do not substantially impair the methods, since this condition is only required in the first-step regression and not in the second-step rank regression. To further illustrate this point, we also compare the results from Example 2 using local linear mean regression (Fan and Gijbels, 1996; the method used in the main article) and local linear median regression (Koenker, 2005) for estimating  $m_0(\cdot)$ , as presented in Table 8. Here, DSRLasso(Mean) and DSRLasso<sup>sub</sup>(Mean) refer to the methods using local mean regression, while DSRLasso(Median) and DSRLasso<sup>sub</sup>(Median) denote those using local median regression. It is worth noting that, under certain regularity conditions (Guerre and Sabbah, 2012), uniform convergence rates for local median regression have been established. In our simulation settings, Condition (C1) holds even for  $t(2)$  and Cauchy errors when local median regression is employed. We observe that the interval lengths for DSRLasso(Mean) and DSRLasso<sup>sub</sup>(Mean) are slightly smaller under normal, mixture-normal, and  $t(2)$  distributions, but larger under the Cauchy distribution. Nevertheless, the coverage probabilities across all methods remain roughly comparable.

In summary, although the interval lengths of the confidence intervals may vary depending on the choice of nonparametric estimator, the coverage probabilities of the proposed methods are quite robust to this choice. Moreover, while Condition (C1) is important for theoretical justification, the methods perform well even when it is mildly violated.

| Error                 | Method                           | CP <sub>1</sub> | CP <sub>500</sub> | ACP <sub>B</sub> | ACP <sub>B<sup>c</sup></sub> | ARP <sub>B</sub> | RCP <sub>1</sub> | AIL <sub>B</sub> | AIL <sub>B<sup>c</sup></sub> |
|-----------------------|----------------------------------|-----------------|-------------------|------------------|------------------------------|------------------|------------------|------------------|------------------------------|
| $\mathcal{N}(0, 2^2)$ | DSRLasso(Mean)                   | 0.955           | 0.955             | 0.942            | 0.949                        | 1.000            | 0.955            | 0.483            | 0.476                        |
|                       | DSRLasso <sup>sub</sup> (Mean)   | 0.960           | 0.950             | 0.943            | 0.949                        | 1.000            | 0.960            | 0.483            | 0.477                        |
|                       | DSRLasso(Median)                 | 0.950           | 0.950             | 0.938            | 0.948                        | 1.000            | 0.950            | 0.518            | 0.522                        |
|                       | DSRLasso <sup>sub</sup> (Median) | 0.940           | 0.950             | 0.942            | 0.947                        | 1.000            | 0.940            | 0.518            | 0.522                        |
| MN                    | DSRLasso(Mean)                   | 0.945           | 0.950             | 0.937            | 0.948                        | 1.000            | 0.945            | 0.329            | 0.326                        |
|                       | DSRLasso <sup>sub</sup> (Mean)   | 0.955           | 0.950             | 0.942            | 0.948                        | 1.000            | 0.955            | 0.330            | 0.326                        |
|                       | DSRLasso(Median)                 | 0.950           | 0.955             | 0.942            | 0.951                        | 1.000            | 0.950            | 0.367            | 0.363                        |
|                       | DSRLasso <sup>sub</sup> (Median) | 0.955           | 0.960             | 0.935            | 0.951                        | 1.000            | 0.955            | 0.367            | 0.363                        |
| $t(2)$                | DSRLasso(Mean)                   | 0.950           | 0.960             | 0.940            | 0.949                        | 1.000            | 0.950            | 0.555            | 0.560                        |
|                       | DSRLasso <sup>sub</sup> (Mean)   | 0.965           | 0.960             | 0.943            | 0.949                        | 1.000            | 0.965            | 0.556            | 0.562                        |
|                       | DSRLasso(Median)                 | 0.950           | 0.955             | 0.942            | 0.951                        | 1.000            | 0.950            | 0.561            | 0.566                        |
|                       | DSRLasso <sup>sub</sup> (Median) | 0.945           | 0.960             | 0.935            | 0.951                        | 1.000            | 0.945            | 0.562            | 0.568                        |
| Cauchy                | DSRLasso(Mean)                   | 0.950           | 0.945             | 0.937            | 0.951                        | 0.992            | 0.940            | 1.104            | 1.116                        |
|                       | DSRLasso <sup>sub</sup> (Mean)   | 0.945           | 0.935             | 0.933            | 0.952                        | 0.987            | 0.930            | 1.137            | 1.150                        |
|                       | DSRLasso(Median)                 | 0.965           | 0.940             | 0.948            | 0.950                        | 1.000            | 0.965            | 0.610            | 0.616                        |
|                       | DSRLasso <sup>sub</sup> (Median) | 0.965           | 0.940             | 0.948            | 0.950                        | 1.000            | 0.965            | 0.611            | 0.617                        |

Table 8: A Comparison of Local Mean Regression and Local Median Regression for Non-parametric Functions. Empirical Coverage Probabilities, Power and Lengths of Confidence Intervals Based on 200 Simulations.

### B.5 Comparison with the Median-based Method

For comparison, we constructed a debiased data-splitting (DS) Lasso-penalized median estimator, following ideas similar to our proposed debiased DS rank Lasso estimator, as follows:

Using the partial prediction residuals  $\{\hat{Y}_i^{(1)}, \hat{\mathbf{x}}_i^{(1)}\}_{i=1}^n$  and  $\{\hat{Y}_i^{(2)}, \hat{\mathbf{x}}_i^{(2)}\}_{i=1}^n$  obtained in Section 2.1, we define

$$\begin{aligned}
 (\hat{\beta}_{\text{med}}^{(1)}, \hat{\beta}_{\text{med},0}^{(1)}) &= \arg \min_{\beta \in \mathbb{R}^p, \beta_0 \in \mathbb{R}} \left\{ n^{-1} \sum_{i=1}^n |\hat{Y}_i^{(2)} - \beta^T \hat{\mathbf{x}}_i^{(2)} - \beta_0| + \lambda^{(1)} \|\beta\|_1 \right\}, \\
 (\hat{\beta}_{\text{med}}^{(2)}, \hat{\beta}_{\text{med},0}^{(2)}) &= \arg \min_{\beta \in \mathbb{R}^p, \beta_0 \in \mathbb{R}} \left\{ n^{-1} \sum_{i=1}^n |\hat{Y}_i^{(1)} - \beta^T \hat{\mathbf{x}}_i^{(1)} - \beta_0| + \lambda^{(2)} \|\beta\|_1 \right\},
 \end{aligned}$$

where  $\lambda^{(1)}$  and  $\lambda^{(2)}$  are tuning parameters selected via cross-validation. The debiased estimator is then defined as

$$\hat{\beta}_{\text{med}}^{\text{d},(l)} = \hat{\beta}_{\text{med}}^{(l)} + \hat{\Theta}^T n^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(l)} \text{sign}\{\hat{Y}_i^{(l)} - (\hat{\beta}_{\text{med}}^{(l)})^T \hat{\mathbf{x}}_i^{(l)} - \hat{\beta}_{\text{med},0}^{(l)}\} / (2\hat{\kappa}_{\text{med}}),$$

for each  $l \in \{1, 2\}$ . Here,  $\hat{\Theta}$  is the estimator of  $\Theta$  as in Section 3, and  $\hat{\kappa}_{\text{med}}$  is the estimator of  $\kappa_{\text{med}} = f\{\text{med}(\epsilon_i)\}$ , with  $f(\cdot)$  and  $\text{med}(\epsilon_i)$  denoting the probability density function and the population median of  $\epsilon_i$ , respectively. Specifically, we define  $\hat{\kappa}_{\text{med}}^{(l)} = n^{-1} \sum_{i=1}^n K(\hat{\epsilon}_{\text{med},i}^{(l)}/h)/h$ , where  $K(\cdot)$  is a kernel function,  $h > 0$  is the bandwidth, and

$\hat{\epsilon}_{\text{med},i}^{(l)} = \hat{Y}_i^{(l)} - (\hat{\mathbf{x}}_i^{(l)})^\top \hat{\boldsymbol{\beta}}_{\text{med}}^{(l)} - \hat{\beta}_{\text{med},0}^{(l)}$ . The final estimator for  $\kappa_{\text{med}}$  is  $\hat{\kappa}_{\text{med}} = (\hat{\kappa}_{\text{med}}^{(1)} + \hat{\kappa}_{\text{med}}^{(2)})/2$ . The final debiased DS Lasso-penalized median estimator is

$$\hat{\boldsymbol{\beta}}_{\text{med}}^{\text{d}} = \frac{\hat{\boldsymbol{\beta}}_{\text{med}}^{\text{d},(1)} + \hat{\boldsymbol{\beta}}_{\text{med}}^{\text{d},(2)}}{2}.$$

Following arguments similar to Bradic and Kolar (2017) and the proof of Theorem 5, we can show that, under certain conditions,

$$\sqrt{N}(\hat{\beta}_{\text{med},k}^{\text{d}} - \beta_{0,k})/\hat{\omega}_{\text{med},k} \xrightarrow{d} \mathcal{N}(0,1),$$

for each  $k \in \{1, \dots, p\}$ , where  $\hat{\beta}_{\text{med},k}^{\text{d}}$  is the  $k$ -th component of  $\hat{\boldsymbol{\beta}}_{\text{med}}^{\text{d}}$ ,

$$\hat{\omega}_{\text{med},k}^2 = \hat{\boldsymbol{\theta}}_k^\top \left( \frac{\hat{\boldsymbol{\Sigma}}^{(1)} + \hat{\boldsymbol{\Sigma}}^{(2)}}{2} \right) \hat{\boldsymbol{\theta}}_k / (4\hat{\kappa}_{\text{med}}^2),$$

with  $\hat{\boldsymbol{\Sigma}}^{(l)} = n^{-1} \sum_{i=1}^n (\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)})(\hat{\mathbf{x}}_i^{(l)} - \bar{\mathbf{x}}^{(l)})^\top$  and  $\bar{\mathbf{x}}^{(l)} = n^{-1} \sum_{i=1}^n \hat{\mathbf{x}}_i^{(l)}$ . The  $(1 - \alpha)$ -confidence interval for  $\beta_{0,k}$  is thus

$$\left[ \hat{\beta}_{\text{med},k}^{\text{d}} - \Phi^{-1}(1 - \alpha/2)\hat{\omega}_{\text{med},k}/\sqrt{N}, \quad \hat{\beta}_{\text{med},k}^{\text{d}} + \Phi^{-1}(1 - \alpha/2)\hat{\omega}_{\text{med},k}/\sqrt{N} \right].$$

For  $\omega_{\text{med},k}^2 = \Theta_{k,k}/(4\kappa_{\text{med}}^2)$ , we have  $\hat{\omega}_{\text{med},k}^2/\omega_{\text{med},k}^2$  probability to 1, and hence the asymptotic relative efficiency (ARE) of  $\beta_k^{\text{d}}$  with respect to  $\beta_{\text{med},k}^{\text{d}}$  is  $3\{f^*(0)\}^2/\kappa_{\text{med}}^2$ . When  $\epsilon_i$  follows a normal distribution, the ARE is 1.5, indicating significant efficiency loss when using the debiased DS Lasso-penalized median estimator compared to the debiased DS rank Lasso estimator.

Table 9 compares the performance of the proposed methods with the median-based method described above (denoted DSMLasso). The evaluation metrics are the same as those in Example 2, and the results are based on 200 simulations. Overall, the confidence intervals for the median-based method are significantly wider than those for the proposed rank-based methods when the random errors follow a normal distribution. As the error distributions become more heavy-tailed, the difference diminishes, and the median-based intervals even become slightly shorter than those of the proposed methods under a Cauchy error distribution.

## B.6 An Additional Real Example

We apply the proposed methods to analyze workers' wage data from Berndt (1991), which includes 534 workers' wages, years of experience, education, living region, gender, race, occupation, and marital status. This data set has also been studied in Xie and Huang (2009) and Zhu et al. (2019), both of which concern high-dimensional partially linear models. We consider the following model:

$$Y_i = \sum_{k=1}^{14} X_{i,k} \beta_{0,k} + g_0(Z_i) + \epsilon_i, \quad i = 1, \dots, 534,$$

| Error                 | Method                  | CP <sub>1</sub> | CP <sub>500</sub> | ACP <sub>B</sub> | ACP <sub>B<sup>c</sup></sub> | ARP <sub>B</sub> | RCP <sub>1</sub> | AIL <sub>B</sub> | AIL <sub>B<sup>c</sup></sub> |
|-----------------------|-------------------------|-----------------|-------------------|------------------|------------------------------|------------------|------------------|------------------|------------------------------|
| $\mathcal{N}(0, 2^2)$ | DSRLasso                | 0.955           | 0.955             | 0.942            | 0.949                        | 1.000            | 0.955            | 0.483            | 0.476                        |
|                       | DSRLasso <sup>sub</sup> | 0.960           | 0.950             | 0.943            | 0.949                        | 1.000            | 0.960            | 0.483            | 0.477                        |
|                       | DSMLasso                | 0.930           | 0.955             | 0.945            | 0.951                        | 1.000            | 0.930            | 0.596            | 0.588                        |
| MN                    | DSRLasso                | 0.945           | 0.950             | 0.937            | 0.948                        | 1.000            | 0.945            | 0.329            | 0.326                        |
|                       | DSRLasso <sup>sub</sup> | 0.955           | 0.950             | 0.942            | 0.948                        | 1.000            | 0.955            | 0.330            | 0.326                        |
|                       | DSMLasso                | 0.960           | 0.960             | 0.943            | 0.953                        | 1.000            | 0.960            | 0.369            | 0.366                        |
| $t(2)$                | DSRLasso                | 0.950           | 0.960             | 0.940            | 0.949                        | 1.000            | 0.950            | 0.555            | 0.560                        |
|                       | DSRLasso <sup>sub</sup> | 0.965           | 0.960             | 0.943            | 0.949                        | 1.000            | 0.965            | 0.556            | 0.562                        |
|                       | DSMLasso                | 0.940           | 0.950             | 0.947            | 0.952                        | 1.000            | 0.940            | 0.584            | 0.577                        |
| Cauchy                | DSRLasso                | 0.950           | 0.945             | 0.937            | 0.951                        | 0.992            | 0.940            | 1.104            | 1.116                        |
|                       | DSRLasso <sup>sub</sup> | 0.945           | 0.935             | 0.933            | 0.952                        | 0.987            | 0.930            | 1.137            | 1.150                        |
|                       | DSMLasso                | 0.950           | 0.965             | 0.955            | 0.962                        | 0.990            | 0.945            | 1.067            | 1.079                        |

Table 9: A Comparison of the Proposed Methods and the Median-based Method. Empirical Coverage Probabilities, Power and Lengths of Confidence Intervals Based on 200 Simulations.

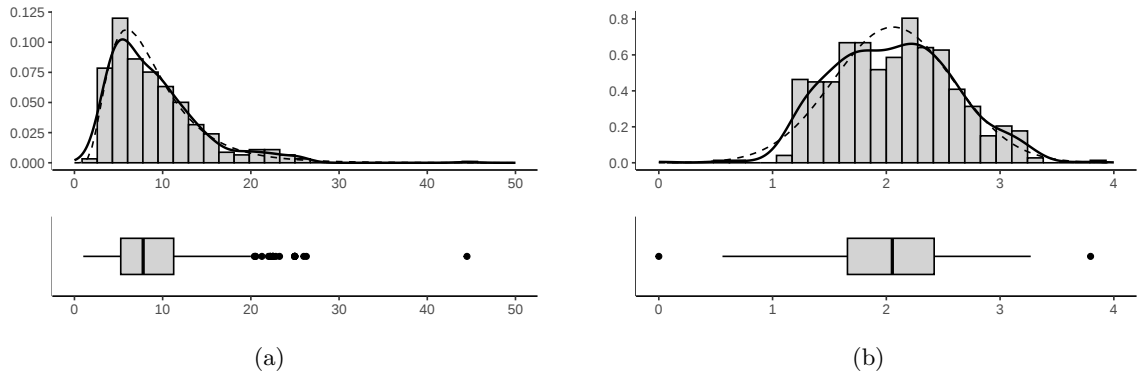


Figure 5: (a) Histogram of Wages ( $Y_i$ ) for 534 Workers, Overlaid with a Boxplot, a Kernel Density Curve (Solid Line), and a Log-normal Density Curve (Dashed Line); (b) Histogram of Natural Logarithms of Wages ( $\log(Y_i)$ ), Overlaid with a Boxplot, a Kernel Density Curve (Solid Line), and a Normal Density Curve (Dashed Line).

| Description |                                 | TSLasso                      |                  | DSRLasso                     |                  |
|-------------|---------------------------------|------------------------------|------------------|------------------------------|------------------|
|             |                                 | $\hat{\beta}_k^d(\text{SE})$ | $p\text{-value}$ | $\hat{\beta}_k^d(\text{SE})$ | $p\text{-value}$ |
| $X_1$       | Years of education              | 0.620(0.100)                 | < 0.001          | 0.486(0.088)                 | < 0.001          |
| $X_2$       | 1 = southern region, 0 = other  | -0.505(0.403)                | 0.209            | -0.882(0.353)                | 0.013            |
| $X_3$       | 1 = female, 0 = other           | -2.002(0.411)                | < 0.001          | -1.833(0.367)                | < 0.001          |
| $X_4$       | 1 = union member, 0 = nonmember | 1.582(0.479)                 | 0.001            | 1.767(0.413)                 | < 0.001          |
| $X_5$       | 1 = nonwhite, 0 = other         | -0.838(0.554)                | 0.130            | -0.547(0.480)                | 0.255            |
| $X_6$       | 1 = Hispanic, 0 = other         | -0.545(0.838)                | 0.515            | -1.089(0.741)                | 0.142            |
| $X_7$       | 1 = management, 0 = other       | 3.401(0.767)                 | < 0.001          | 2.710(0.660)                 | < 0.001          |
| $X_8$       | 1 = sales, 0 = other            | -0.588(0.793)                | 0.458            | -0.195(0.681)                | 0.775            |
| $X_9$       | 1 = clerical, 0 = other         | 0.084(0.630)                 | 0.894            | 0.620(0.542)                 | 0.253            |
| $X_{10}$    | 1 = service, 0 = other          | -0.522(0.643)                | 0.417            | -0.265(0.561)                | 0.637            |
| $X_{11}$    | 1 = professional, 0 = other     | 2.096(0.689)                 | 0.002            | 1.919(0.599)                 | 0.001            |
| $X_{12}$    | 1 = manufacturing, 0 = other    | 1.171(0.537)                 | 0.029            | 1.398(0.465)                 | 0.003            |
| $X_{13}$    | 1 = construction, 0 = other     | 0.691(0.930)                 | 0.458            | 0.922(0.797)                 | 0.247            |
| $X_{14}$    | 1 = married, 0 = other          | 0.026(0.410)                 | 0.949            | 0.428(0.357)                 | 0.230            |

Table 10: Estimates of  $\beta_{0,k}$  ( $k = 1, \dots, 14$ ) for the Real Data Example. The debiased estimators ( $\hat{\beta}_k^d$ ) are presented alongside their standard errors (in parentheses) and corresponding  $p$ -values.

where  $Y_i$  denotes the wage of the  $i$ -th worker,  $Z_i$  represents the years of experience, and  $\epsilon_i$  ( $i = 1, \dots, 534$ ) are independent and identically distributed random errors. The design matrix  $\mathbf{X} = (X_{i,k})_{534 \times 14}$  contains covariates other than the years of experience, with  $X_{i,k}$  being the  $k$ -th covariate for the  $i$ -th worker.

The histogram of wages for 534 workers and their corresponding natural logarithm transformations are presented in Figures 5 (a) and (b), respectively. Boxplots and kernel density curves are included in both panels. Additionally, a log-normal density curve is overlaid in Figure 5 (a), while a normal density curve is overlaid in Figure 5 (b). Notably, the wage distribution shows a close resemblance to a log-normal distribution, which motivates the application of the proposed debiased DS rank Lasso (DSRLasso) to this data set.

Table 10 presents the debiased DSRLasso estimators of  $\beta_{0,k}$  ( $k = 1, \dots, 14$ ), along with their standard errors and  $p$ -values. For comparison, we also include results from the debiased two-step projection based Lasso (TSLasso; Zhu, 2017; Zhu et al., 2019). The signs of the debiased DSRLasso and TSLasso estimators align exactly. However, debiased DSRLasso yields consistently smaller standard errors. Furthermore, debiased DSRLasso identifies one additional significant covariate ( $X_2$ , whether living in the south or not) at the 5% significance level, compared to debiased TSLasso.

To evaluate the performance of debiased DSRLasso in high-dimensional settings, we expand the original design matrix  $\mathbf{X}$  into an augmented matrix  $\mathbf{X}_{\text{aug}} = (\mathbf{X}, \mathbf{X}^{(1)}, \dots, \mathbf{X}^{(40)}) \in \mathbb{R}^{534 \times 574}$ . Each submatrix  $\mathbf{X}^{(b)}$  ( $b = 1, \dots, 40$ ) is constructed by randomly sampling rows from  $\mathbf{X}$ , ensuring each row of  $\mathbf{X}^{(b)}$  has the same distribution as the original covariates

|          | Description                     | TSLasso | DSRLasso |
|----------|---------------------------------|---------|----------|
| $X_1$    | Years of education              | 1.000   | 1.000    |
| $X_2$    | 1 = southern region, 0 = other  | 0.110   | 0.985    |
| $X_3$    | 1 = female, 0 = other           | 1.000   | 1.000    |
| $X_4$    | 1 = union member, 0 = nonmember | 1.000   | 1.000    |
| $X_5$    | 1 = nonwhite, 0 = other         | 0.015   | 0.005    |
| $X_6$    | 1 = Hispanic, 0 = other         | 0.000   | 0.160    |
| $X_7$    | 1 = management, 0 = other       | 1.000   | 1.000    |
| $X_8$    | 1 = sales, 0 = other            | 0.050   | 0.195    |
| $X_9$    | 1 = clerical, 0 = other         | 0.000   | 0.000    |
| $X_{10}$ | 1 = service, 0 = other          | 0.155   | 0.580    |
| $X_{11}$ | 1 = professional, 0 = other     | 1.000   | 0.955    |
| $X_{12}$ | 1 = manufacturing, 0 = other    | 0.830   | 0.880    |
| $X_{13}$ | 1 = construction, 0 = other     | 0.000   | 0.040    |
| $X_{14}$ | 1 = married, 0 = other          | 0.000   | 0.010    |

Table 11: Rejection probabilities of the first 14 variables at the 5% significance level in 200 replications.

| Method | TSLasso |        | DSRLasso |        |
|--------|---------|--------|----------|--------|
|        | Power   | Length | Power    | Length |
| NST    | 1.000   | 3.812  | 1.000    | 3.558  |
| ST     | 1.000   | 2.913  | 1.000    | 2.713  |

Table 12: Empirical Powers and Average Interval Lengths of NST-TSLasso, ST-TSLasso, NST-DSRLasso and ST-DSRLasso for Testing the Simultaneous Hypothesis  $H_{0,\mathcal{G}_0} : \beta_{0,1} = \dots = \beta_{0,14} = 0$ . “Power” and “Length” denote the empirical powers and average interval lengths, respectively. The results are based on 200 replications.

$\{X_{i,k}\}_{k=1}^{14}$ . The augmented model is:

$$Y_i = \sum_{k=1}^{574} X_{i,k} \beta_{0,k} + g_0(Z_i) + \epsilon_i, \quad i = 1, \dots, 534,$$

where  $\beta_{0,15} = \dots = \beta_{0,574} = 0$  by construction.

Let  $p_k$  denote the  $p$ -value for testing the significance of the  $k$ -th covariate. The average type I error rate, defined as  $\sum_{k=15}^{574} \text{pr}(p_k \leq 0.05)/560$ , is estimated based on 200 replications. The resulting values are 0.051 for debiased DSRLasso and 0.052 for debiased TSLasso, demonstrating that both methods effectively control type I errors at the 5% significance level.

To evaluate the relative importance of the first 14 variables, we compare their rejection probabilities at the 5% significance level across 200 replications. As shown in Table 11, covariates with smaller  $p$ -values in Table 10 generally demonstrate higher relative importance. Notably, the living region covariate ( $X_2$ ) is identified as highly important by debiased DSR-

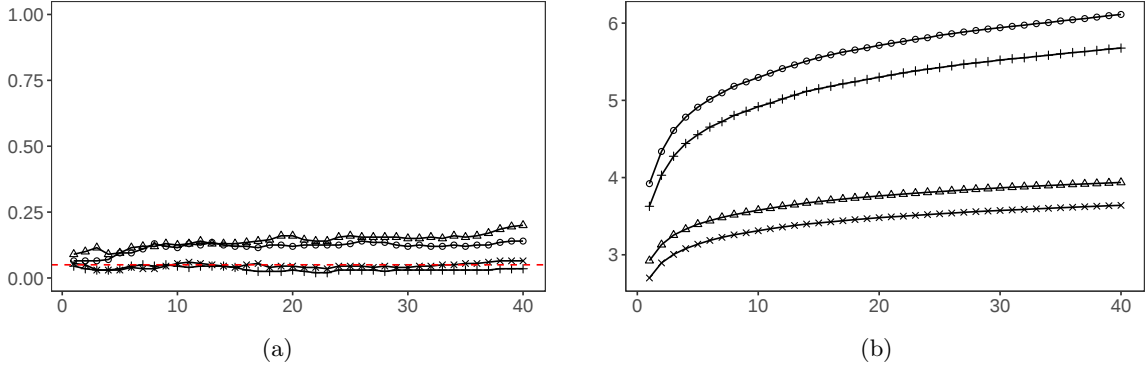


Figure 6: Empirical Type I Error Rates (a) and Average Interval Lengths (b) for Testing the Simultaneous Hypothesis  $H_{0,\mathcal{G}_B} : \beta_{0,15} = \dots = \beta_{0,14 \times (B+1)} = 0$  Across  $B \in \{1, \dots, 40\}$ . The results compare four methods: NST-TSLasso ( $\circ$ ), ST-TSLasso ( $\triangle$ ), NST-DSRLasso ( $+$ ), and ST-DSRLasso ( $\times$ ). A red dashed line indicating the significance level  $\alpha = 0.05$  is overlaid in (a). All results are based on 200 replications.

Lasso but not by debiased TSLasso, which echoes our prior results in the low-dimensional setting.

Finally, we apply the proposed debiased DSRLasso-based simultaneous testing procedure to test  $H_{0,\mathcal{G}} : \beta_{0,k} = 0$  for all  $k \in \mathcal{G}$ . Here,  $\mathcal{G}$  is either  $\mathcal{G}_0 = \{1, \dots, 14\}$  or  $\mathcal{G}_B = \{15, \dots, 14 \times (B+1)\}$  for  $B \in \{1, \dots, 40\}$ . In the case of  $\mathcal{G}_0$ , the null hypothesis holds, while for  $\mathcal{G}_B$  the alternative hypothesis holds. We use the debiased TSLasso-based simultaneous testing procedure from (Zhu et al., 2019) as a benchmark for comparison. Both non-studentized (NST) and studentized (ST) test statistics are considered, leading to four methods: NST-TSLasso, ST-TSLasso, NST-DSRLasso, and ST-DSRLasso. The empirical type I error rates for testing  $H_{0,\mathcal{G}_0}$ , along with the average lengths of the corresponding simultaneous confidence intervals, are presented in Table 12. These results are based on 200 replications. On the other hand, empirical powers for testing  $H_{0,\mathcal{G}_0}$ , along with average interval lengths, are shown in Figure 6. The results indicate that NST-type methods generally achieve better control of type I errors for testing  $H_{0,\mathcal{G}_B}$  compared to ST-type methods, although with wider interval lengths. Moreover, all four methods achieve an empirical power of one when testing  $H_{0,\mathcal{G}_0}$ , and our proposed simultaneous testing procedure provides better type I error control and narrower interval lengths than the debiased TSLasso-based procedure.

## References

- Miguel A. Arcones and Evarist Gine. Limit theorems for  $U$ -processes. *The Annals of Probability*, 21(3):1494–1542, 1993.
- Alexandre Belloni, Victor Chernozhukov, and Christian Hansen. Inference on treatment effects after selection among high-dimensional controls. *The Review of Economic Studies*, 81(2):608–650, 2014.

- Ernst R. Berndt. *The Practice of Econometrics: Classic and Contemporary*. Addison-Wesley Publishing Company, 1991.
- Jelena Bradic and Mladen Kolar. Uniform inference for high-dimensional quantile regression: linear functionals and regression rank scores. *arXiv preprint arXiv:1702.06209*, 2017.
- Leheng Cai, Xu Guo, Heng Lian, and Liping Zhu. Statistical inference for high-dimensional convoluted rank regression. *Journal of the American Statistical Association*, 0(0):1–25, 2025.
- T. Tony Cai and Zijian Guo. Confidence intervals for high-dimensional linear regression: Minimax rates and adaptivity. *The Annals of Statistics*, 45(2):615–646, 2017.
- Jinyuan Chang, Song Xi Chen, Cheng Yong Tang, and Tong Tong Wu. High-dimensional empirical likelihood inference. *Biometrika*, 108(1):127–147, 2021.
- Patrick Chène. Inhibiting the p53–mdm2 interaction: an important target for cancer therapy. *Nature Reviews Cancer*, 3(2):102–109, 2003.
- Victor Chernozhukov, Denis Chetverikov, and Kengo Kato. Gaussian approximations and multiplier bootstrap for maxima of sums of high-dimensional random vectors. *The Annals of Statistics*, 41(6):2786–2819, 2013.
- Victor Chernozhukov, Denis Chetverikov, and Kengo Kato. Comparison and anti-concentration bounds for maxima of Gaussian random vectors. *Probability Theory and Related Fields*, 162(1):47–70, 2015.
- Ruben Dezeure, Peter Bühlmann, and Cun-Hui Zhang. High-dimensional simultaneous inference with the bootstrap. *Test*, 26(4):685–719, 2017.
- Robert F. Engle, C. W. J. Granger, John Rice, and Andrew Weiss. Semiparametric estimates of the relation between weather and electricity sales. *Journal of the American Statistical Association*, 81(394):310–320, 1986.
- Jianqing Fan and Irene Gijbels. *Local Polynomial Modelling and Its Applications*. Chapman & Hall, 1996.
- Jianqing Fan, Runze Li, Cun-Hui Zhang, and Hui Zou. *Statistical Foundations of Data Science*. CRC Press, 2020a.
- Jianqing Fan, Cong Ma, and Kaizheng Wang. Comment on “a tuning-free robust and efficient approach to high-dimensional regression”. *Journal of the American Statistical Association*, 115(532):1720–1725, 2020b.
- Long Feng, Changliang Zou, Zhaojun Wang, and Bin Chen. Rank-based score tests for high-dimensional regression coefficients. *Electronic Journal of Statistics*, 7:2131–2149, 2013.

- Hui Gong, Yixuan Li, Wen Yang, Zishan Xie, Jiahao Hu, Peihang Li, Rou Xu, Yifan Li, Tianyu Tao, Riqing Li, Shuguang Liu, Yefeng Zhu, Libing Song, and Lishan Fang. Prodh2-mediated metabolism in the bone microenvironment promotes breast cancer metastasis. *Cancer Research*, 85(21):4198–4211, 2025.
- Karolina Gronkowska and Agnieszka Robaszkiewicz. Genetic dysregulation of ep300 in cancers in light of cancer epigenome control–targeting of p300-proficient and -deficient cancers. *Molecular Therapy: Oncology*, 32(4):200871, 2024.
- Emmanuel Guerre and Camille Sabbah. Uniform bias study and bahadur representation for local polynomial estimators of the conditional quantile function. *Econometric Theory*, 28(1):87–129, 2012.
- Thomas P. Hettmansperger and Joseph W. McKean. *Robust Nonparametric Statistical Methods*. Chapman & Hall/CRC Monographs on Statistics & Applied Probability. CRC Press, 2010.
- Jean Honorio and Tommi Jaakkola. Tight bounds for the expected risk of linear classifiers and PAC-Bayes finite-sample guarantees. In *Proceedings of the Seventeenth International Conference on Artificial Intelligence and Statistics*, volume 33 of *Proceedings of Machine Learning Research*, pages 384–392. PMLR, 2014.
- Shangshang Hu, Muzi Ding, Jinwei Lou, Jian Qin, Yuhan Chen, Zixuan Liu, Yue Li, Junjie Nie, Mu Xu, Huiling Sun, Xinliang Gu, Tao Xu, Shukui Wang, and Yuqin Pan. Col10a1+ fibroblasts promote colorectal cancer metastasis and m2 macrophage polarization with pan-cancer relevance. *Journal of Experimental & Clinical Cancer Research*, 44(1):243, 2025.
- Adel Javanmard and Andrea Montanari. Confidence intervals and hypothesis testing for high-dimensional regression. *Journal of Machine Learning Research*, 15(1):2869–2909, 2014.
- Adel Javanmard and Andrea Montanari. Debiasing the lasso: Optimal sample size for Gaussian designs. *The Annals of Statistics*, 46(6A):2593 – 2622, 2018.
- Roger Koenker. *Quantile Regression*. Econometric Society Monographs. Cambridge University Press, 2005.
- Chrysoula Komini, Irini Theohari, Andromachi Lambrianidou, Lydia Nakopoulou, and Theoni Trangas. Papola contributes to cyclin d1 mrna alternative polyadenylation and promotes breast cancer cell proliferation. *Journal of Cell Science*, 134(7):jcs252304, 2021.
- M. Ledoux and M. Talagrand. *Probability in Banach Spaces: Isoperimetry and Processes*. A Series of Modern Surveys in Mathematics Series. Springer, 1991.
- Qi Li and Jeffrey Scott Racine. *Nonparametric Econometrics: Theory and Practice*. Princeton University Press, 2007.
- Elias Masry. Multivariate local polynomial regression for time series: Uniform strong consistency and rates. *Journal of Time Series Analysis*, 17(6):571–599, 1996.

- Pascal Massart. About the constants in Talagrand’s concentration inequalities for empirical processes. *The Annals of Probability*, 28(2):863–884, 2000.
- Léa Montégut, Carlos López-Otín, and Guido Kroemer. Aging and cancer. *Molecular Cancer*, 23(1):106, May 2024.
- Pavel Moudry, Katarina Chroma, Sladana Bursac, Sinisa Volarevic, and Jiri Bartek. Rna-interference screen for p53 regulators unveils a role of wdr75 in ribosome biogenesis. *Cell Death & Differentiation*, 29(3):687–696, 2022.
- Whitney K. Newey. Convergence rates and asymptotic normality for series estimators. *Journal of Econometrics*, 79(1):147–168, 1997.
- Matey Neykov, Yang Ning, Jun S Liu, and Han Liu. A unified theory of confidence regions and testing for high-dimensional estimating equations. *Statistical Science*, 33(3):427–443, 2018.
- Yang Ning and Han Liu. A general theory of hypothesis tests and confidence regions for sparse high dimensional models. *The Annals of Statistics*, 45(1):158–195, 2017.
- Andrew G. Renehan, Margaret Tyson, Matthias Egger, Richard F. Heller, and Marcel Zwahlen. Body-mass index and incidence of cancer: a systematic review and meta-analysis of prospective observational studies. *The Lancet*, 371(9612):569–578, 2008.
- Sidney Resnick. *A Probability Path*. Birkhäuser Boston, 2019.
- Charles J. Stone. Optimal global rates of convergence for nonparametric regression. *The Annals of Statistics*, 10(4):1040–1053, 1982.
- Sara van de Geer, Peter Bühlmann, Ya’acov Ritov, and Ruben Dezeure. On asymptotically optimal confidence regions and tests for high-dimensional models. *The Annals of Statistics*, 42(3):1166–1202, 2014.
- Roman Vershynin. *High-Dimensional Probability: An Introduction with Applications in Data Science*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press, 2018.
- Martin J. Wainwright. *High-Dimensional Statistics: A Non-Asymptotic Viewpoint*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press, 2019.
- Han Wang, Ziling Zhou, Hanyi Zhong, Shoutang Wang, Kunwei Shen, Renhong Huang, Ruowang, and Zheng Wang. A pyrimidine metabolism-related gene signature for prognosis prediction and immune microenvironment description of breast cancer. *Journal of Translational Medicine*, 23(1):698, 2025.
- Jun Wang, Xuehong Zhang, Andrew H. Beck, Laura C. Collins, Wendy Y. Chen, Rulla M. Tamimi, Aditi Hazra, Myles Brown, Bernard Rosner, and Susan E. Hankinson. Alcohol consumption and risk of breast cancer by tumor receptor expression. *Hormones and Cancer*, 6(5):237–246, 2015.

- Lan Wang, Bo Peng, Jelena Bradic, Runze Li, and Yunan Wu. A tuning-free robust and efficient approach to high-dimensional regression. *Journal of the American Statistical Association*, 115(532):1700–1714, 2020.
- Huiliang Xie and Jian Huang. SCAD-penalized regression in high-dimensional partially linear models. *The Annals of Statistics*, 37(2):673–696, 2009.
- Limeng Yang, Vanessa A. N. Kraft, Susanne Pfeiffer, Juliane Merl-Pham, Xuanwen Bao, Yu An, Stefanie M. Hauck, and Joel A. Schick. Nonsense-mediated decay factor smg7 sensitizes cells to tnfa-induced apoptosis via cyld tumor suppressor and the noncoding oncogene pvt1. *Molecular Oncology*, 14(10):2420–2435, 2020.
- Jung Eun Yoo, Kyungdo Han, Dong Wook Shin, Dahye Kim, Bong-seong Kim, Sohyun Chun, Keun Hye Jeon, Wonyoung Jung, Jinsung Park, Jin Ho Park, Kui Son Choi, and Joo Sung Kim. Association between changes in alcohol consumption and cancer risk. *JAMA Network Open*, 5(8):e2228544–e2228544, 2022.
- Nikolaos Zacharakis, Harshini Chinnasamy, Mary Black, Hui Xu, Yong-Chen Lu, Zhili Zheng, Anna Pasetto, Michelle Langan, Thomas Shelton, Todd Prickett, Jared Gartner, Li Jia, Katarzyna Trebska-McGowan, Robert P. Somerville, Paul F. Robbins, Steven A. Rosenberg, Stephanie L. Goff, and Steven A. Feldman. Immune recognition of somatic mutations leading to complete durable regression in metastatic breast cancer. *Nature Medicine*, 24(6):724–730, 2018.
- Cun-Hui Zhang and Stephanie S. Zhang. Confidence intervals for low dimensional parameters in high dimensional linear models. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 76(1):217–242, 2014.
- Huijun Zhang, An Wang, Yulong Tan, Shaohua Wang, Qinyun Ma, Xiaofeng Chen, and Zelai He. Ncbp1 promotes the development of lung adenocarcinoma through up-regulation of cul4b. *Journal of Cellular and Molecular Medicine*, 23(10):6965–6977, 2019.
- Xianyang Zhang and Guang Cheng. Simultaneous inference for high-dimensional linear models. *Journal of the American Statistical Association*, 112(518):757–768, 2017.
- Huanyu Zhao, Ruoyu Dang, Yipan Zhu, Baijian Qu, Yasra Sayyed, Ying Wen, Xicheng Liu, Jianping Lin, and Luyuan Li. Hub genes associated with immune cell infiltration in breast cancer, identified through bioinformatic analyses of multiple datasets. *Cancer Biology & Medicine*, 19(9):1352–1374, 2022. doi: 10.20892/j.issn.2095-3941.2021.0586.
- Yue Zhou, Tongjia Zhang, Shujie Wang, Ruixiang Yang, Zitao Jiao, Kejia Lu, Hui Li, Wei Jiang, and Xiaowei Zhang. Targeting of hbp1/timp3 axis as a novel strategy against breast cancer. *Pharmacological Research*, 194:106846, 2023.
- Yinchu Zhu and Jelena Bradic. Linear hypothesis testing in dense high-dimensional linear models. *Journal of the American Statistical Association*, 113(524):1583–1600, 2018a.
- Yinchu Zhu and Jelena Bradic. Significance testing in non-sparse high-dimensional linear models. *Electronic Journal of Statistics*, 12(2):3312–3364, 2018b.

- Ying Zhu. Nonasymptotic analysis of semiparametric regression models with high-dimensional parametric coefficients. *The Annals of Statistics*, 45(5):2274–2298, 2017.
- Ying Zhu, Zhuqing Yu, and Guang Cheng. High dimensional inference in partially linear models. In *Proceedings of the Twenty-Second International Conference on Artificial Intelligence and Statistics*, volume 89 of *Proceedings of Machine Learning Research*, pages 2760–2769. PMLR, 2019.