



A Tuning-free Robust and Efficient Approach to High-dimensional Regression

Lan Wang , Bo Peng , Jelena Bradic , Runze Li & Yunan Wu

To cite this article: Lan Wang , Bo Peng , Jelena Bradic , Runze Li & Yunan Wu (2020) A Tuning-free Robust and Efficient Approach to High-dimensional Regression, Journal of the American Statistical Association, 115:532, 1700-1714, DOI: [10.1080/01621459.2020.1840989](https://doi.org/10.1080/01621459.2020.1840989)

To link to this article: <https://doi.org/10.1080/01621459.2020.1840989>



View supplementary material [↗](#)



Published online: 18 Dec 2020.



Submit your article to this journal [↗](#)



Article views: 319



View related articles [↗](#)



View Crossmark data [↗](#)



Citing articles: 1 View citing articles [↗](#)



A Tuning-Free Robust and Efficient Approach to High-Dimensional Regression

Lan Wang^a, Bo Peng^b, Jelena Bradic^c, Runze Li^d, and Yunan Wu^e

^aDepartment of Management Science, University of Miami, Coral Gables, FL; ^bAdobe Systems, Inc., San Jose, CA; ^cDepartment of Mathematics, Halicioglu Data Science Institute, University of California at San Diego, La Jolla, CA; ^dDepartment of Statistics, Pennsylvania State University, University Park, PA; ^eSchool of Statistics, University of Minnesota, Minneapolis, MN

ABSTRACT

We introduce a novel approach for high-dimensional regression with theoretical guarantees. The new procedure overcomes the challenge of tuning parameter selection of Lasso and possesses several appealing properties. It uses an easily simulated tuning parameter that automatically adapts to both the unknown random error distribution and the correlation structure of the design matrix. It is robust with substantial efficiency gain for heavy-tailed random errors while maintaining high efficiency for normal random errors. Comparing with other alternative robust regression procedures, it also enjoys the property of being equivariant when the response variable undergoes a scale transformation. Computationally, it can be efficiently solved via linear programming. Theoretically, under weak conditions on the random error distribution, we establish a finite-sample error bound with a near-oracle rate for the new estimator with the simulated tuning parameter. Our results make useful contributions to mending the gap between the practice and theory of Lasso and its variants. We also prove that further improvement in efficiency can be achieved by a second-stage enhancement with some light tuning. Our simulation results demonstrate that the proposed methods often outperform cross-validated Lasso in various settings. Supplementary materials for this article are available online.

ARTICLE HISTORY

Received August 2018
Accepted October 2020

KEYWORDS

Efficiency; Heavy-tailed error;
High dimension; Linear
regression; Robustness;
Tuning parameter

1. Introduction

Consider a linear regression model

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta}_0 + \boldsymbol{\epsilon}, \quad (1)$$

where $\mathbf{Y} = (Y_1, \dots, Y_n)^T$ is an $n \times 1$ vector of responses, $\mathbf{X}$ is an $n \times p$ centered matrix of covariates, $\boldsymbol{\beta}_0 = (\beta_{01}, \dots, \beta_{0p})^T$ is a $p \times 1$ vector of unknown parameters, and $\boldsymbol{\epsilon} = (\epsilon_1, \dots, \epsilon_n)^T$ is an $n \times 1$ vector of independent and identically distributed random errors. We are interested in estimating $\boldsymbol{\beta}_0$ in the ultrahigh-dimensional setting, where the number of covariates (features) p can grow exponentially fast with the sample size n .

In the last decade, substantial progress has been achieved in high-dimensional regression analysis. In particular, L_1 regularized least squares regression techniques as represented by Lasso (Tibshirani 1996; Chen, Donoho, and Saunders 2001), Dantzig selector (Candes and Tao 2007), and their variants such as SCAD (Fan and Li 2001), MCP (C. H. Zhang 2010), Capped L_1 (T. Zhang 2010), among others, have become popular tools. The literature in this area is vast. We refer to Bühlmann and van de Geer (2011) and the reviews of Fan and Lv (2010) and Zhang and Zhang (2012) for a fuller list of references. Despite the significant advances in algorithm and theory development, at least two challenges remain.

The first challenge is to determine the right amount of regularization in a computationally efficient way with proper theoretical justification. Practical performance of regularized high-dimensional regression depends crucially on the choice

of tuning parameter λ , which prescribes the level of penalty or shrinkage. For Lasso, it is well understood that the optimal choice of λ depends on both the random error distribution and the design (see, e.g., Meinshausen and Bühlmann 2006; Zhao and Yu 2006; Bunea, Tsybakov, and Wegkamp 2007; Van de Geer 2008; Zhang and Huang 2008; Bickel, Ritov, and Tsybakov 2009; Wainwright 2009). However, theory is often derived while fixing λ at an ideal value $\tau\sigma\sqrt{\log p/n}$, where σ is the standard deviation of the random error distribution and τ is some positive constant. Estimation of σ in high dimensions is itself a very difficult problem (Fan, Guo, and Hao 2012; Sun and Zhang 2012; Dicker 2014; Yu and Bien 2019). To circumvent this difficulty, practitioners often employ cross-validation to select λ . Several recent work shed light on the properties of cross-validated Lasso (e.g., Homrighausen and McDonald 2013; Chatterjee and Jafarov 2015; Chetverikov, Liao, and Chernozhukov 2016; Homrighausen and McDonald 2017; Feng and Yu 2019). Comparing with the corresponding theory for Lasso with fixed theoretical choice of λ , these results, however, generally require stronger regularity conditions and have less sharp bounds. There still exists an important gap between the theory and practice of Lasso. See also Wu and Wang (2020) for a recent survey on tuning parameter selection for high-dimensional regression.

The second challenge is concerned with how to properly handle heavy-tailed error contamination in high dimensions so that one achieves robustness while maintaining efficiency at the normal error setting. Heavy-tailed error contamination is

ubiquitous in high-dimensional microarray data, climate data, insurance claim data, e-commerce data, and many other applications. For such data, Gaussian or sub-Gaussian error assumption is rarely justified. Direct application of standard procedures can result in biased estimation and misleading conclusions. Heavy-tailed errors also affect the choice of tuning parameter.

It is important to note that the above two challenges are usually intertwined. Solutions focusing on only one aspect of the two issues run the risk of making the other aspect more challenging. Several authors (Fan, Li, and Wang 2017; Loh 2017; Sun, Zhou, and Fan 2020) recently made important contributions to high-dimensional M -estimation based on Huber's loss which possesses desirable robustness properties. Lozano, Meinshausen, and Yang (2016) proposed a minimum distance estimator. Wang et al. (2013) investigated robust regression based on exponential squared loss. Avella-Medina and Ronchetti (2018) studied robust penalized quasilielihood. Prasad et al. (2020) considered robust variant of gradient descent and proved theoretical guarantees. However, these work have not fully addressed the challenge of selecting λ and may even require some additional tuning parameter to achieve robustness. For example, Huber's loss function contains an additional tuning parameter and in general penalized Huber regression requires cross-validation for tuning parameter selection. An alternative robust loss function is the least absolute deviation loss (or more generally, quantile loss) (see, e.g., Belloni, Chernozhukov, and Wang 2011; Bradic, Fan, and Wang 2011; Wang, Wu, and Li 2012; Wang 2013; Fan, Fan, and Barut 2014). Although being robust, this loss may incur significant efficiency loss for normal random errors. On the other hand, an interesting stream of research has investigated how to alleviate the difficulty of selecting λ for Lasso. The scaled Lasso of Sun and Zhang (2012) iteratively estimates the regression parameter and σ . The square-root Lasso (Belloni, Chernozhukov, and Wang 2011) eliminates the need to calibrate λ for σ but does not adjust for the design matrix. TREX (Lederer and Müller 2015; Bien et al. 2016, 2018) automatically adjusts λ for both the tail of the error distribution and the design matrix but the modified loss function is no longer convex. Sabourin, Valdar, and Nobel (2015) adopted a permutation approach and Chichignoud, Lederer, and Wainwright (2016) developed a novel testing procedure to select λ . Yu and Bien (2019) proposed the organic Lasso and derived a prediction error bound under weaker assumptions on the design matrix with a theoretical choice of the tuning parameter that does not depend on σ . These, however, have not addressed the robustness challenge and may have suboptimal performance for heavy-tailed errors. For example, the theory of scaled Lasso requires the Gaussian error assumption. Estimating σ is particularly challenging for high-dimensional regression with heavy-tailed errors.

This article makes a useful contribution to the high-dimensional regression literature by showing that a rather simple solution exists that addresses these two challenges simultaneously. Our new procedure is inspired by Jaeckel's dispersion function with Wilcoxon scores (Jaeckel 1972), which plays an important role in classical nonparametric statistics due to its robustness and efficiency properties. In the low-dimensional setting ($p < n$), regression with Wilcoxon loss function was investigated by Wang and Li (2009), Wang, Kai,

and Li (2009), Leng (2010), Feng, Zou, and Wang (2012), among others. However, they required computationally intensive tuning and did not study the theory in high dimensions as is done in this article. Here, we carefully develop a new L_1 regularized estimator based on this dispersion function with theoretical guarantees in the ultrahigh-dimensional setting. We demonstrate that the new estimator achieves several goals simultaneously.

- The new estimator is convenient to implement. It is the solution to a convex optimization problem and can be obtained by linear programming. Its tuning parameter λ can be easily simulated and automatically adjusts to both the random error distribution and the design matrix correlation structure without sacrificing the severity of vanilla assumptions.
- Theoretically, we derive a nonasymptotic L_2 estimation error bound for the L_1 regularized new estimator with simulated tuning parameter in ultrahigh dimensions under mild regularity conditions. Let q be the number of nonzero regression coefficients of the underlying model. We prove that with high probability the L_2 estimation error bound achieves the rate $O(\sqrt{q \log p/n})$, the same near-oracle bound for Lasso with the theoretical tuning parameter (see Theorem 1 in Section 2.3). However, we do not need the sub-Gaussian error assumption as Lasso does or the lower-order moment assumption as Huber-loss based procedures require.
- For random errors with distributions symmetric about zero, our modeling parameter $\mathbf{x}_i^T \boldsymbol{\beta}_0$ coincides with the conditional mean. However, we do not require the symmetric random error assumption. For iid random errors, since Jaeckel's dispersion function is invariant to a location change, $\mathbf{x}_i^T \boldsymbol{\beta}_0$ differs from the conditional mean only by a constant. In particular, the slopes in $\boldsymbol{\beta}_0$ still bear the same interpretation as the effects of the covariates on the conditional mean. This is different from some alternative robust methods which may imply an altered interpretation on what parameters are directly estimated due to the modified loss function.

The new estimator is very close to Lasso for normal random errors and is robust with substantial efficiency gain for heavy-tailed errors. Figure 1, corresponding to Example 1 in Section 4, provides an example of the robust and efficient behavior of the proposed new estimators. The left panel of Figure 1 displays the boxplots of the L_2 estimation error of four different estimators for normal error distribution; while the plots in the right panel are for the heavy-tailed mixture normal error distribution. It is worth emphasizing that our conditions on the random error distribution are much weaker than the sub-Gaussian condition and permit heavy-tailed distributions such as Cauchy distribution. Due to the strong nonsmoothness of Jaeckel's dispersion function, advanced empirical processes techniques are needed to establish the theory for the new estimator comparing with that for the L_1 regularized least squares estimator.

As another contribution of the article, we show that a second-stage enhancement can further improve the estimation efficiency due to the reduction of the bias induced by the L_1

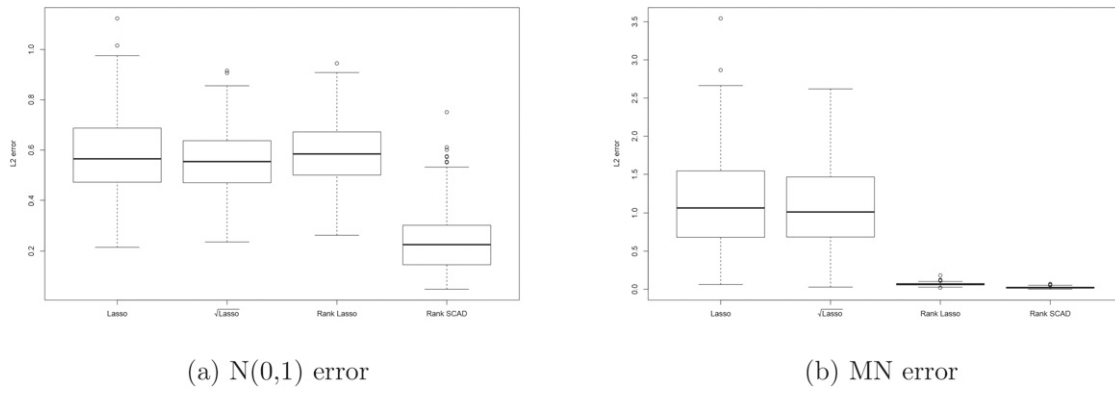


Figure 1. Boxplots for L_2 estimation error for four different procedures in Example 1.

penalty. This is motivated by the work on nonconvex penalized regression such as SCAD (Fan and Li 2001), MCP (C. H. Zhang 2010), Capped L_1 (T. Zhang 2010), which recognize that Lasso tends to overshrink large coefficients and leads to biased estimate. This second step does require some tuning but fairly light. The tuning parameter for this step can be efficiently computed via a high-dimensional BIC procedure (Chen and Chen 2008). Theoretically, we prove that with high probability the second-stage estimator possess the strong oracle property, that is, it is exactly equal to what one would obtain if the underlying data generative model is known in advance. With high probability, the zero coefficients are estimated to be exactly zero. For estimating nonzero coefficients, we derive interesting efficiency results: the resulted estimator is almost as efficient as the oracle least squares estimator for normal random errors; and can be substantially more efficient for heavy-tailed random errors. Indeed, the asymptotic relative efficiency (ARE) is shown to be the same as that of the one-sample Wilcoxon test with respect to the t -test. This implies that the ARE is as high as 0.955 for normal error distribution, and can be significantly higher than one for many heavier-tailed distributions. In particular, the efficiency of using the Jaeckel's dispersion loss function with Wilcoxon score is about 1.5 times that of using the absolute deviation (or L_1 loss) function if the random error distribution is normal. In fact, the asymptotic efficiency of Jaeckel's dispersion loss function with Wilcoxon score amounts to what would be achieved when one combines quantile losses at infinitely many quantiles, see Section 3 for more discussions. We also rigorously establish the proposed high-dimensional BIC is consistent for variable selection.

The rest of the article is organized as follows. In Section 2, we introduce the L_1 regularized new estimator based on Jaeckel's dispersion function with Wilcoxon scores and derive the nonasymptotic near-oracle L_2 estimation error bound. Section 3 introduces a second-stage enhancement for further bias reduction and efficiency improvement. It presents a strong oracle property and a consistency result for a high-dimensional BIC procedure. Monte Carlo simulations and a real data example are reported in Section 4 to demonstrate the superior performance of the proposed estimators. Section 5 concludes the article with some discussions. The Appendix provides the proofs of the main theory. The supplementary materials contain additional technical and numerical results.

2. The Methodology

2.1. Background

In model (1), all parameters are allowed to depend on the sample size n , but the dependence is suppressed in the notation for simplicity. The regression parameter β_0 is sparse in the sense that most of its components are zero.

The Lasso estimator is the solution to the regularized least squares minimization problem

$$\hat{\beta}^{\text{Lasso}}(\lambda) = \arg \min_{\beta} \left\{ (2n)^{-1} \sum_{i=1}^n (Y_i - \mathbf{x}_i^T \beta)^2 + \lambda \|\beta\|_1 \right\}, \quad (2)$$

where $\mathbf{x}_i^T = (x_{i1}, \dots, x_{ip})$ is the i th row of $\mathbf{X}$, $\|\beta\|_1$ denotes the L_1 -norm of β and λ denotes the tuning parameter. The magnitude of λ controls the complexity of the model. A larger value of λ indicates heavier shrinkage.

The optimal choice of λ involves a careful trade-off. Motivated by the Karush–Kuhn–Tucker condition for convex optimization (Boyd and Vandenberghe 2004), the general principal of tuning parameter selection for penalized regression (Bickel, Ritov, and Tsybakov 2009) suggests that λ should satisfy $\lambda \geq n^{-1} \|\mathbf{X}^T \epsilon\|_{\infty}$, where $\|\cdot\|_{\infty}$ denotes the infinity norm. As an example, if the random errors are independent $N(0, \sigma^2)$ variables and the design matrix is normalized such that each column has L_2 -norm equal to $\sqrt{n}$, then the above lower bound is satisfied with high probability by the choice $\lambda = \tau \sigma \sqrt{\log p/n}$ for some positive constant τ .

On the other hand, the theory of Lasso reveals that λ is an important factor appearing in its estimation error bound. Motivated by the subgradient condition of Lasso (Bickel, Ritov, and Tsybakov 2009; Bühlmann and van de Geer 2011), it is known that on the event

$$\left\{ 2n^{-1} \|\mathbf{X}^T \epsilon\|_{\infty} \leq \lambda \right\}. \quad (3)$$

Lasso enjoys the near-oracle error bound: $\|\hat{\beta}^{\text{Lasso}}(\lambda) - \beta_0\|_2 \leq \gamma_0 \sqrt{q} \lambda$, where $\|\cdot\|_2$ denotes the Euclidean norm of a vector, γ_0 is a constant depending on n and p only through the structure of the scaled Gram matrix $n^{-1} \mathbf{X}^T \mathbf{X}$, and q is the sparsity index or the number of nonzero coefficients in β_0 . This suggests it is desirable to choose a small λ such that the event (3) holds with high probability.

2.2. The New Method

We will study the following L_1 regularized estimator of β_0 :

$$\hat{\beta}(\lambda) = \arg \min_{\beta \in \mathcal{R}^p} \left\{ [n(n-1)]^{-1} \sum_{i \neq j} \sum_{k=1}^p |(Y_i - \mathbf{x}_i^T \beta) - (Y_j - \mathbf{x}_j^T \beta)| + \lambda \sum_{k=1}^p |\beta_k| \right\}. \quad (4)$$

The loss function in (4) has its origin in classical nonparametric statistics (see, e.g., Hettmansperger and McKean 1998 for an introduction). Minimizing this loss is equivalent to minimizing Jaeckel's dispersion function (Jaeckel 1972) with Wilcoxon scores:

$$\sqrt{12} \sum_{i=1}^n \left[\frac{R(Y_i - \mathbf{x}_i^T \beta)}{n+1} - \frac{1}{2} \right] (Y_i - \mathbf{x}_i^T \beta),$$

where $R(Y_i - \mathbf{x}_i^T \beta)$ denotes the rank of $Y_i - \mathbf{x}_i^T \beta$ among $Y_1 - \mathbf{x}_1^T \beta, \dots, Y_n - \mathbf{x}_n^T \beta$. This article emphasizes other interesting and important features of Jaeckel's dispersion function in the high-dimensional regression setting. The new estimator $\hat{\beta}(\lambda)$ is the solution to a convex optimization problem and can be obtained by linear programming (see Section 4.2). In terms of statistical performance, we will show that the new estimator behaves very similarly as Lasso for normal random errors and remains robust with potential significant efficiency gains under heavy-tailed errors for which the cross-validated Lasso could break down.

First, the loss function in (4) is invariant to a location change. Under weak conditions, β_0 is the minimizer of the population version of the loss function. This is different from most other robust methods which alter the least squares loss function by truncation or downweighting. The corresponding population parameter of the altered loss function may no longer be the regression parameter in the original model (see Fan, Li, and Wang 2017). This property can be easily seen by noting that the population version of our loss function can be expressed as $E|(\epsilon_i - \epsilon_j) - (\mathbf{x}_i - \mathbf{x}_j)^T(\beta - \beta_0)|$, which is minimized at β_0 , as $\epsilon_i - \epsilon_j$ has a symmetric distribution about zero whenever the random errors are independent and identically distributed (iid). In model (1), the intercept term is absorbed into ϵ_i and can be identified with an additional location constraint on ϵ_i , however, what we estimate using this loss function does not depend on what location constraint is imposed on ϵ_i . In particular, for random errors with distributions symmetric about zero, $\mathbf{x}_i^T \beta_0$ coincides with the conditional mean. However, symmetric random error distribution assumption is not required. For iid random errors, β_0 still bears the interpretation as the effects of the covariates on the conditional mean. This is different from Huber's loss function which combines L_2 and L_1 loss with an additional tuning parameter. The minimizer of Huber loss approximates β_0 when the additional tuning parameter is carefully tuned to diverge.

Second, it was noted in Parzen, Wei, and Ying (1994) in a different setting that the gradient function of our loss function is *completely pivotal*, which as we will show, leads to an appealing tuning-free property of the new estimator in the high-dimensional setting. This allows the procedure to circumvent

the difficulty of tuning parameter selection. We write

$$Q_n(\gamma) = [n(n-1)]^{-1} \sum_{i \neq j} |(\epsilon_i - \epsilon_j) - (\mathbf{x}_i - \mathbf{x}_j)^T \gamma|. \quad (5)$$

Then, the loss function in (4) is $Q_n(\beta - \beta_0)$. Denote the subgradient of $Q_n(\gamma)$ at $\gamma = \mathbf{0}$ (or equivalently $\beta = \beta_0$) by $\mathbf{S}_n = \frac{\partial Q_n(\gamma)}{\partial \gamma} \Big|_{\gamma=\mathbf{0}}$. Motivated by the general principal of tuning parameter selection discussed in Section 2.1, we choose λ such that

$$P(\lambda > c \|\mathbf{S}_n\|_\infty) \geq 1 - \alpha_0, \quad (6)$$

for a given small $\alpha_0 > 0$ and a theoretical constant $c > 1$, where c is a theoretical constant that does not depend n or p .

Direct computation yields

$$\begin{aligned} \mathbf{S}_n &= [n(n-1)]^{-1} \sum_{i \neq j} (\mathbf{x}_j - \mathbf{x}_i) \text{sign}(\epsilon_i - \epsilon_j) \\ &= -2[n(n-1)]^{-1} \sum_{j=1}^n \mathbf{x}_j \left(\sum_{i=1, i \neq j}^n \text{sign}(\epsilon_j - \epsilon_i) \right), \end{aligned}$$

where $\text{sign}(t) = 1$ if $t > 0$, $= -1$ if $t < 0$, and $= 0$ if $t = 0$. Denote $\xi_j = \sum_{i=1, i \neq j}^n \text{sign}(\epsilon_j - \epsilon_i)$, $j = 1, \dots, n$. It is important to observe that ξ_j is closely related to the rank of ϵ_j among $\{\epsilon_1, \dots, \epsilon_n\}$. Denote $\text{rank}(\epsilon_j) = r_j$, then

$$\begin{aligned} \xi_j &= \sum_{i=1, i \neq j}^n \text{sign}(\epsilon_j - \epsilon_i) = (r_j - 1) + (-1)(n - r_j) \\ &= 2r_j - (n + 1). \end{aligned}$$

Lemma 1 characterizes the distribution of the subgradient $\mathbf{S}_n$.

Lemma 1. Assume model (1) holds, then $\mathbf{S}_n = -2[n(n-1)]^{-1} \mathbf{X}^T \boldsymbol{\xi}$, where the n -dimensional random vector $\boldsymbol{\xi} = 2\mathbf{r} - (n+1)$, with $\mathbf{r}$ following the uniform distribution on the permutations of the integers $\{1, 2, \dots, n\}$.

We refer to this property of the gradient function as the *completely pivotal* property, as it does not depend on the error distribution. It is straightforward to simulate the distribution of $\|\mathbf{S}_n\|_\infty$ since $\{r_1, \dots, r_n\}$ can be simulated by generating a random permutation of the integers between 1 and n . Moreover, the simulations can be easily paralleled. For any given $c > 1$ and α_0 , we suggest to take λ equal to

$$\lambda^* = c G_{\|\mathbf{S}_n\|_\infty}^{-1}(1 - \alpha_0), \quad (7)$$

where $G_{\|\mathbf{S}_n\|_\infty}^{-1}(1 - \alpha_0)$ denotes the $(1 - \alpha_0)$ -quantile of the distribution of $\|\mathbf{S}_n\|_\infty$.

The simulated λ^* does not depend on the pre-estimation of any unknown population quantity and automatically adjusts to both the random error distribution and the structure of the design matrix $\mathbf{X}$. In the numerical studies in Section 4.1, we considered data generative models corresponding to a variety of error distributions and different covariance matrices for the distribution of $\mathbf{X}$. It is interesting to compare the above *completely pivotal* property of $\mathbf{S}_n$ with the *partial pivotal* property of square-root Lasso. The gradient function of the loss function of square-root Lasso, evaluated at β_0 , has the form

$(\sum_{i=1}^n \epsilon_i^2)^{1/2} \sum_{i=1}^n \mathbf{x}_i \epsilon_i$, which does not depend on σ but its distribution still depends on the distribution of ϵ_i . As a result, square-root Lasso circumvents the difficulty of tuning λ for σ but does not adjust to other aspects of the error distribution nor the design matrix. Moreover, Hebiri and Lederer (2013) reveals that the standard tuning parameters that do not depend on the design matrix $\mathbf{X}$ may lead to suboptimal performance of Lasso.

Remark 1. The work of Parzen, Wei, and Ying (1994) also revealed that the *completely pivotal* property also holds for the L_1 loss function. This was also recognized in Belloni, Chernozhukov, and Wang (2011) and Wang (2013) for penalized quantile regression. However, direct use of quantile loss may potentially result in significant efficiency loss for normal random errors. Indeed, the asymptotic efficiency analysis in Section 3 reveals that when the second-stage enhancement is implemented, the efficiency of using the Wilcoxon loss function is about 1.5 times that of using the absolute deviation (or L_1 loss) function if the random error distribution is normal for estimating the nonzero coefficients. Alternative robust loss functions such as Huber loss usually do not possess the pivotal property.

Remark 2. Let $\hat{\beta}(\lambda^*, \mathbf{Y}, \mathbf{X})$ be the estimator in (4) with the simulated tuning parameter λ^* in (7), given the response vector $\mathbf{Y}$ and the design matrix $\mathbf{X}$. It is easy to see $\hat{\beta}(\lambda^*, b\mathbf{Y}, \mathbf{X}) = b\hat{\beta}(\lambda^*, \mathbf{Y}, \mathbf{X})$ for any nonzero constant b . This equivariance property is convenient for coherent interpretation of results from regularized regression. This property is not shared by robust high-dimensional regression procedure based on Huber's loss function. Note that when $\mathbf{Y}$ has a scale change, the simulated tuning parameter λ^* remains the same when it is computed using the transformed data.

2.3. Near-Oracle Rate of the L_2 Error Bound

We consider the estimator $\hat{\beta}(\lambda^*)$, which is obtained by setting $\lambda = \lambda^*$ in (4). The main result of this subsection establishes that $\hat{\beta}(\lambda^*)$ enjoys the same near-oracle rate for the L_2 error bound as Lasso does when its λ is fixed at a theoretical value.

Let $A = \{j : \beta_{0j} \neq 0, j = 1, \dots, p\}$ be the index set of nonzero coefficients in β_0 . The cardinality $|A|_0 = q$ is the sparsity size of the underlying data generative model, where $\|\cdot\|_0$ denotes the L_0 norm and q is allowed to depend on the sample size n . For a given index set $B \in \{1, 2, \dots, p\}$, let $\mathbf{x}_B$ denote the p -dimensional vector that has the same coordinates as $\mathbf{x}$ on the index set B and zero coordinates on B^c . For a matrix D , we use $\xi_{\min}(D)$ and $\xi_{\max}(D)$ to denote the smallest eigenvalue and the largest eigenvalue of D , respectively.

For the constant c in (6), define $\bar{c} = \frac{c+1}{c-1}$ and consider the following cone set

$$\Gamma = \left\{ \boldsymbol{\gamma} \in \mathbb{R}^p : \|\boldsymbol{\gamma}_{B^c}\|_1 \leq \bar{c} \|\boldsymbol{\gamma}_B\|_1, B \subset \{1, 2, \dots, p\} \text{ and } \|B\|_0 \leq q \right\}.$$

We impose the following regularity conditions to facilitate our technical derivation.

(C1) (Conditions on the design) There exists a positive constant b_1 such that $|x_{ij}| \leq b_1$ for all $1 \leq i \leq n, 1 \leq j \leq p$.

The covariates are empirically centered in the sense that $n^{-1} \sum_{i=1}^n x_{ij} = 0$, for $1 \leq j \leq p$.

(C2) (Lower restricted eigenvalue condition) There exist some positive constant b_2 such that

$$\inf_{\boldsymbol{\gamma} \in \Gamma} \frac{n^{-1} \boldsymbol{\gamma}^T \mathbf{X}^T \mathbf{X} \boldsymbol{\gamma}}{\|\boldsymbol{\gamma}\|_2^2} \geq b_2.$$

(C3) (Conditions on the random error) The random errors ϵ_i are independent and identically distributed with density function $f(\cdot)$. Let $\zeta_{ij} = \epsilon_i - \epsilon_j$, $1 \leq i \neq j \leq n$. Let $F^*(\cdot)$ denote the distribution function of ζ_{ij} and let $f^*(\cdot)$ denote the corresponding probability density function. There exists a positive constants b_3 such that $f^*(t) \geq b_3$ uniformly in $\{t : |t| \leq q\sqrt{\log p/n}\}$.

The above conditions are mild for theoretical analysis of high-dimensional regression. Condition (C1) is common for fixed design regression and can be relaxed under additional technicality. For example, we can allow b_1 to diverge to ∞ at an appropriate rate and the error bound derived in this article still holds with probability approaching one. Alternatively, we can consider random designs and the results of the article hold for sub-Gaussian design matrix. The lower restricted eigenvalue condition is also standard to studying the error bound for Lasso-type estimators. The lower bound only needs to hold for $\boldsymbol{\gamma}$ in the cone set Γ .

Remark 3. It is worth emphasizing that our assumptions on the random error distribution (condition (C3)) are considerably weaker than what are usually imposed in the literature for high-dimensional regression. Existing theoretical work on Lasso often requires ϵ_i to have a sub-Gaussian distribution. Although the class of sub-Gaussian distribution is large, it excludes many commonly encountered heavy-tailed distributions. For example, the χ^2 -distribution is not sub-Gaussian. Square-root Lasso requires ϵ_i to have finite variance. Existing work on high-dimensional robust regression based on Huber loss also imposes moment conditions on ϵ_i , for example, Fan, Li, and Wang (2017) assume $E(|\epsilon_i|^k)$ is bounded for some $k \geq 2$, and Sun, Zhou, and Fan (2020) assume $E(|\epsilon_i|^{1+\delta})$ is bounded for some $\delta > 0$. These assumptions exclude heavy-tailed error distributions such as Cauchy distribution, which is not sub-Gaussian and does not have a finite mean.

Before presenting the main theorem, we first state a useful lemma.

Lemma 2.

- (i) Let $\hat{\boldsymbol{\gamma}}(\lambda) = \hat{\beta}(\lambda) - \beta_0$. For any $\lambda \geq c\|\mathbf{S}_n\|_\infty$, we have $\hat{\boldsymbol{\gamma}}(\lambda) \in \Gamma$.
- (ii) Assume condition (C1) is satisfied. There exists a universal constant $c_0 = 4\sqrt{2}b_1c$, where c is the constant in (7), such that for any positive constant $l > 1$,

$$P(c\|\mathbf{S}_n\|_\infty < lc_0\sqrt{\log p/n}) \geq 1 - 2 \exp(-l^2 - 1) \log p. \quad (8)$$

- (iii) If $p > (2/\alpha_0)^{1/3}$ where α_0 is the positive constant in (7), then $\lambda^* < 2c_0\sqrt{\log p/n}$.

Remark 4. Part (i) of Lemma 2 states that $\hat{\boldsymbol{\gamma}}(\lambda)$ belongs to the cone set Γ with an appropriate choice of λ . Part (ii) implies that with high probability $c\|\mathbf{S}_n\|_\infty$ has an upper bound $lc_0\sqrt{\log p/n}$. Part (iii) implies that under mild conditions on p , the simulated tuning parameter λ^* has an upper bound $2c_0\sqrt{\log p/n}$. It is worth emphasizing that the statement in (iii) is deterministic, which follows by combining (ii) with the observation that λ^* is defined as the $(1 - \alpha_0)$ -quantile of the distribution of $c\|\mathbf{S}_n\|_\infty$. The proof of Lemma 2 is given in the Appendix. Although the upper bound in (iii) has a simple form and is of the same order of the theoretical value of λ , it tends to be larger than the simulated tuning parameter value, which we recommend for real data applications as the latter automatically adjusts for the error distribution or the design matrix.

In the following theorem, we provide a nonasymptotic bound for the L_2 estimation error of $\hat{\boldsymbol{\beta}}(\lambda^*)$.

Theorem 1. Suppose conditions (C1)–(C3) hold. If $p > (2/\alpha_0)^{1/3}$, then the estimated L_1 -penalized Wilcoxon rank regression estimator $\hat{\boldsymbol{\beta}}(\lambda^*)$, where λ^* is defined in (7), satisfies

$$\|\hat{\boldsymbol{\beta}}(\lambda^*) - \boldsymbol{\beta}_0\|_2 \leq \frac{8(1 + \bar{c})c_0}{b_2b_3} \sqrt{\frac{q \log p}{n}} \quad (9)$$

with probability at least $1 - \alpha_0 - \exp(-2 \log p)$.

Theorem 1 proves that the L_2 error bound of $\hat{\boldsymbol{\beta}}(\lambda^*)$ achieves the same near-oracle rate as Lasso has with theoretical choice of λ . The requirement $p > (2/\alpha_0)^{1/3}$ is very weak. For $\alpha_0 = 0.01$, this amounts to requiring $p \geq 6$. The proof of Theorem 1 is given in the Appendix, and uses advanced empirical process theory techniques to overcome the challenge of the nonsmoothness of Jaeckel's dispersion function. The results are of independent interest as the techniques can be applied to handle other nonsmooth loss functions.

3. Bias Reduction and Efficiency Improvement for High-Dimensional Heavy-Tailed Data

3.1. A Second-Stage Enhancement

We next consider a second-stage enhancement by using $\hat{\boldsymbol{\beta}}(\lambda^*)$ as an initial estimator. The major goal is to further reduce the mean-squared error in the high-dimensional setting at the presence of heavy-tailed errors, comparing with standard least-squares based penalized regression procedures. This is achieved by a combination of an appropriate loss function (rank loss introduced earlier) and the use of a nonconvex penalty function. It is now widely recognized that L_1 penalty tends to over-penalize large coefficients, since the magnitude of L_1 penalty increases linearly with the magnitude of the coefficient. We prove that with high probability the second-stage estimator possesses the strong oracle property, that is, it is equal to the estimator one would obtain if the underlying model is known in advance. This implies that: (1) one can recover the support of the generative model with high probability; (2) one can estimate the nonzero coefficients more efficiently.

Let the initial estimator $\tilde{\boldsymbol{\beta}}^{(0)} = (\tilde{\beta}_1^{(0)}, \dots, \tilde{\beta}_p^{(0)})^T$ be $\hat{\boldsymbol{\beta}}(\lambda^*)$, the L_1 regularized estimator defined in (4) with simulated tun-

ing parameter λ^* . The second-stage estimator is defined as

$$\begin{aligned} \tilde{\boldsymbol{\beta}}^{(1)} = \arg \min_{\boldsymbol{\beta}} \Big\{ [n(n-1)]^{-1} \sum_{i \neq j} \sum_{k=1}^p p'_\eta(|\tilde{\beta}_k^{(0)}|) |\beta_j| \Big\} \\ \times |(Y_i - \mathbf{x}_i^T \boldsymbol{\beta}) - (Y_j - \mathbf{x}_j^T \boldsymbol{\beta})| \end{aligned} \quad (10)$$

where $p'_\eta(\cdot)$ denotes the derivative of some nonconvex penalty function $p_\eta(\cdot)$, where $\eta > 0$ is a tuning parameter. The second stage estimator is motivated by the local linear approximation algorithm of Zou and Li (2008) in the lower-dimensional case for penalized likelihood setting. The penalty function is only assumed to satisfy some general conditions. More specifically, $p_\eta(t)$ is increasing and concave for $t \in [0, +\infty)$ with a continuous derivative $p'_\eta(t)$ on $(0, +\infty)$. It has a singularity at the origin, that is, $p'_\eta(0+) > 0$. Without loss of generality, the penalty function can be standardized such that $p'_\eta(0+) = \eta$. Furthermore, there exist constants $a_1 > 0$ and $a_2 > 1$ such that $p'_\eta(t) \geq a_1\eta$, $\forall 0 < t < a_2\eta$; and $p'_\eta(t) = 0$, $\forall t > a_2\eta$. Two popular choices of nonconvex penalty functions satisfying these conditions are the SCAD penalty function (Fan and Li 2001) and the MCP penalty function (C. H. Zhang 2010). The SCAD penalty function is given by

$$\begin{aligned} p_\eta(|\beta|) = \eta |\beta| I(0 \leq |\beta| < \eta) \\ + \frac{a\eta |\beta| - (\beta^2 + \eta^2)/2}{a-1} I(\eta \leq |\beta| \leq a\eta) \\ + \frac{(a+1)\eta^2}{2} I(|\beta| > a\eta), \end{aligned}$$

for some $a > 2$. The MCP function has the form

$$p_\eta(|\beta|) = \eta \left(|\beta| - \frac{\beta^2}{2a\eta} \right) I(0 \leq |\beta| < a\eta) + \frac{a\eta^2}{2} I(|\beta| \geq a\eta),$$

for some $a > 1$. In practice, these two popular choices lead to similar performance.

3.2. Statistical Properties of Second-Stage Estimation

We first consider the property of the oracle estimator while allowing the underlying model dimension to diverge. Without loss of generality, we assume that the first q components of $\boldsymbol{\beta}_0$ are nonzero and the remaining $p - q$ components are zero. Hence, we can write $\boldsymbol{\beta}_0 = (\boldsymbol{\beta}_{01}^T, \mathbf{0}_{p-q}^T)^T$, where $\mathbf{0}_{p-q}$ denotes a $(p - q)$ -vector of zeros. Let $\mathbf{x}_{1i}$ be the subvector of $\mathbf{x}_i$ that consists its first q components, $i = 1, \dots, n$. It is assumed that $(\mathbf{x}_{1i}, Y_i)$ are in general positions (Koenker 2005) and that there is at least one continuous covariate in the true underlying model. Condition (C2) implies that $\xi_{\min}(n^{-1} \mathbf{X}_A^T \mathbf{X}_A) \geq b_2$, where $\mathbf{X}_A$ denotes the $n \times q$ matrix that consists of the columns of $\mathbf{X}$ whose indices are in A .

Let

$$L_n(\boldsymbol{\beta}_1) = \sum_{i \neq j} |(Y_i - \mathbf{x}_{1i}^T \boldsymbol{\beta}_1) - (Y_j - \mathbf{x}_{1j}^T \boldsymbol{\beta}_1)|$$

and $\tilde{\boldsymbol{\beta}}_1^{(o)} = \arg \min_{\boldsymbol{\beta}_1} L_n(\boldsymbol{\beta}_1)$. The oracle estimator for $\boldsymbol{\beta}_0$ is $\tilde{\boldsymbol{\beta}}^o = (\tilde{\boldsymbol{\beta}}_1^{(o)T}, \mathbf{0}_{p-q}^T)^T$. In other words, this is the estimator we

would obtain if the support of the generative model is known. [Lemma 3](#) provides the convergence rate of the oracle estimator when the number of nonzero coefficients $q = |A|$ is diverging with the sample size n .

Lemma 3. Assume conditions (C1) and (C3) hold and that $q = o(n)$. Then,

$$\|\widehat{\beta}_1^{(o)} - \beta_{01}\|_2 = O_p(q^{1/2}n^{-1/2}).$$

Note that if q is fixed, then this reduces to the classical rate of convergence. The proof of [Lemma 3](#) is given in the Appendix. The k th subgradient of the unpenalized loss function $[n(n-1)]^{-1} \sum_{i \neq j} |(Y_i - \mathbf{x}_i^T \beta) - (Y_j - \mathbf{x}_j^T \beta)|$ is given by

$$\begin{aligned} \delta_k(\beta) &= [n(n-1)]^{-1} \sum_{i \neq j} (x_{jk} - x_{ik}) \\ &\quad \times \text{sign}(Y_i - Y_j - (\mathbf{x}_i - \mathbf{x}_j)^T \beta) \\ &\quad - [n(n-1)]^{-1} \sum_{i \neq j} (x_{jk} - x_{ik}) \\ &\quad \times v_{ij} I(Y_i - Y_j = (\mathbf{x}_i - \mathbf{x}_j)^T \beta), \end{aligned}$$

where $v_{ij} \in [-1, 1]$, $k = 1, 2, \dots, p$. Lemma B in the supplementary materials characterizes the important properties of the gradient functions when being evaluated at the oracle estimator.

[Theorem 2](#) states that the second-stage estimator $\widehat{\beta}^{(1)}$ possesses the strong oracle property with high probability.

Theorem 2. Assume the conditions of [Theorem 1](#) are satisfied. Suppose $q = O(n^{c_1})$, $\eta = O(n^{-(1-c_2)/2})$, $\min_{1 \leq j \leq q} |\beta_{0j}| \geq bn^{-(1-c_3)/2}$, $p = \exp(n^{c_4})$ for some positive constants b and c_i ($i = 1, \dots, 4$) such that $2c_1 < c_2 < c_3 \leq 1$ and $c_1 + c_4 < c_2$. We have

$$P(\widehat{\beta}^{(1)} = \widehat{\beta}^{(o)}) \geq 1 - \alpha_0 - h_n,$$

where $h_n \rightarrow 0$ as $n \rightarrow \infty$.

Remark 5. Let $\widetilde{\beta}_1^{(1)}$ be the subvector containing the first q elements of $\widetilde{\beta}^{(1)}$. [Theorem 2](#) indicates that $\sqrt{n}(\widetilde{\beta}_1^{(1)} - \beta_{01})$ follows an asymptotic normal distribution in the case q is fixed. It follows from the theory of the classical nonparametric estimator based on Jaeckel's Wilcoxon-type dispersion function (Hettmansperger and McKean 1998) that the relative efficiency (ARE) of $\widetilde{\beta}_1^{(1)}$ with respect to the least-squares oracle for estimating β_{01} has the form $\text{ARE} = 12\sigma^2 \left(\int f^2(u) du \right)^2$, where σ^2 is the random error variance. It is worth noting that this ARE is the same as that of the one-sample Wilcoxon rank test with respect to the t -test. The ARE is as high as 0.955 for normal error distribution, and can be significantly higher than one for many heavier-tailed error distributions. For instance, ARE is 1.5 for the double exponential distribution, and is 1.9 for the t distribution with 3 degrees of freedom. For symmetric error distributions with finite Fisher information, the ARE is known to have a lower bound equal to 0.864 (see, e.g., Hettmansperger and McKean 1998, Theorem 1.7.6). Although theoretically it is possible to introduce nonconvex penalized rank-based loss with adaptive optimal weights to obtain an efficient estimator

(first-order equivalent to MLE), the weights nonetheless depend on the unknown error density function, see the discussions in Naranjo and McKean (1997) for the classical low-dimensional setting. In the high-dimensional setting, it is usually challenging to estimate the random error density function.

Remark 6. It is interesting to note that the above ARE is equivalent to that of composite quantile regression (Zou and Yuan 2008) when K , the number of quantiles, goes to infinity. The objective function of composite quantile regression involves a mixture of quantile objective functions at different quantiles (the suggested value of K for practical use is 19). As a result, besides the regression parameters one also needs to estimate K additional parameters corresponding to K different quantiles of the error distribution.

3.3. High-Dimensional BIC for Tuning Parameter Selection

As the main objective of the second step is to alleviate the bias due to the over-fitting of L_1 penalty, it is intuitive to not tune η by cross-validation which aims for prediction optimality. Motivated by Chen and Chen (2008), we propose a modified high-dimensional BIC criterion for selecting η . Let $\widetilde{\beta}_n^{(1)}(\eta) = (\widetilde{\beta}_{n1}^{(1)}(\eta), \dots, \widetilde{\beta}_{nn}^{(1)}(\eta))^T$ denote the estimator defined in (10) with the tuning parameter value η . Let $A_\eta = \{j : \widetilde{\beta}_{nj}^{(1)}(\eta) \neq 0, j = 1, \dots, p\}$ be the index set of the selected model and A_η^C be its complement; and let

$$\begin{aligned} \widehat{\beta}_\eta &= \arg \min_{\beta \in \mathbb{R}^p, \beta_{A_\eta^C} = 0} \left\{ [n(n-1)]^{-1} \sum_{i \neq j} \right. \\ &\quad \left. \times |(Y_i - \mathbf{x}_i^T \beta) - (Y_j - \mathbf{x}_j^T \beta)| \right\}. \end{aligned}$$

Define the high-dimensional BIC (HBIC) for selecting η as

$$\begin{aligned} \text{HBIC}(\eta) &= \log \left\{ \sum_{i \neq j} |(Y_i - \mathbf{x}_i^T \widehat{\beta}_\eta) - (Y_j - \mathbf{x}_j^T \widehat{\beta}_\eta)| \right\} \\ &\quad + |A_\eta| \frac{\log \log n}{n} \log p, \end{aligned} \quad (11)$$

where $|A_\eta|$ denotes the cardinality of set A_η . We select the value of η that minimizes $\text{HBIC}(\eta)$.

As we observe in [Section 4.2](#), HBIC can be computationally fast. High-dimensional BIC-type criterion for nonconvex penalized regression has been recently investigated by Chen and Chen (2008), Wang and Li (2009), Wang, Kim, and Li (2013), Lee, Noh, and Park (2014), Peng and Wang (2015), among others. For nonconvex penalized least-squares regression with fixed p , Wang, Li, and Tsai (2007) proved that cross-validation leads to a tuning parameter that would yield an over-fitted model with a positive probability. Wang, Kim, and Li (2013) established the consistency of high-dimensional BIC for penalized least squares regression. Lee, Noh, and Park (2014) established the consistency of high-dimensional BIC for penalized quantile regression. However, the results in these earlier work do not apply to our setting. To rigorously prove the consistency of the HBIC in (11), the main challenge is to establish a uniform approximation of a U -process that involves nonsmooth functions over a class of models.

Table 1. Simulation results for Example 1.

Error	Method	L1 error	L2 error	ME	FP	FN
$N(0, 0.25)$	Lasso	0.83 (0.03)	0.28 (0.00)	0.05 (0.00)	13.48 (0.46)	0 (0)
	$\sqrt{\text{Lasso}}$	0.70 (0.01)	0.26 (0.00)	0.05 (0.00)	11.58 (0.29)	0 (0)
	SCAD	0.24 (0.01)	0.14 (0.00)	0.02 (0.00)	0 (0)	0 (0)
	Rank Lasso	0.59 (0.01)	0.28 (0.00)	0.10 (0.00)	6.23 (0.23)	0 (0)
	Rank SCAD	0.16 (0.00)	0.11 (0.00)	0.01 (0.00)	0 (0)	0 (0)
$N(0, 1)$	Lasso	1.54 (0.04)	0.57 (0.01)	0.20 (0.01)	13.08 (0.43)	0 (0)
	$\sqrt{\text{Lasso}}$	1.41 (0.03)	0.54 (0.01)	0.21 (0.01)	11.46 (0.33)	0 (0)
	SCAD	0.46 (0.02)	0.28 (0.01)	0.06 (0.00)	0 (0)	0 (0)
	Rank Lasso	1.24 (0.02)	0.59 (0.01)	0.37 (0.00)	6.32 (0.23)	0 (0)
	Rank SCAD	0.41 (0.02)	0.28 (0.01)	0.04 (0.00)	0 (0)	0 (0)
$N(0, 2)$	Lasso	2.16 (0.05)	0.80 (0.01)	0.41 (0.01)	12.32 (0.31)	0 (0)
	$\sqrt{\text{Lasso}}$	2.07 (0.04)	0.78 (0.01)	0.42 (0.01)	11.52 (0.27)	0 (0)
	SCAD	0.66 (0.03)	0.38 (0.01)	0.11 (0.01)	0 (0)	0 (0)
	Rank Lasso	1.79 (0.04)	0.82 (0.01)	0.76 (0.02)	6.65 (0.23)	0 (0)
	Rank SCAD	0.81 (0.03)	0.52 (0.02)	0.10 (0.00)	0 (0)	0 (0)
MN	Lasso	3.02 (0.12)	1.12 (0.04)	0.81 (0.04)	12.00 (0.35)	0 (0)
	$\sqrt{\text{Lasso}}$	3.10 (0.11)	1.09 (0.04)	0.93 (0.05)	14.23 (0.27)	0 (0)
	SCAD	0.91 (0.05)	0.54 (0.03)	0.21 (0.02)	0.38 (0.04)	0 (0)
	Rank Lasso	0.14 (0.00)	0.07 (0.00)	0.04 (0.00)	6.85 (0.22)	0 (0)
	Rank SCAD	0.03 (0.00)	0.02 (0.00)	0.03 (0.00)	0 (0)	0 (0)
$\sqrt{2}t_4$	Lasso	3.42 (1.21)	1.18 (0.02)	0.77 (0.02)	15.52 (0.54)	0 (0)
	$\sqrt{\text{Lasso}}$	3.01 (0.06)	1.10 (0.02)	0.79 (0.02)	12.58 (0.30)	0 (0)
	SCAD	0.86 (0.03)	0.52 (0.02)	0.21 (0.01)	0 (0)	0 (0)
	Rank Lasso	2.20 (0.04)	1.02 (0.02)	1.44 (0.03)	7.39 (0.25)	0 (0)
	Rank SCAD	1.21 (0.04)	0.78 (0.03)	0.18 (0.01)	0 (0)	0 (0)
Cauchy	Lasso	7.32 (0.16)	3.16 (0.03)	10.92 (0.52)	5.84 (0.38)	2.17 (0.08)
	$\sqrt{\text{Lasso}}$	9.31 (0.20)	3.45 (0.06)	10.83 (0.52)	6.93 (0.28)	2.16 (0.08)
	SCAD	6.73 (0.11)	3.45 (0.05)	20.39 (0.39)	0.00 (0.00)	3.00 (0.00)
	Rank Lasso	2.60 (0.05)	1.27 (0.02)	3.73 (0.20)	6.00 (0.20)	0 (0)
	Rank SCAD	1.89 (0.07)	1.21 (0.04)	2.32 (0.19)	0 (0)	0 (0)

Let $\Lambda_n = \{\eta : |A_\eta| \leq k_n\}$, where $k_n > q$ represents a rough estimate of an upper bound of the sparsity size of the underlying model and is allowed to diverge to ∞ . We select the tuning parameter $\hat{\eta} = \arg \min_{\eta \in \Lambda_n} \text{HBIC}(\eta)$. [Theorem 3](#) shows that under some general regularity conditions, HBIC achieves model selection consistency. The proof of [Theorem 3](#) is given in the supplementary materials.

Theorem 3 (Consistency of HBIC). Assume the conditions of [Theorem 2](#) are satisfied, and that $k_n \log(p \vee n) = o(\sqrt{n})$.

Assume $\beta_{\min}^* \gg \max \left\{ \sqrt{\frac{\log(\log n)}{n}} \log p, \sqrt{\frac{q \log q}{n}} \right\}$, where $\beta_{\min}^* = \min\{|\beta_{0j}| : j \in A\}$. Then

$$P(A_{\hat{\eta}} = A) \rightarrow 1, \quad \text{as } n \rightarrow \infty,$$

where $A = \{j : \beta_{0j} \neq 0, j = 1, \dots, p\}$.

4. Numerical Studies

4.1. Monte Carlo Examples

This subsection summarizes the simulation results from three experiments. Additional simulation results are reported in the supplementary materials. We compare the performance of the cross-validated Lasso (standard Lasso with cross-validated choice of tuning parameter), the square root Lasso (denoted by $\sqrt{\text{Lasso}}$), SCAD (Fan and Li 2001), Rank Lasso (the L_1 regularized estimator in (4)), and Rank SCAD (the two-stage estimator in (10) with the SCAD penalty). The cross-validated Lasso is computed using the R package `glmnet`

(Friedman, Hastie, and Tibshirani 2010). The square-root Lasso is computed using the R package `flare` (Li et al. 2018). For Lasso or square root Lasso, the tuning parameter is chosen based on a 5-fold cross-validation as usually recommended in the literature. We compute the SCAD estimator using the function “`glmnet`” in R package `glmnet`. The initial estimator is selected as the Lasso estimator, by “`cv.glmnet`” function. For Rank Lasso, the tuning parameter λ^* is obtained by simulation based on 500 repetitions using (7) with $\alpha_0 = 0.10$ and $c = 1.01$.

Example 1 (Comparison under different random error distributions). We simulate random data from the regression model $Y_i = \mathbf{X}_i^T \boldsymbol{\beta}_0 + \epsilon_i$, $i = 1, \dots, n$, where $\mathbf{X}_i$ is generated from a p -dimensional multivariate normal distribution $N_p(0, \Sigma)$ and is independent of ϵ_i . In this example, we take $n = 100$, $p = 400$, and $\boldsymbol{\beta}_0 = (\sqrt{3}, \sqrt{3}, \sqrt{3}, 0, \dots, 0)^T$. The correlation matrix Σ has a compound symmetry structure: $\Sigma_{(ij)} = 0.5$ for $i \neq j$; and $\Sigma_{(ij)} = 1$ for $i = j$. We consider six different distributions for ϵ_i : (1) normal distribution with mean 0 and variance 0.25 (denoted by $N(0, 0.25)$); (2) normal distribution with mean 0 and variance 1 (denoted by $N(0, 1)$); (3) normal distribution with mean 0 and variance 2 (denoted by $N(0, 2)$); (4) mixture normal distribution $\epsilon \sim 0.95N(0, 1) + 0.05N(0, 100)$ (denoted by MN); (5) $\epsilon \sim \sqrt{2}t(4)$, where $t(4)$ denotes the t distribution with 4 degree of freedom; and (6) $\epsilon \sim \text{Cauchy}(0, 1)$, where $\text{Cauchy}(0, 1)$ denotes the standard Cauchy distribution.

[Table 1](#) summarizes the average (and standard error) of the L_1 estimation error, the L_2 estimation error, the model error

Table 2. Simulation results for *Example 2*.

Error	Method	L1 error	L2 error	ME	FP	FN
$N(0, 0.25)$	Lasso	12.91 (0.08)	2.06 (0.01)	2.69 (0.03)	46.78 (0.27)	0.73 (0.05)
	$\sqrt{\text{Lasso}}$	9.27 (0.10)	1.40 (0.01)	1.01 (0.02)	53.56 (0.29)	0.52 (0.04)
	SCAD	6.66 (0.12)	1.47 (0.02)	1.12 (0.03)	7.37 (0.27)	3.87 (0.12)
	Rank Lasso	5.81 (0.05)	1.08 (0.01)	0.71 (0.01)	33.19 (0.40)	0 (0)
	Rank SCAD	4.40 (0.05)	1.05 (0.01)	0.57 (0.01)	1.70 (0.11)	2.34 (0.05)
$N(0, 1)$	Lasso	16.48 (0.15)	2.56 (0.02)	3.65 (0.06)	50.06 (0.31)	2.48 (0.07)
	$\sqrt{\text{Lasso}}$	16.52 (0.16)	2.50 (0.02)	3.23 (0.06)	53.45 (0.36)	2.65 (0.08)
	SCAD	9.94 (0.16)	2.11 (0.03)	2.27 (0.05)	5.87 (0.27)	6.23 (0.10)
	Rank Lasso	9.15 (0.09)	1.68 (0.02)	1.78 (0.03)	32.95 (0.34)	0.27 (0.03)
	Rank SCAD	6.95 (0.09)	1.63 (0.02)	1.40 (0.03)	3.99 (0.17)	2.64 (0.06)
$N(0, 2)$	Lasso	20.31 (0.15)	3.14 (0.02)	5.11 (0.07)	50.78 (0.28)	4.01 (0.10)
	$\sqrt{\text{Lasso}}$	20.64 (0.16)	3.16 (0.02)	5.12 (0.07)	50.9 (0.33)	4.25 (0.1)
	SCAD	15.31 (0.25)	3.12 (0.04)	5.08 (0.13)	8.30 (0.30)	8.32 (0.12)
	Rank Lasso	12.61 (0.12)	2.3 (0.02)	3.37 (0.06)	33.92 (0.35)	0.59 (0.05)
	Rank SCAD	10.11 (0.13)	2.33 (0.03)	2.89 (0.07)	5.74 (0.21)	3.25 (0.08)
MN	Lasso	25.12 (0.47)	3.87 (0.07)	8.11 (0.26)	49.51 (0.37)	6.18 (0.19)
	$\sqrt{\text{Lasso}}$	25.02 (0.44)	3.88 (0.07)	7.93 (0.27)	48.19 (0.49)	6.12 (0.20)
	SCAD	20.40 (0.57)	4.03 (0.10)	8.90 (0.39)	9.05 (0.30)	10.57 (0.23)
	Rank Lasso	5.36 (0.10)	0.97 (0.02)	0.61 (0.02)	36.7 (0.48)	0 (0)
	Rank SCAD	4.62 (0.09)	1.10 (0.02)	0.64 (0.02)	1.15 (0.10)	2.53 (0.05)
$\sqrt{2}t_4$	Lasso	24.76 (0.20)	3.78 (0.03)	7.38 (0.11)	50.82 (0.32)	6.13 (0.11)
	$\sqrt{\text{Lasso}}$	24.56 (0.20)	3.77 (0.03)	7.42 (0.11)	48.68 (0.34)	6.14 (0.11)
	SCAD	20.46 (0.31)	4.07 (0.05)	8.66 (0.21)	8.81 (0.24)	10.76 (0.13)
	Rank Lasso	16.33 (0.17)	2.96 (0.03)	5.58 (0.12)	35.06 (0.36)	1.17 (0.07)
	Rank SCAD	14.38 (0.23)	3.21 (0.05)	5.61 (0.16)	8.17 (0.23)	5.30 (0.15)
Cauchy	Lasso	48.07 (0.53)	8.46 (0.14)	42.38 (1.65)	25.73 (0.98)	19.35 (0.31)
	$\sqrt{\text{Lasso}}$	45.66 (0.46)	8.15 (0.13)	44.69 (1.90)	23.12 (0.87)	19.57 (0.30)
	SCAD	36.53 (0.30)	7.40 (0.08)	249.49 (13.23)	11.19 (0.75)	21.20 (0.31)
	Rank Lasso	30.59 (0.51)	5.33 (0.08)	19.97 (0.65)	35.48 (0.35)	6.81 (0.26)
	Rank SCAD	32.92 (0.57)	6.78 (0.12)	25.7 (0.82)	8.68 (0.27)	13.81 (0.25)

(denoted by ME), number of false positive variables (FP), and number of false negative variables (FN) for the six methods based on 200 simulation runs. More specifically, for an estimator $\hat{\beta}$ from a given method, its L_1 error is $\|\hat{\beta} - \beta_0\|_1$; its L_2 error is $\|\hat{\beta} - \beta_0\|_2$; its model error is $(\hat{\beta} - \beta_0)^T \Sigma_X (\hat{\beta} - \beta_0)$ where Σ_X is the population covariance matrix of $\mathbf{X}$. FP is the number of noise covariates that are selected in the model; and FN is the number of active variables that are not selected in the model.

We have the following important observations. (1) Even for normal errors, Rank Lasso performs slightly better (particularly with respect to the false positives) compared with the cross-validated Lasso and square-root Lasso. This is probably due to the fact its tuning parameter is optimally chosen. It is worth pointing out that Rank Lasso uses the same tuning parameter (around 0.39) for all six error distributions. This choice is observed to have uniform good performance across different error distributions. SCAD and Rank SCAD perform similarly, and both outperform other methods for normal errors. (2) For heavy-tailed errors, the robustness and efficiency gain of Rank Lasso is substantial. Rank SCAD is observed to further improve the performance of Rank Lasso with respect to estimation error and variable selection performance. For example, for normal mixture error distribution, the average L_1 estimator error for cross-validated Lasso or square-root Lasso is above 3, while that of Rank Lasso is 0.14 and that of Rank SCAD is 0.03. The new procedures also yield smaller false positives and false negatives rates for heavy-tailed errors.

Example 2 (More challenging setting with a denser model and weaker signals). Here, we consider the same data generative model as in *Example 1* except that $\beta_0 = (2, 2, 2, 2, 1.75, 1.75, 1.75, 1.5, 1.5, 1.5, 1.5, 1.25, 1.25, 1.25, 1, 1, 1, 0.75, 0.75, 0.75, 0.5, 0.5, 0.5, 0.25, 0.25, 0.25, 0.25, \mathbf{0}_{p-25})^T$, where $\mathbf{0}_{p-25}$ is a $(p - 25)$ -dimensional vector of zeros. Comparing with *Example 1*, this is a considerably more challenging scenario with 25 active variables and a number of weak signals. *Table 2* summarizes the simulation results. We observe similar results as in *Example 1*. Rank Lasso improve both the estimation and prediction accuracy in all cases. Rank SCAD further improves the model selection performance.

Example 3 (Comparisons under different design matrices). We consider the same data generative model as in *Example 1* but with the following three different choices of Σ : (1) the compound symmetry correlated correlation matrix with correlation coefficient 0.8 (Σ_1); (2) the compound symmetry correlation matrix with correlation coefficient 0.2 (Σ_2); and (3) the AR(1) correlation matrix with autocorrelation coefficient 0.5 (Σ_3). For each choice of Σ , we consider three different error distributions: $N(0, 1)$, the mixture normal distribution in *Example 1*, and Cauchy distribution. The simulation results are summarized in *Table 3*. Note that *Table 1* contains the simulation results on compound symmetry correlation matrix with correlation coefficient 0.5. The tuning parameter selection of L_1 regularized new estimator automatically adapts to the design matrix. In all cases considered in this example, the new procedures display superior

Table 3. Simulation results for Example 3.

Σ	Error	Method	L1 error	L2 error	ME	FP	FN
Σ_1	$N(0, 1)$	Lasso	2.42 (0.06)	0.88 (0.02)	0.17 (0.00)	14.42 (0.46)	0 (0)
		$\sqrt{\text{Lasso}}$	2.22 (0.04)	0.83 (0.01)	0.15 (0.00)	13.09 (0.30)	0 (0)
		SCAD	0.68 (0.03)	0.40 (0.02)	0.06 (0.00)	0 (0)	0 (0)
		Rank Lasso	2.23 (0.04)	0.91 (0.01)	0.25 (0.01)	11.48 (0.28)	0 (0)
		Rank SCAD	0.55 (0.02)	0.34 (0.01)	0.04 (0.00)	0 (0)	0 (0)
	MN	Lasso	4.87 (0.18)	1.75 (0.06)	0.76 (0.04)	12.30 (0.25)	0 (0)
		$\sqrt{\text{Lasso}}$	4.79 (0.17)	1.72 (0.06)	0.75 (0.04)	12.23 (0.26)	0 (0)
		SCAD	2.87 (0.17)	1.63 (0.09)	0.78 (0.06)	0.61 (0.05)	0.56 (0.05)
		Rank Lasso	0.27 (0.01)	0.11 (0.00)	0.00 (0.00)	12.29 (0.27)	0 (0)
		Rank SCAD	0.05 (0.00)	0.04 (0.00)	0.00 (0.00)	0 (0)	0 (0)
	Cauchy	Lasso	8.70 (0.17)	3.65 (0.05)	5.33 (0.25)	5.54 (0.30)	2.73 (0.04)
		$\sqrt{\text{Lasso}}$	11.14 (0.17)	4.31 (0.09)	4.56 (0.18)	7.85 (0.23)	2.73 (0.54)
		SCAD	6.48 (0.08)	3.25 (0.03)	25.31 (0.23)	0.00 (0.00)	3.00 (0.00)
		Rank Lasso	5.11 (0.11)	2.03 (0.04)	1.27 (0.04)	11.52 (0.26)	0 (0)
		Rank SCAD	4.07 (0.17)	2.16 (0.08)	1.19 (0.01)	1.12 (0.08)	0.91 (0.05)
Σ_2	$N(0, 1)$	Lasso	1.21 (0.04)	0.45 (0.01)	0.18 (0.01)	13.08 (0.62)	0 (0)
		$\sqrt{\text{Lasso}}$	0.99 (0.02)	0.44 (0.01)	0.17 (0.00)	6.50 (0.23)	0 (0)
		SCAD	0.42 (0.01)	0.26 (0.01)	0.07 (0.00)	0 (0)	0 (0)
		Rank Lasso	0.94 (0.01)	0.55 (0.01)	0.39 (0.01)	1.34 (0.09)	0 (0)
		Rank SCAD	0.32 (0.01)	0.20 (0.01)	0.04 (0.00)	0.56 (0.07)	0 (0)
	MN	Lasso	2.49 (0.10)	0.92 (0.03)	0.82 (0.05)	11.37 (0.35)	0 (0)
		$\sqrt{\text{Lasso}}$	2.18 (0.08)	0.93 (0.03)	0.81 (0.05)	7.65 (0.23)	0 (0)
		SCAD	0.87 (0.04)	0.48 (0.02)	0.25 (0.02)	0.41 (0.04)	0 (0)
		Rank Lasso	0.11 (0.00)	0.07 (0.00)	0.01 (0.00)	1.62 (0.11)	0 (0)
		Rank SCAD	0.03 (0.00)	0.02 (0.00)	0.00 (0.00)	0.40 (0.04)	0 (0)
	Cauchy	Lasso	5.90 (0.10)	2.97 (0.01)	10.00 (0.24)	3.78 (0.32)	2.15 (0.08)
		$\sqrt{\text{Lasso}}$	12.00 (0.37)	3.75 (0.11)	11.19 (0.52)	18.63 (0.33)	1.86 (0.08)
		SCAD	6.33 (0.10)	3.21 (0.03)	13.87 (0.24)	0.00 (0.00)	3.00 (0.00)
		Rank Lasso	2.36 (0.05)	1.35 (0.03)	2.31 (0.08)	1.45 (0.09)	0 (0)
		Rank SCAD	1.15 (0.05)	0.76 (0.03)	0.56 (0.04)	0.72 (0.06)	0 (0)
Σ_3	$N(0, 1)$	Lasso	0.80 (0.03)	0.36 (0.01)	0.14 (0.00)	9.38 (0.60)	0 (0)
		$\sqrt{\text{Lasso}}$	0.71 (0.01)	0.35 (0.01)	0.12 (0.00)	4.47 (0.15)	0 (0)
		SCAD	0.48 (0.02)	0.29 (0.01)	0.08 (0.00)	0.39 (0.04)	0 (0)
		Rank Lasso	0.64 (0.01)	0.43 (0.01)	0.25 (0.01)	0 (0)	0 (0)
		Rank SCAD	0.41 (0.02)	0.25 (0.01)	0.04 (0.00)	1.11 (0.13)	0 (0)
	MN	Lasso	1.53 (0.06)	0.67 (0.02)	0.59 (0.04)	7.12 (0.33)	0 (0)
		$\sqrt{\text{Lasso}}$	1.55 (0.06)	0.68 (0.02)	0.57 (0.04)	6.32 (0.19)	0 (0)
		SCAD	1.07 (0.05)	0.60 (0.03)	0.34 (0.03)	0.45 (0.05)	0 (0)
		Rank Lasso	0.08 (0.00)	0.05 (0.00)	0.00 (0.00)	0 (0)	0 (0)
		Rank SCAD	0.04 (0.00)	0.02 (0.00)	0.00 (0.00)	0.62 (0.05)	0 (0)
	Cauchy	Lasso	4.96 (0.04)	2.69 (0.04)	11.44 (0.40)	2.57 (0.24)	1.75 (0.09)
		$\sqrt{\text{Lasso}}$	8.99 (0.37)	3.33 (0.12)	11.91 (0.67)	9.78 (0.23)	1.49 (1.17)
		SCAD	6.63 (0.13)	3.41 (0.06)	18.61 (0.45)	0.00 (0.00)	3.00 (0.00)
		Rank Lasso	1.51 (0.03)	0.99 (0.02)	1.37 (0.05)	0 (0)	0 (0)
		Rank SCAD	1.27 (0.06)	0.82 (0.04)	0.38 (0.03)	0.63 (0.05)	0 (0)

performance with notable improvement even for normal error distribution and substantial improvements for the heavy-tailed error distributions.

4.2. Computational Aspects

The new estimator in (4) is the minimizer of a convex objective function and can be conveniently solved via linear programming. With the aid of slack variables ξ_{ij}^+ , ξ_{ij}^- , and ζ_k , the convex optimization problem in (4) can be equivalently rewritten as

$$\min_{\beta, \xi, \zeta} \left\{ [n(n-1)]^{-1} \sum_{i \neq j} (\xi_{ij}^+ + \xi_{ij}^-) + \lambda \sum_{k=1}^p \zeta_k \right\}$$

$$\text{subject to } \xi_{ij}^+ - \xi_{ij}^- = (Y_i - Y_j) - (\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\beta},$$

$$i, j = 1, 2, \dots, n;$$

$$\xi_{ij}^+ \geq 0, \xi_{ij}^- \geq 0, i, j = 1, 2, \dots, n;$$

$$\zeta_k \geq \beta_k, \zeta_k \geq -\beta_k, k = 1, 2, \dots, p.$$

This is a linear programming problem and can be solved using existing optimization software packages. The second-stage estimator in (10) can be computed similarly by incorporating weights. The major computational barrier is due to the U -statistics structure of the loss function in (4), where the sum consists of $O(n^2)$ terms. An effective approach to alleviate this challenge is to adopt a resampling technique called incomplete U -statistic (Cléménçon, Colin, and Bellet 2016), which reduces the computational complexity of the loss function to $O(n)$. Our numerical results below demonstrate that this approximation scheme substantially improves the computation time without deteriorating the quality of the final estimators.

Table 4. Comparison of computation time.

Error	Method	L1 error	L2 error	ME	FP	FN	Time (sec)
$N(0, 1)$	Lasso	0.83 (0.03)	0.28 (0.00)	0.05 (0.00)	13.48 (0.46)	0 (0)	0.44
	$\sqrt{\text{Lasso}}$	0.70 (0.01)	0.26 (0.00)	0.05 (0.00)	11.58 (0.29)	0 (0)	8.08
	SCAD	0.24 (0.01)	0.14 (0.00)	0.02 (0.00)	0 (0)	0 (0)	0.93
	Rank Lasso	1.07 (0.02)	0.55 (0.01)	0.35 (0.01)	4.31 (0.19)	0 (0)	0.54
	Rank SCAD	0.47 (0.01)	0.26 (0.01)	0.05 (0.00)	0 (0)	0 (0)	3.72
Cauchy	Lasso	7.32 (0.16)	3.16 (0.03)	10.92 (0.52)	5.84 (0.38)	2.17 (0.08)	0.87
	$\sqrt{\text{Lasso}}$	9.31 (0.20)	3.45 (0.06)	10.83 (0.52)	6.93 (0.28)	2.16 (0.08)	10.20
	SCAD	6.73 (0.11)	3.45 (0.05)	20.39 (0.39)	0.00 (0.00)	3.00 (0.00)	0.79
	Rank Lasso	3.84 (0.11)	1.82 (0.05)	3.45 (0.19)	7.05 (0.25)	0 (0)	0.49
	Rank SCAD	2.86 (0.12)	1.64 (0.07)	2.41 (0.17)	0.38 (0.05)	0 (0)	3.72

Table 5. Analysis of eQTL data: results based on 100 random partitions.

Method	L_1 error	L_2 error	Model size
Lasso	0.075 (0.001)	0.011 (0.000)	19.50 (1.09)
$\sqrt{\text{Lasso}}$	0.074 (0.001)	0.011 (0.000)	19.09 (0.88)
SCAD	0.083 (0.001)	0.015 (0.000)	5.04 (0.27)
Rank Lasso	0.080 (0.001)	0.014 (0.001)	6.72 (0.27)
Rank SCAD	0.077 (0.001)	0.012 (0.001)	8.17 (0.39)

We consider the simulation setup in [Example 1](#) with $N(0, 1)$ error and Cauchy random error. The results are summarized in [Table 4](#). The last column of the table reports the computational time per simulation run (measured by seconds) for each procedure with tuning parameter computation time included. We observe that Rank Lasso has computational time comparable to cross-validated Lasso (implemented by the R package `glmnet` with default options) for normal error distribution, and can be slightly faster than cross-validated Lasso for heavy-tailed Cauchy error distribution. Furthermore, Rank SCAD can be implemented quite efficiently.

4.3. A Real Data Example

We use a genetic dataset to illustrate the performance of the new Wilcoxon rank based procedures. Scheetz et al. (2006) investigated gene regulation in the mammalian eye to identify genetic variation relevant to human eye disease based on expression quantitative trait locus (eQTL) mapping data. The dataset we analyze contains expression values on 300 probes (after pre-processing) from 120 twelve-week-old male offspring of rats. The response variable is the expression of gene TRIM32, a gene identified to be associated with human hereditary diseases of the retina, corresponding to probe 1389163_at. The sample standard deviation of TRIM32 is 0.14.

We conducted 100 random partitions of the dataset. For each partition, we randomly select 60 rats as the training data and the other 60 as the testing data. Regularized regression is fitted using the training data with the performance being evaluated on the testing dataset. [Table 5](#) summarizes the average (with the standard error in the parenthesis) of the L_1 prediction error, L_2 prediction error and model size, respectively, across 100 partitions. We observe that the new rank based procedures tends to select a sparser model with similar predictive performance comparing with Lasso and square-root Lasso.

5. Conclusions and Discussions

The article proposes a new approach for high-dimensional regression. Comparing to Lasso, the proposed L_1 regularized new estimator achieves several goals simultaneously: it keeps the convex structure for convenient computation, has a tuning parameter that can be easily simulated and automatically adapts to both the error distribution and the design matrix, and is equivariant to scale transformation of the response variable. Moreover, the L_2 estimation error bound of the new estimator achieves the same near-oracle rate as Lasso does. It has similar performance as Lasso does with normal random error distribution and can be substantially more efficient with heavy-tailed error distribution. Rank Lasso enjoys the tuning-free property in the sense that its tuning parameter selection does not depend on unknown population quantities. In particular, it does not depend on the variance of the noise. The efficiency of Rank Lasso can be further improved via a second-stage enhancement with some light tuning.

There is also different line of work on model selection on the Lasso solution path with the goal of asymptotically identifying the true model (see [Chen and Chen 2008](#); [Wang, Li, and Leng 2009](#); [Fan and Tang 2013](#), among others). These methods, however, are not robust to heavy-tailed errors.

Rather than choosing λ to control the L_∞ bound of the noises ([Bickel, Ritov, and Tsybakov 2009](#)), [Sun and Zhang \(2013\)](#) recently investigated an interesting alternative that chooses λ to control a sparse L_2 measure of the noises. For scaled Lasso, they showed that a penalty level with order smaller than $\sqrt{\log p/n}$ can still be valid and that it can lead to faster convergence rate in some settings. It will be interesting to investigate this alternative in our proposed setting in the future.

Appendix: Proofs of Theorem 1 and Theorem 2

The results in [Lemma 1](#) are straightforward. The proofs of [Lemmas 2](#) and [3](#) are given in the supplementary materials.

Proof of Theorem 1. Write $L_n(\boldsymbol{\gamma}) = Q_n(\boldsymbol{\gamma}) + \lambda^* \|\boldsymbol{\beta}_0 + \boldsymbol{\gamma}\|_1$, where $Q_n(\boldsymbol{\gamma})$ is defined in (5). Then

$$\hat{\boldsymbol{\gamma}}(\lambda^*) = \hat{\boldsymbol{\beta}}(\lambda^*) - \boldsymbol{\beta}_0 = \arg \min_{\boldsymbol{\gamma}} L_n(\boldsymbol{\gamma}).$$

It follows from [Lemma 2](#) that $P(\hat{\boldsymbol{\gamma}}(\lambda^*) \in \Gamma) \geq 1 - \alpha_0$ and $\lambda^* < 2c_0\sqrt{\log p/n}$. Let $h_n = \sqrt{n^{-1}q \log p}$. Denote

$$\Gamma^* = \{\boldsymbol{\gamma} \in R^p : \boldsymbol{\gamma} \in \Gamma, \|\boldsymbol{\gamma}\|_2 = \Delta h_n\}, \quad (\text{A.1})$$

where $\Delta = \frac{8(1+\bar{c})c_0}{b_2b_3}$ with c_0 being the universal positive constant in Lemma 2. We have

$$\begin{aligned} P(\|\widehat{\boldsymbol{\gamma}}(\lambda^*)\|_2 \leq \Delta h_n) &\geq 1 - P(\|\widehat{\boldsymbol{\gamma}}(\lambda^*)\|_2 \leq \Delta h_n | \widehat{\boldsymbol{\gamma}}(\lambda^*) \in \Gamma) \\ &\quad - P(\widehat{\boldsymbol{\gamma}}(\lambda^*) \notin \Gamma) \\ &\geq P\left(\inf_{\|\boldsymbol{\gamma}\|_2 \geq \Delta h_n, \boldsymbol{\gamma} \in \Gamma} L_n(\boldsymbol{\gamma}) > L_n(\mathbf{0})\right) - \alpha_0 \\ &\geq P\left(\inf_{\boldsymbol{\gamma} \in \Gamma^*} L_n(\boldsymbol{\gamma}) > L_n(\mathbf{0})\right) - \alpha_0, \end{aligned} \quad (\text{A.2})$$

where the second inequality is due to the convexity of $L_n(\boldsymbol{\gamma})$ (e.g., Hjort and Pollard 2011). To see this, for an arbitrary point $\mathbf{s}$ outside the L_2 ball centered at $\mathbf{0}$ with radius Δh_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 $L_n(\boldsymbol{\gamma})$ implies that $(1 - l^{-1}\Delta h_n)L_n(\mathbf{0}) + l^{-1}\Delta h_n L_n(\mathbf{s}) \geq L_n(\Delta h_n \mathbf{u})$. Hence, $l^{-1}\Delta h_n(L_n(\mathbf{s}) - L_n(\mathbf{0})) \geq L_n(\Delta h_n \mathbf{u}) - L_n(\mathbf{0}) \geq \inf_{\|\boldsymbol{\gamma}\|_2 = \Delta h_n} (L_n(\boldsymbol{\gamma}) - L_n(\mathbf{0}))$. This implies (A.2).

We first note that $\forall \boldsymbol{\gamma} \in \Gamma^*$,

$$\begin{aligned} \lambda^* \|\boldsymbol{\beta}_0 + \boldsymbol{\gamma}\|_1 - \|\boldsymbol{\beta}_0\|_1 &\leq \lambda^* \|\boldsymbol{\gamma}\|_1 = \lambda^* (\|\boldsymbol{\gamma}_A\|_1 + \|\boldsymbol{\gamma}_{A^c}\|_1) \leq \lambda^* (1 + \bar{c}) \|\boldsymbol{\gamma}_A\|_1 \\ &\leq \lambda^* (1 + \bar{c}) \sqrt{q} \|\boldsymbol{\gamma}_A\|_2 \leq \lambda^* (1 + \bar{c}) \sqrt{q} \Delta h_n \\ &\leq 2c_0(1 + \bar{c}) \Delta h_n^2, \end{aligned} \quad (\text{A.3})$$

where the second inequality follows because $\boldsymbol{\gamma} \in \Gamma^*$, and the third inequality follows by applying the Cauchy-Schwartz inequality. Let $Q(\boldsymbol{\gamma}) = E\{Q_n(\boldsymbol{\gamma})\}$. We have

$$\begin{aligned} &\inf_{\boldsymbol{\gamma} \in \Gamma^*} \{L_n(\boldsymbol{\gamma}) - L_n(\mathbf{0})\} \\ &\geq \inf_{\boldsymbol{\gamma} \in \Gamma^*} \{Q(\boldsymbol{\gamma}) - Q(\mathbf{0})\} - \sup_{\boldsymbol{\gamma} \in \Gamma^*} |Q_n(\boldsymbol{\gamma}) - Q_n(\mathbf{0}) - Q(\boldsymbol{\gamma}) + Q(\mathbf{0})| \\ &\quad - \lambda^* (1 + \bar{c}) \sqrt{q} \Delta h_n. \end{aligned} \quad (\text{A.4})$$

By Knight's identity (Koenker 2005),

$$\begin{aligned} Q_n(\boldsymbol{\gamma}) - Q_n(\mathbf{0}) &= [n(n-1)]^{-1} \sum_{i \neq j} (\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma} [I(\zeta_{ij} < 0) - 1/2] \\ &\quad + [n(n-1)]^{-1} \sum_{i \neq j} \int_0^{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}} [I(\zeta_{ij} \leq s) - I(\zeta_{ij} \leq 0)] ds. \end{aligned}$$

In the above expression, the first term has mean 0 and the second term is always nonnegative. Hence,

$$\begin{aligned} Q(\boldsymbol{\gamma}) - Q(\mathbf{0}) &= [n(n-1)]^{-1} \sum_{i \neq j} \int_0^{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}} [F^*(s) - F^*(0)] ds \\ &= [n(n-1)]^{-1} \sum_{i \neq j} \int_0^{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}} \\ &\quad \times [F^*(s) - F^*(0)] ds I\{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma} > 0\} \\ &\quad + [n(n-1)]^{-1} \sum_{i \neq j} \int_0^{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}} \\ &\quad \times [F^*(s) - F^*(0)] ds I\{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma} \leq 0\} \\ &= I_1 + I_2, \end{aligned}$$

where the definition of I_i , $i = 1, 2$, is clear from the context. By the mean value theorem, for some ξ_{ij} between 0 and $(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}$, we have

$$I_1 = [n(n-1)]^{-1} \sum_{i \neq j} \int_0^{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}} f^*(\xi_{ij}) s ds I\{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma} > 0\}$$

$$\geq 0.5b_3[n(n-1)]^{-1} \sum_{i \neq j} [(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}]^2 I\{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma} > 0\},$$

by condition (C3). To evaluate I_2 , we apply the transformation of variable $t = -s$ and obtain for some ξ_{ij} between 0 and $|(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}|$,

$$\begin{aligned} I_2 &= [n(n-1)]^{-1} \sum_{i \neq j} \int_0^{-(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}} \\ &\quad \times [F^*(0) - F^*(-t)] dt I\{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma} \leq 0\} \\ &\geq [n(n-1)]^{-1} \sum_{i \neq j} \int_0^{|(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}|} f^*(\xi_{ij}) t dt I\{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma} \leq 0\} \\ &\geq 0.5b_3[n(n-1)]^{-1} \sum_{i \neq j} [(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}]^2 I\{(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma} \leq 0\}. \end{aligned}$$

Hence, by condition (C2),

$$\begin{aligned} Q(\boldsymbol{\gamma}) - Q(\mathbf{0}) &\geq 0.5b_3[n(n-1)]^{-1} \sum_{i \neq j} [(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}]^2 \\ &= b_3 n^{-1} \sum_{i=1}^n \boldsymbol{\gamma}^T \mathbf{x}_i \mathbf{x}_i^T \boldsymbol{\gamma} \geq b_2 b_3 \|\boldsymbol{\gamma}\|_2^2 = b_2 b_3 \Delta^2 h_n^2. \end{aligned} \quad (\text{A.5})$$

Write $W_n(\boldsymbol{\gamma}) = \sup_{\boldsymbol{\gamma} \in \Gamma^*} |Q_n(\boldsymbol{\gamma}) - Q_n(\mathbf{0}) - Q(\boldsymbol{\gamma}) + Q(\mathbf{0})|$ and $h(\epsilon_i, \epsilon_j) = |(\epsilon_i - \epsilon_j) - (\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}| - |\epsilon_i - \epsilon_j|$. Note that

$$\begin{aligned} |h(\epsilon_i, \epsilon_j)| &\leq |(\mathbf{x}_i - \mathbf{x}_j)^T \boldsymbol{\gamma}| \leq 2b_1 \|\boldsymbol{\gamma}\|_1 \leq 2b_1(1 + \bar{c}) \|\boldsymbol{\gamma}_A\|_1 \\ &\leq 2b_1(1 + \bar{c}) \sqrt{q} \|\boldsymbol{\gamma}_A\|_2 \leq 2b_1(1 + \bar{c}) \sqrt{q} \Delta h_n. \end{aligned}$$

Hence, if we perturb one observation of the dataset, the value of $W_n(\boldsymbol{\gamma})$ changes at most $c_0(1 + \bar{c})n^{-1}\sqrt{q}\Delta h_n$, where $c_0 = 4\sqrt{2}b_1c$. By the bounded difference inequality, $\forall t > 0$,

$$P(W_n(\boldsymbol{\gamma}) - E\{W_n(\boldsymbol{\gamma})\} > t) \leq \exp\left(-\frac{2nt^2}{c_0^2(1 + \bar{c})^2 q \Delta^2 h_n^2}\right). \quad (\text{A.6})$$

Let $M_n = \lfloor n/2 \rfloor$, the smallest integer greater than or equal to $n/2$. Let σ_i , $i = 1, \dots, n$, denote a Rademacher sequence independent of $\epsilon_1, \dots, \epsilon_n$, such that $P(\sigma_i = 1) = P(\sigma_i = -1) = 1/2$. We have

$$\begin{aligned} E\{W_n(\boldsymbol{\gamma})\} &= E\left(\sup_{\boldsymbol{\gamma} \in \Gamma^*} \frac{1}{n!} \left| \sum_{\pi} M_n^{-1} \sum_{i=1}^{M_n} (h(\epsilon_{\pi(i)}, \epsilon_{\pi(M_n+i)} \right. \right. \\ &\quad \left. \left. - E\{h(\epsilon_{\pi(i)}, \epsilon_{\pi(M_n+i)})\}) \right| \right) \\ &\leq \frac{2}{n!} \sum_{\pi} E\left(\sup_{\boldsymbol{\gamma} \in \Gamma^*} \left| M_n^{-1} \sum_{i=1}^{M_n} \sigma_i h(\epsilon_{\pi(i)}, \epsilon_{\pi(M_n+i)}) \right| \right) \\ &\leq \frac{4}{n!} \sum_{\pi} E\left(\sup_{\boldsymbol{\gamma} \in \Gamma^*} \left| M_n^{-1} \sum_{i=1}^{M_n} \sigma_i (\mathbf{x}_{\pi(i)} - \mathbf{x}_{\pi(M_n+i)})^T \boldsymbol{\gamma} \right| \right) \\ &\leq \frac{4}{n!} \sum_{\pi} \sup_{\boldsymbol{\gamma} \in \Gamma^*} \|\boldsymbol{\gamma}\|_1 E\left(\left\| M_n^{-1} \sum_{i=1}^{M_n} \sigma_i \mathbf{x}_{\pi(i)} - \mathbf{x}_{\pi(M_n+i)} \right\|_{\infty} \right) \\ &\leq 4(1 + \bar{c}) \sqrt{q} \Delta h_n \frac{1}{n!} \sum_{\pi} E\left(\max_{1 \leq j \leq p} \left| M_n^{-1} \sum_{i=1}^{M_n} \sigma_i x_{\pi(i)j} - x_{\pi(M_n+i)j} \right| \right), \end{aligned}$$

where the equality is a result of Lemma A.1 of Cl  men  on, Lugosi, and Vayatis (2008) on U -statistic with the first sum taking over all

permutations π of $\{1, \dots, n\}$; the second inequality follows from the symmetrization theorem (van der Vaart and Wellner 1996); the third inequality is due to the contraction theorem (Ledoux and Talagrand 2013); while the last inequality follows because $\sup_{\gamma \in \Gamma^*} \|\gamma\|_1 \leq (1 + \bar{c})\sqrt{q}\|\gamma_A\|_2$. Applying Lemma 14.12 of Bühlmann and van de Geer (2011), we have

$$\begin{aligned} & \mathbb{E} \left(\max_{1 \leq j \leq p} \left| M_n^{-1} \sum_{i=1}^{M_n} \sigma_i(x_{\pi(i)j} - x_{\pi(M_n+i)j}) \right| \right) \\ & \leq 4b_1 \{n^{-1} \log(p+1) + \sqrt{2n^{-1} \log(p+1)}\} \\ & \leq 8b_1 \sqrt{n^{-1} \log p}, \end{aligned}$$

for all n sufficiently large. This implies $\mathbb{E}\{W_n(\gamma)\} \leq 6(1 + \bar{c})c_0\Delta h_n^2$. Now we take $t = (1 + \bar{c})c_0\Delta h_n^2$ in (A.6). This gives

$$\begin{aligned} & P(W_n(\gamma) - \mathbb{E}\{W_n(\gamma)\} > (1 + \bar{c})c_0\Delta h_n^2 \sqrt{n^{-1}q \log p}) \\ & \leq \exp(-2 \log p). \end{aligned}$$

Hence, $W_n(\gamma) \leq 7(1 + \bar{c})c_0\Delta h_n^2$ with probability at least $1 - \exp(-2 \log p)$. Combining this result with (A.4) and (A.5), we have with probability at least $1 - \exp(-2 \log p)$,

$$\begin{aligned} \inf_{\gamma \in \Gamma^*} (L_n(\gamma) - L_n(\mathbf{0})) & \geq b_2 b_3 \Delta^2 h_n^2 - 7(1 + \bar{c})c_0\Delta h_n^2 \\ & \quad - 2c_0(1 + \bar{c})\Delta h_n^2 \\ & = (b_2 b_3 \Delta - 7(1 + \bar{c})c_0)\Delta h_n^2 > 0. \end{aligned}$$

Hence, $P(\|\hat{\gamma}(\lambda^*)\|_2 \leq \Delta h_n) > 1 - \alpha_0 - \exp(-2 \log p)$. $\square$

Proof of Theorem 2. Let $\tilde{\beta}^{(0)} = (\tilde{\beta}_1^{(0)}, \dots, \tilde{\beta}_p^{(0)})^T$ be the initial estimator obtained from the L_1 penalized Wilcoxon rank regression. By Theorem 1, $\sup_{1 \leq j \leq p} |\tilde{\beta}_j^{(0)} - \beta_{0j}| \leq \tilde{b}\sqrt{q \log p/n}$ for some positive constant $\tilde{b}$ with probability at least $1 - \alpha_0 - \exp(-2 \log p)$. By the conditions of Theorem 2, we have $\sqrt{q \log p/n} = O(n^{-(1-c_1-c_4)/2})$. By Lemma B in the supplementary materials, for the oracle estimator $\hat{\beta}^{(o)}$ there exist v_{ij}^* which satisfies $v_{ij}^* = 0$ if $Y_i - Y_j \neq (\mathbf{x}_i - \mathbf{x}_j)^T \hat{\beta}^{(o)}$, and $v_{ij}^* \in [-1, 1]$ if $Y_i - Y_j = (\mathbf{x}_i - \mathbf{x}_j)^T \hat{\beta}^{(o)}$, such that for $\delta_k(\hat{\beta}^{(o)})$ with $v_{ij} = v_{ij}^*$, we have with probability approaching one, $\delta_k(\hat{\beta}^{(o)}) = 0$, $k = 1, \dots, q$; and $|\delta_k(\hat{\beta}^{(o)})| < a_1\eta$, $k = q+1, \dots, p$. Consider $\delta_k(\hat{\beta}^{(o)})$ with $v_{ij} = v_{ij}^*$. Define the following two events,

$$\begin{aligned} F_{n1} &= \{|\tilde{\beta}_j^{(0)} - \beta_{0j}| > \eta, \text{ for some } 1 \leq j \leq p\}, \\ F_{n2} &= \{|\delta_k(\hat{\beta}^{(o)})| \geq a_1\eta, \text{ for some } q+1 \leq k \leq p; \text{ or } \delta_k(\hat{\beta}^{(o)}) \neq 0 \\ & \quad \text{for some } 0 \leq k \leq q\}, \end{aligned}$$

where $\hat{\beta}^{(o)} = (\hat{\beta}_1^{(o)}, \dots, \hat{\beta}_p^{(o)})^T$ is the oracle estimator. Define $G_n = F_{n1}^c \cap F_{n2}^c$. By Theorem 1 and Lemma C in the supplementary materials, for all n sufficiently large, we have $P(G_n) \geq 1 - \alpha_0 - h_n$, where $h_n = o(1)$.

We observe, for all n sufficiently large on the event G_n , $|\tilde{\beta}_j^{(0)}| < a_2\eta$ for $q+1 \leq j \leq p$, and $|\tilde{\beta}_j^{(0)}| \geq |\beta_{0j}| - |\tilde{\beta}_j^{(0)} - \beta_{0j}| \geq a_2\eta$, for $1 \leq j \leq q$. By the assumptions on the penalty function, we have $p'_\eta(|\tilde{\beta}_j^{(0)}|) = 0$ for $1 \leq j \leq q$; and $p'_\eta(|\tilde{\beta}_j^{(0)}|) \geq a_1\eta$ for $q+1 \leq j \leq p$. Therefore, on the event G_n , the second-stage estimator $\tilde{\beta}^{(1)}$ can be expressed as the solution to the following convex optimization problem:

$$\begin{aligned} \tilde{\beta}^{(1)} &= \arg \min_{\beta} [n(n-1)]^{-1} \sum_{i \neq j} |(Y_i - \mathbf{x}_i^T \beta) - (Y_j - \mathbf{x}_j^T \beta)| \\ & \quad + \sum_{k=q+1}^p p'_\eta(|\tilde{\beta}_k^{(0)}|) |\beta_k|. \end{aligned} \quad (\text{A.7})$$

By the property of the subgradient of a convex function, we have, $\forall \beta \in \mathcal{R}^p$,

$$\begin{aligned} & [n(n-1)]^{-1} \sum_{i \neq j} |(Y_i - \mathbf{x}_i^T \beta) - (Y_j - \mathbf{x}_j^T \beta)| \\ & \geq [n(n-1)]^{-1} \sum_{i \neq j} |(Y_i - \mathbf{x}_i^T \hat{\beta}^{(o)}) - (Y_j - \mathbf{x}_j^T \hat{\beta}^{(o)})| \\ & \quad + \sum_{k=1}^p \delta_k(\hat{\beta}^{(o)}) (\beta_k - \hat{\beta}_k^{(o)}) \\ & = [n(n-1)]^{-1} \sum_{i \neq j} |(Y_i - \mathbf{x}_i^T \hat{\beta}^{(o)}) - (Y_j - \mathbf{x}_j^T \hat{\beta}^{(o)})| \\ & \quad + \sum_{k=q+1}^p \delta_k(\hat{\beta}^{(o)}) (\beta_k - \hat{\beta}_k^{(o)}). \end{aligned}$$

Hence, $\forall \beta \in \mathcal{R}^p$,

$$\begin{aligned} & \left\{ [n(n-1)]^{-1} \sum_{i \neq j} |(Y_i - \mathbf{x}_i^T \beta) - (Y_j - \mathbf{x}_j^T \beta)| \right. \\ & \quad \left. + \sum_{k=q+1}^p p'_\eta(|\tilde{\beta}_k^{(0)}|) |\beta_k| \right\} \\ & - \left\{ [n(n-1)]^{-1} \sum_{i \neq j} |(Y_i - \mathbf{x}_i^T \hat{\beta}^{(o)}) - (Y_j - \mathbf{x}_j^T \hat{\beta}^{(o)})| \right. \\ & \quad \left. + \sum_{k=q+1}^p p'_\eta(|\tilde{\beta}_k^{(0)}|) |\hat{\beta}_k^{(o)}| \right\} \\ & \geq \sum_{k=q+1}^p \left\{ p'_\eta(|\tilde{\beta}_k^{(0)}|) + \delta_k(\hat{\beta}^{(o)}) \text{sign}(\beta_k) \right\} |\beta_k| \\ & \geq \sum_{k=q+1}^p \left\{ a_1\eta - |\delta_k(\hat{\beta}^{(o)})| \right\} |\beta_k| \geq 0 \end{aligned}$$

since $\hat{\beta}_k^{(o)} = 0$ for $k = q+1, \dots, p$. The inequality is strict unless $\beta_k = 0$ for $k = q+1, \dots, p$. This implies on G_n , $\tilde{\beta}^{(1)} = (\tilde{\beta}_1^{(1)T}, \mathbf{0}_{p-q}^T)^T$ with $\tilde{\beta}_1^{(1)} = \arg \min_{\beta_1} [n(n-1)]^{-1} \sum_{i \neq j} |(Y_i - \mathbf{x}_{i1}^T \beta_1) - (Y_j - \mathbf{x}_{j1}^T \beta_1)|$. Hence, $\tilde{\beta}^{(1)}$ is the oracle estimator. $\square$

Supplementary Materials

The supplementary material contains additional technical proofs and numerical results.

Acknowledgments

The authors are indebted to the referees, the associate editor, and the co-editors for their valuable comments, which have significantly improved the article.

Funding

Wang and Wu's research was supported by NSF DMS-1712706, NSF OAC-1940160, and FRGMS-1952373. Bradic's research was supported by NSF DMS-1712481. Li's research was supported by NSF DMS 1820702, DMS 1953196, and DMS 2015539, NIDA grant P50 DA039838, R01CA229542, and R01 ES019672. The content is solely the responsibility of the authors and does not necessarily represent the official views of the NIDA, the NIH or the NSF.

References

- Avella-Medina, M., and Ronchetti, E. (2018), "Robust and Consistent Variable Selection in High-Dimensional Generalized Linear Models," *Biometrika*, 105, 31–44. [1701]
- Belloni, A., Chernozhukov, V., and Wang, L. (2011), "Square-Root Lasso: Pivotal Recovery of Sparse Signals via Conic Programming," *Biometrika*, 98, 791–806. [1701,1704]
- Bickel, P. J., Ritov, Y., and Tsybakov, A. B. (2009), "Simultaneous Analysis of Lasso and Dantzig Selector," *The Annals of Statistics*, 37, 1705–1732. [1700,1702,1710]
- Bien, J., Gaynanova, I., Lederer, J., and Müller, C. (2016), "Non-Convex Global Minimization and False Discovery Rate Control for the TREX," arXiv no. 1604.06815. [1701]
- (2018), "Prediction Error Bounds for Linear Regression With the TREX," *Test*, 28, 451–474. [1701]
- Boyd, S., and Vandenberghe, L. (2004), *Convex Optimization*, Cambridge: Cambridge University Press. [1702]
- Bradic, J., Fan, J., and Wang, W. (2011), "Penalized Composite Quasi-Likelihood for Ultrahigh Dimensional Variable Selection," *Journal of the Royal Statistical Society, Series B*, 73, 325–349. [1701]
- Bühlmann, P., and van de Geer, S. (2011), *Statistics for High-Dimensional Data: Methods, Theory and Applications*, Berlin, Heidelberg: Springer. [1700,1702,1712]
- Bunea, F., Tsybakov, A., and Wegkamp, M. (2007), "Sparsity Oracle Inequalities for the Lasso," *Electronic Journal of Statistics*, 1, 169–194. [1700]
- Candes, E., and Tao, T. (2007), "The Dantzig Selector: Statistical Estimation When p Is Much Larger Than n ," *The Annals of Statistics*, 35, 2313–2351. [1700]
- Chatterjee, S., and Jafarov, J. (2015), "Prediction Error of Cross-Validated Lasso," arXiv no. 1502.06291. [1700]
- Chen, J., and Chen, Z. (2008), "Extended Bayesian Information Criteria for Model Selection With Large Model Spaces," *Biometrika*, 95, 759–771. [1702,1706,1710]
- Chen, S. S., Donoho, D. L., and Saunders, M. A. (2001), "Atomic Decomposition by Basis Pursuit," *SIAM Review*, 43, 129–159. [1700]
- Chetverikov, D., Liao, Z., and Chernozhukov, V. (2016), "On Cross-Validated Lasso," arXiv no. 1605.02214. [1700]
- Chichignoud, M., Lederer, J., and Wainwright, M. J. (2016), "A Practical Scheme and Fast Algorithm to Tune the Lasso With Optimality Guarantees," *Journal of Machine Learning Research*, 17, 1–20. [1701]
- Cléménçon, S., Colin, I., and Bellet, A. (2016), "Scaling-Up Empirical Risk Minimization: Optimization of Incomplete U -Statistics," *The Journal of Machine Learning Research*, 17, 2682–2717. [1709]
- Cléménçon, S., Lugosi, G., and Vayatis, N. (2008), "Ranking and Empirical Minimization of U -Statistics," *The Annals of Statistics*, 36, 844–874. [1711]
- Dicker, L. H. (2014), "Variance Estimation in High-Dimensional Linear Models," *Biometrika*, 101, 269–284. [1700]
- Fan, J., Fan, Y., and Barut, E. (2014), "Adaptive Robust Variable Selection," *The Annals of Statistics*, 42, 324. [1701]
- Fan, J., Guo, S., and Hao, N. (2012), "Variance Estimation Using Refitted Cross-Validation in Ultrahigh Dimensional Regression," *Journal of the Royal Statistical Society, Series B*, 74, 37–65. [1700]
- Fan, J., Li, Q., and Wang, Y. (2017), "Estimation of High Dimensional Mean Regression in the Absence of Symmetry and Light Tail Assumptions," *Journal of the Royal Statistical Society, Series B*, 79, 247–265. [1701,1703,1704]
- Fan, J., and Li, R. (2001), "Variable Selection via Nonconcave Penalized Likelihood and Its Oracle Property," *Journal of the American Statistical Association*, 96, 1348–1360. [1700,1702,1705,1707]
- Fan, J., and Lv, J. (2010), "A Selective Overview of Variable Selection in High Dimensional Feature Space," *Statistica Sinica*, 20, 101–148. [1700]
- Fan, Y., and Tang, C. Y. (2013), "Tuning Parameter Selection in High Dimensional Penalized Likelihood," *Journal of the Royal Statistical Society, Series B*, 75, 531–552. [1710]
- Feng, L., Zou, C., and Wang, Z. (2012), "Rank-Based Inference for the Single-Index Model," *Statistics & Probability Letters*, 82, 535–541. [1701]
- Feng, Y., and Yu, Y. (2019), "The Restricted Consistency Property of Leave- n_v -Out Cross-Validation for High-Dimensional Variable Selection," *Statistica Sinica*, 29, 1607–1630. [1700]
- Friedman, J., Hastie, T., and Tibshirani, R. (2010), "Regularization Paths for Generalized Linear Models via Coordinate Descent," *Journal of Statistical Software*, 33, 1–22. [1707]
- Hebiri, M., and Lederer, J. (2013), "How Correlations Influence Lasso Prediction," *IEEE Transactions on Information Theory*, 59, 1846–1854. [1704]
- Hettmansperger, T. P., and McKean, J. W. (1998), *Robust Nonparametric Statistical Methods*, London: Arnold. [1703,1706]
- Hjort, N. L., and Pollard, D. (2011), "Asymptotics for Minimisers of Convex Processes," arXiv no. 1107.3806. [1711]
- Homrighausen, D., and McDonald, D. (2013), "The Lasso, Persistence, and Cross-Validation," in *International Conference on Machine Learning*, pp. 1031–1039. [1700]
- (2017), "Risk Consistency of Cross-Validation With Lasso-Type Procedures," *Statistica Sinica*, 27, 1017–1036. [1700]
- Jaekel, L. A. (1972), "Estimating Regression Coefficients by Minimizing the Dispersion of the Residuals," *The Annals of Mathematical Statistics*, 43, 1449–1458. [1701,1703]
- Koenker, R. (2005), *Quantile Regression*, New York: Cambridge University Press. [1705,1711]
- Lederer, J., and Müller, C. (2015), "Don't Fall for Tuning Parameters: Tuning-Free Variable Selection in High Dimensions With the TREX," in *AAAI Conference on Artificial Intelligence*, pp. 2729–2735. [1701]
- Ledoux, M., and Talagrand, M. (2013), *Probability in Banach Spaces: Isoperimetry and Processes*, Berlin, Heidelberg: Springer. [1712]
- Lee, E. R., Noh, H., and Park, B. U. (2014), "Model Selection via Bayesian Information Criterion for Quantile Regression Models," *Journal of the American Statistical Association*, 109, 216–229. [1706]
- Leng, C. (2010), "Variable Selection and Coefficient Estimation via Regularized Rank Regression," *Statistica Sinica*, 20, 167–181. [1701]
- Li, X., Zhao, T., Wang, L., Yuan, X., and Liu, H. (2018), "flare: Family of Lasso Regression," R Package Version 1.6.0. [1707]
- Loh, P.-L. (2017), "Statistical Consistency and Asymptotic Normality for High-Dimensional Robust M-Estimators," *The Annals of Statistics*, 45, 866–896. [1701]
- Lozano, A. C., Meinshausen, N., and Yang, E. (2016), "Minimum Distance Lasso for Robust High-Dimensional Regression," *Electronic Journal of Statistics*, 10, 1296–1340. [1701]
- Meinshausen, N., and Bühlmann, P. (2006), "High-Dimensional Graphs and Variable Selection With the Lasso," *The Annals of Statistics*, 34, 1436–1462. [1700]
- Naranjo, J. D., and McKean, J. W. (1997), "Rank Regression With Estimated Scores," *Statistics & Probability Letters*, 33, 209–216. [1706]
- Parzen, M., Wei, L., and Ying, Z. (1994), "A Resampling Method Based on Pivotal Estimating Functions," *Biometrika*, 81, 341–350. [1703,1704]
- Peng, B., and Wang, L. (2015), "An Iterative Coordinate Descent Algorithm for High-Dimensional Nonconvex Penalized Quantile Regression," *Journal of Computational and Graphical Statistics*, 24, 676–694. [1706]
- Prasad, A., Suggala, A. S., Balakrishnan, S., and Ravikumar, P. (2020), "Robust Estimation via Robust Gradient Estimation," *Journal of the Royal Statistical Society, Series B*, 82, 601–627. [1701]
- Sabourin, J. A., Valdar, W., and Nobel, A. B. (2015), "A Permutation Approach for Selecting the Penalty Parameter in Penalized Model Selection," *Biometrics*, 71, 1185–1194. [1701]
- Scheetz, T. E., Kim, K. Y. A., Swiderski, R. E., Philp, A. R., Braun, T. A., Knudtson, K. L., Dorrance, A. M., DiBona, G. F., Huang, J., Casavant, T. L., and Sheffield, V. C. (2006), "Regulation of gene expression in the

- mammalian eye and its relevance to eye disease," *Proceedings of the National Academy of Sciences*, 103, 14429–14434. [1710]
- Sun, Q., Zhou, W.-X., and Fan, J. (2020), "Adaptive Huber Regression," *Journal of the American Statistical Association*, 115, 254–265. [1701, 1704]
- Sun, T., and Zhang, C.-H. (2012), "Scaled Sparse Linear Regression," *Biometrika*, 99, 879–898. [1700, 1701]
- (2013), "Sparse Matrix Inversion With Scaled Lasso," *The Journal of Machine Learning Research*, 14, 3385–3418. [1710]
- Tibshirani, R. (1996), "Regression Shrinkage and Selection via the Lasso," *Journal of the Royal Statistical Society, Series B*, 58, 267–288. [1700]
- Van de Geer, S. A. (2008), "High-Dimensional Generalized Linear Models and the Lasso," *The Annals of Statistics*, 36, 614–645. [1700]
- van der Vaart, A., and Wellner, J. (1996), *Weak Convergence and Empirical Processes: With Applications to Statistics*, New York: Springer-Verlag. [1712]
- Wainwright, M. J. (2009), "Sharp Thresholds for High-Dimensional and Noisy Sparsity Recovery Using l_1 -Constrained Quadratic Programming (Lasso)," *IEEE Transactions on Information Theory*, 55, 2183–2202. [1700]
- Wang, H., Li, B., and Leng, C. (2009), "Shrinkage Tuning Parameter Selection With a Diverging Number of Parameters," *Journal of the Royal Statistical Society, Series B*, 71, 671–683. [1710]
- Wang, H., Li, R., and Tsai, C.-L. (2007), "Tuning Parameter Selectors for the Smoothly Clipped Absolute Deviation Method," *Biometrika*, 94, 553–568. [1706]
- Wang, L. (2013), "The l_1 Penalized LAD Estimator for High Dimensional Linear Regression," *Journal of Multivariate Analysis*, 120, 135–151. [1701, 1704]
- Wang, L., Kai, B., and Li, R. (2009), "Local Rank Inference for Varying Coefficient Models," *Journal of the American Statistical Association*, 104, 1631–1645. [1701]
- Wang, L., Kim, Y., and Li, R. (2013), "Calibrating Non-Convex Penalized Regression in Ultra-High Dimension," *The Annals of Statistics*, 41, 2505–2536. [1706]
- Wang, L., and Li, R. (2009), "Weighted Wilcoxon-Type Smoothly Clipped Absolute Deviation Method," *Biometrics*, 65, 564–571. [1701, 1706]
- Wang, L., Wu, Y., and Li, R. (2012), "Quantile Regression for Analyzing Heterogeneity in Ultra-High Dimension," *Journal of the American Statistical Association*, 107, 214–222. [1701]
- Wang, X., Jiang, Y., Huang, M., and Zhang, H. (2013), "Robust Variable Selection With Exponential Squared Loss," *Journal of the American Statistical Association*, 108, 632–643. [1701]
- Wu, Y., and Wang, L. (2020), "A Survey of Tuning Parameter Selection for High-Dimensional Regression," *Annual Review of Statistics and Its Application*, 7, 209–226. [1700]
- Yu, G., and Bien, J. (2019), "Estimating the Error Variance in a High-Dimensional Linear Model," *Biometrika*, 106, 533–546. [1700, 1701]
- Zhang, C. H. (2010), "Nearly Unbiased Variable Selection Under Minimax Concave Penalty," *The Annals of Statistics*, 38, 894–942. [1700, 1702, 1705]
- Zhang, C.-H., and Huang, J. (2008), "The Sparsity and Bias of the Lasso Selection in High-Dimensional Linear Regression," *The Annals of Statistics*, 36, 1567–1594. [1700]
- Zhang, C.-H., and Zhang, T. (2012), "A General Theory of Concave Regularization for High-Dimensional Sparse Estimation Problems," *Statistical Science*, 27, 576–593. [1700]
- Zhang, T. (2010), "Analysis of Multi-Stage Convex Relaxation for Sparse Regularization," *Journal of Machine Learning Research*, 11, 1081–1107. [1700, 1702]
- Zhao, P., and Yu, B. (2006), "On Model Selection Consistency of Lasso," *Journal of Machine Learning Research*, 7, 2541–2563. [1700]
- Zou, H., and Li, R. (2008), "One-Step Sparse Estimates in Nonconcave Penalized Likelihood Models," *The Annals of Statistics*, 36, 1509–1566. [1705]
- Zou, H. and Yuan, M. (2008), "Composite quantile regression and the oracle model selection theory," *The Annals of Statistics*, 36, 1108–1126. [1706]