

SORTED CONCAVE PENALIZED REGRESSION

BY LONG FENG AND CUN-HUI ZHANG¹

City University of Hong Kong and Rutgers University

The Lasso is biased. Concave penalized least squares estimation (PLSE) takes advantage of signal strength to reduce this bias, leading to sharper error bounds in prediction, coefficient estimation and variable selection. For prediction and estimation, the bias of the Lasso can be also reduced by taking a smaller penalty level than what selection consistency requires, but such smaller penalty level depends on the sparsity of the true coefficient vector. The sorted ℓ_1 penalized estimation (Slope) was proposed for adaptation to such smaller penalty levels. However, the advantages of concave PLSE and Slope do not subsume each other. We propose sorted concave penalized estimation to combine the advantages of concave and sorted penalizations. We prove that sorted concave penalties adaptively choose the smaller penalty level and at the same time benefits from signal strength, especially when a significant proportion of signals are stronger than the corresponding adaptively selected penalty levels. A local convex approximation for sorted concave penalties, which extends the local linear and quadratic approximations for separable concave penalties, is developed to facilitate the computation of sorted concave PLSE and proven to possess desired prediction and estimation error bounds. Our analysis of prediction and estimation errors requires the restricted eigenvalue condition on the design, not beyond, and provides selection consistency under a required minimum signal strength condition in addition. Thus, our results also sharpens existing results on concave PLSE by removing the upper sparse eigenvalue component of the sparse Riesz condition.

1. Introduction. The purpose of this paper is twofold. First, we provide a unified treatment of prediction, coefficient estimation and variable selection properties of concave penalized least squares estimation (PLSE) in high-dimensional linear regression under the restrictive eigenvalue (RE) condition on the design matrix. Second, we propose sorted concave PLSE to combine the advantages of concave and sorted penalties, and to prove its superior theoretical properties and computational feasibility under the RE condition. Local convex approximation (LCA) is proposed and studied as a solution for the computation of sorted concave PLSE.

Consider the linear model

$$(1.1) \quad y = X\beta^* + \varepsilon,$$

Received November 2017; revised June 2018.

¹Supported in part by NSF Grants IIS-1407939, DMS-1513378, DMS-1721495 and IIS-1741390. *MSC2010 subject classifications.* Primary 62J05, 62J07; secondary 62H12.

Key words and phrases. Penalized least squares, sorted penalties, concave penalties, Slope, local convex approximation, restricted eigenvalue, minimax rate, signal strength.

where $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_p) \in \mathbb{R}^{n \times p}$ is a design matrix, $\mathbf{y} \in \mathbb{R}^n$ is a response vector, $\boldsymbol{\varepsilon} \in \mathbb{R}^n$ is a noise vector and $\boldsymbol{\beta}^* \in \mathbb{R}^p$ is an unknown coefficient vector. For simplicity, we assume throughout the paper that the design matrix is column normalized with $\|\mathbf{x}_j\|_2^2 = n$.

Our study focuses on local and approximate solutions for the minimization of penalized loss functions of the form

$$(1.2) \quad \|\mathbf{y} - \mathbf{X}\boldsymbol{\beta}\|_2^2/(2n) + \text{Pen}(\boldsymbol{\beta})$$

with a penalty function $\text{Pen}(\cdot)$ satisfying certain minimum penalty level and maximum concavity conditions. The PLSE can be viewed as a statistical choice among local minimizers of the penalized loss.

Among PLSE methods, the Lasso [27] with the ℓ_1 penalty $\text{Pen}(\boldsymbol{\beta}) = \lambda\|\boldsymbol{\beta}\|_1$ is the most widely used and extensively studied. The Lasso is relatively easy to compute as it is a convex minimization problem, but it is well known that the Lasso is biased. A consequence of this bias is the requirement of a neighborhood stability/strong irrepresentable condition on the design matrix $\mathbf{X}$ for the selection consistency of the Lasso [16, 28, 30, 37]. Fan and Li [11] proposed a concave penalty to remove the bias of the Lasso and proved an oracle property for one of the local minimizers of the resulting penalized loss. Zhang [33] proposed a path finding algorithm PLUS for concave PLSE and proved the selection consistency of the PLUS-computed local minimizer under a rate optimal signal strength condition on the coefficients and the sparse Riesz condition (SRC) [34] on the design. The SRC, which requires bounds on both the lower and upper sparse eigenvalues of the Gram matrix and is closely related to the restricted isometry property (RIP) [8], is substantially weaker than the strong irrepresentable condition. This advantage of concave PLSE over the Lasso has since become well understood.

For prediction and coefficient estimation, the existing literature somehow presents an opposite story. Consider hard sparse coefficient vectors satisfying $|\text{supp}(\boldsymbol{\beta}^*)| \leq s$ with $\log(p/s) \asymp \log p$ and small $(s/n) \log p$. Although rate minimax error bounds were proved under the RIP and SRC for the Dantzig selector and Lasso, respectively, in [7] and [34], Bickel, Ritov and Tsybakov [5] sharpened their results by weakening the RIP and SRC to the RE condition, and van de Geer and Bühlmann [29] proved comparable prediction and ℓ_1 estimation error bounds under an even weaker compatibility or ℓ_1 RE condition. Meanwhile, rate minimax error bounds for concave PLSE still require two-sided sparse eigenvalue conditions like the SRC [12, 31, 33, 36] or a proper known upper bound for the ℓ_1 norm of the true coefficient vector [15]. It turns out that the difference between the SRC and RE conditions are quite significant as Rudelson and Zhou [23] proved that the RE condition is a consequence of a lower sparse eigenvalue condition alone. This seems to suggest a theoretical advantage of the Lasso, in addition to its relative computational simplicity, compared with concave PLSE.

Emerging from the above discussion, an interesting question is whether the RE condition alone on the design matrix is also sufficient for the above discussed

results for concave penalized prediction, coefficient estimation and variable selection, provided proper conditions on the true coefficient vector and the noise. An affirmative answer to this question, which we provide in this paper, amounts to the removal of the upper sparse eigenvalue condition on the design matrix and actually also a relaxation of the lower sparse eigenvalue condition or the restricted strong convexity (RSC) condition [17] imposed in [15]; or equivalently, to the removal of the remaining analytical advantage of the Lasso as far as error bounds for the aforementioned aims are concerned. Specifically, we prove that when the true β is sparse, concave PLSE achieves rate minimaxity in prediction and coefficient estimation under the RE condition on the design. Furthermore, the selection consistency of concave PLSE is also guaranteed under the same RE condition and an additional uniform signal strength condition on the nonzero coefficients, and these results also cover nonseparable multivariate penalties imposed on the vector β as a whole, including sorted and mixed penalties such as the spike-and-slab Lasso [22].

In addition to the above conservative prediction and estimation error bounds for the concave PLSE that are comparable with those for the Lasso in both rates and regularity conditions on the design, we also prove faster rates for concave PLSE when the signal is partially strong. A short version of this result (cf. Corollary 1 in Section 2) can be stated as follows.

THEOREM 1. *Suppose $\varepsilon \sim N(0, \sigma^2 \mathbf{I})$ in (1.1). Let $\hat{\beta}^o$ be the oracle least squares estimator of β^* based on the extra knowledge of the model $S = \text{supp}(\beta^*)$, and $s_1 = \#\{j : 0 < |\beta_j^*| < \gamma \sigma \sqrt{(2/n) \log p}\}$ with $\gamma > 1$. Then, under a restricted eigenvalue condition on the design matrix,*

$$\|X\hat{\beta} - X\hat{\beta}^o\|_2^2/n + \|\hat{\beta} - \hat{\beta}^o\|_2^2 = O_{\mathbb{P}}(\sigma^2(s_1/n) \log p),$$

where $\hat{\beta}$ is a statistical choice of a local minimizer of (1.2) with a proper concave penalty. Consequently, under the “beta-min” condition $s_1 = 0$,

$$\mathbb{P}\{\hat{\beta} = \hat{\beta}^o, \text{sgn}(\hat{\beta}) = \text{sgn}(\beta^*)\} \rightarrow 1.$$

Moreover, the solution $\hat{\beta}$ can be computed in polynomial time.

Because the prediction and ℓ_2 estimation error of the oracle $\hat{\beta}^o$ is of the order $\sigma^2 s/n$, the prediction rate for concave PLSE is $\sigma^2(s + s_1 \log p)/n$ by Theorem 1, which is of smaller order than the better known rate $\sigma^2(s/n) \log p$ for the Lasso when $s_1 \ll s$. Moreover, the ℓ_2 error bound automatically yields selection consistency for the concave PLSE under the beta-min condition. Thus, concave PLSE adaptively benefits from signal strength with no harm to the performance in the worst case scenario where all signals are just below the radar screen. This advantage of concave PLSE is known in the existing literature under the sparse Riesz and comparable conditions, but not under the RE condition as presented in this paper.

The bias of the Lasso can be also reduced by taking a smaller penalty level than those required for variable selection consistency, regardless of signal strength. In the literature, PLSE is typically studied in a standard setting at penalty level $\lambda \geq \lambda_* = (\sigma/\eta)\sqrt{(2/n)\log p}$, with $0 < \eta \leq 1$. This lower bound has been referred to as the universal penalty level. However, as the bias of the Lasso is proportional to its penalty level, rate minimaxity in prediction and coefficient estimation requires smaller $\lambda \asymp \sigma\sqrt{(2/n)\log(p/s)}$ [3, 26]. Unfortunately, this smaller penalty level depends on $s = \|\beta^*\|_0$, which is typically unknown. For the ℓ_1 penalty, a remedy for this issue is to apply the Slope or a Lepski-type procedure [3, 24]. However, it is unclear from the literature whether the same can be done with concave penalties.

We propose a class of sorted concave penalties to combine the advantages of concave and sorted penalties. This extends the Slope beyond ℓ_1 . Under an RE condition, we prove that the sorted concave PLSE inherits the benefits of both concave and sorted PLSE, namely bias reduction through signal strength and adaptation to the smaller penalty level. A short version of the resulting error bounds (cf. Corollary 3 in Section 3) can be stated as follows.

THEOREM 2. *Let $\{\varepsilon, \beta^*, s_1\}$ be as in Theorem 1 and $s = \|\beta^*\|_0$. Then, under a restricted eigenvalue condition on the design matrix,*

$$\|X\hat{\beta} - X\beta^*\|_2^2/n + \|\hat{\beta} - \beta^*\|_2^2 = O_{\mathbb{P}}(\sigma^2\{s + s_1 \log(p/s)\}/n),$$

where $\hat{\beta}$ is a certain statistical choice of an approximate local minimizer of (1.2) with a proper sorted concave penalty. Moreover, the solution $\hat{\beta}$ can be computed in polynomial time.

We recall that under the same condition, the error bound for the Slope or the Lasso with Lepski adaptation is of the order $\sigma(s/n)\log(p/s)$ [3, 24]. We would like to mention here that based on empirical evidence, the least squares estimator (LSE) after model selection by the Lasso is commonly believed to reduce the bias of the Lasso. However, to the best of our knowledge, such bias reduction has not yet been theoretically quantified without imposing strong conditions to guarantee model selection consistency. In fact, existing error bounds for LSE after Lasso [4, 25] is typically no superior to those for the Lasso.

To prove the computational feasibility of our theoretical results in polynomial time, we develop an LCA algorithm for a large class of multivariate concave PLSE to produce approximate local solutions to which our theoretical results apply. The LCA is a majorization-minimization (MM) algorithm and is closely related to the local quadratic approximation (LQA) [11] and the local linear approximation (LLA) [38] algorithms. The development of the LCA is needed as the LLA does not majorize sorted concave penalties in general. Our analysis of the LCA can be viewed as an extension of the results in [1, 12, 14, 15, 17, 31, 36] where separable penalties are considered, typically at larger penalty levels.

The rest of this paper is organized as follows. In Section 2, we develop a unified treatment of prediction, coefficient estimation and variable selection properties of concave PLSE under the RE condition at penalty levels required for variable selection consistency. In Section 3, we provide error bounds for approximate solutions with sorted or smaller penalties. In Section 4, we introduce the LCA for sorted penalties. Section 5 contains discussion and a generalization of the results in Section 2 to multivariate penalty functions. We provide the detailed technical proofs in the Supplementary Material [13].

Notation: We denote by β^* the true regression coefficient vector, $\bar{\Sigma} = X^T X/n$ the sample Gram matrix, $\mathcal{S} = \text{supp}(\beta^*)$ the support set of the coefficient vector, $s = |\mathcal{S}|$ the size of the support and $\Phi(\cdot)$ the standard Gaussian cumulative distribution function. For vectors $\mathbf{v} = (v_1, \dots, v_p)$, we denote by $\|\mathbf{v}\|_q = \sum_j (|v_j|^q)^{1/q}$ the ℓ_q norm, with $\|\mathbf{v}\|_\infty = \max_j |v_j|$ and $\|\mathbf{v}\|_0 = \#\{j : v_j \neq 0\}$, and by $\mathbf{v}^\#$ the vector with components $v_j^\#$ as the j th largest among $\{|v_1|, \dots, |v_p|\}$. For symmetric matrices $\mathbf{M}$, $\phi_{\min}(\mathbf{M})$ and $\phi_{\max}(\mathbf{M})$, respectively, denote the minimum and maximum eigenvalues. Moreover, $x_+ = \max(x, 0)$.

2. PLSE with separable concave penalties. In this section, we present our results for concave PLSE at a sufficiently high penalty level to allow selection consistency. Smaller and sorted penalties are considered in Section 3. We divide the section into three subsections to describe the collection of PLSE under consideration, conditions on the design matrix and error bounds for prediction, coefficient estimation and variable selection.

2.1. Concave PLSE. In general, separable penalty functions can be written as a sum of penalties on individual variables,

$$(2.1) \quad \text{Pen}(\mathbf{b}) = \sum_{j=1}^p \rho_j(b_j; \lambda_j).$$

When $\rho_j(\cdot)$ and λ_j are fixed across $j = 1, \dots, p$, $\text{Pen}(\mathbf{b})$ reduces to usual form

$$(2.2) \quad \text{Pen}(\mathbf{b}) = \rho(\mathbf{b}; \lambda) = \sum_{j=1}^p \rho(b_j; \lambda).$$

We assume that $\rho(t; \lambda)$ is a parametric family of concave penalties that is symmetric about 0, $\rho(t; \lambda) = \rho(-t; \lambda)$, with $\rho(0; \lambda) = 0$, and monotone, $\rho(t; \lambda) \uparrow |t|$. Moreover, we assume that $\rho(t; \lambda)$ is indexed by its penalty level in the sense of $\lambda = \dot{\rho}(0+; \lambda)$. Here, $\dot{\rho}(t; \lambda)$ is defined as any value between the left and right derivatives of $\rho(t; \lambda)$ with respect to t . We denote by $\partial \rho(t; \lambda)$ the set of all $\dot{\rho}(t; \lambda)$ and by $\partial \text{Pen}(\mathbf{b})$ the set of all possible choices of

$$\dot{\text{Pen}}(b_1, \dots, b_p)^T = (\dot{\rho}_1(b_1; \lambda_1), \dots, \dot{\rho}_p(b_p; \lambda_p))^T.$$

Unless otherwise stated, given a mathematical expression where $\dot{\text{Pen}}(\mathbf{b})$ appears, $\dot{\text{Pen}}(\mathbf{b})$ denotes the most favorable member of $\partial \text{Pen}(\mathbf{b})$ for the expression to hold. We define the concavity of $\rho(t; \lambda)$ as

$$(2.3) \quad \bar{\kappa}(t; \rho, \lambda) = \sup_{t' > t} \{ \dot{\rho}(t'; \lambda) - \dot{\rho}(t; \lambda) \} / (t - t'),$$

where the supreme is taken over all possible choices of $\dot{\rho}(t; \lambda)$ and $\dot{\rho}(t'; \lambda)$. Further, we define the overall maximum concavity of $\rho(t; \lambda)$ as

$$(2.4) \quad \bar{\kappa}(\rho) = \max_{t \geq 0, \lambda > 0} \bar{\kappa}(t; \rho, \lambda).$$

Our analysis is applicable to many popular penalty functions, such as the ℓ_1 penalty $\rho(t; \lambda) = \lambda|t|$ for the Lasso with $\bar{\kappa}(\rho) = 0$, the SCAD (smoothly clipped absolute deviation) penalty [11] with $\rho(t; \lambda) = \int_0^{|t|} \{\lambda - \bar{\kappa}(x - \lambda)_+\}_+ dx$ and $\bar{\kappa}(\rho) = \bar{\kappa}$, and the MCP (minimax concave penalty) [33] with $\rho(t; \lambda) = \int_0^{|t|} (\lambda - \bar{\kappa}x)_+ dx$ and $\bar{\kappa}(\rho) = \bar{\kappa}$.

Given a penalty (2.2), local minimizers of the penalized loss (1.2) must satisfy the following (local) Karush–Kuhn–Tucker (KKT) condition

$$(2.5) \quad \mathbf{X}_j^T (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}) / n = \dot{\rho}(\hat{\beta}_j; \lambda)$$

for a certain $\dot{\rho}(\hat{\beta}_j; \lambda) \in \partial \rho(\hat{\beta}_j; \lambda)$. Our analysis also works with approximate local minimizers of (1.2) with general penalties of the form (2.1) and a common penalty level $\lambda = \lambda_j$. Such approximate solutions can be written as

$$(2.6) \quad \mathbf{X}_j^T (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}) / n = \dot{\rho}_j(\hat{\beta}_j; \lambda) + v_j,$$

where $\mathbf{v} = (v_1, \dots, v_p)^T$ is an approximation error satisfying $\hat{\beta}_j v_j \geq 0$. This extension from (2.5), which explicitly quantifies the approximation error, brings true practical benefits as many algorithms only provide approximate solutions for computational efficiency. We note that the condition $\hat{\beta}_j v_j \geq 0$, which the approximate solution (2.6) is required to satisfy, is verifiable when the solution is computed. It requires the approximation error not to reduce the penalty. We also assume in (2.11) below that $\|\mathbf{v}\|_\infty / \lambda$ is not too large. In Section 3, we consider solutions with less restrictive approximation errors. If we confine our attention to penalty functions $\rho_j(\cdot)$ with upper bounded concavity $\bar{\kappa}(\rho_j) \leq \kappa_*$, the collection of local solutions of form (2.6) is characterized by

$$(2.7) \quad \begin{cases} (\lambda - \kappa_* |\hat{\beta}_j|)_+ \leq \text{sgn}(\hat{\beta}_j) \mathbf{X}_j^T (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}) / n \leq \lambda + \text{sgn}(\hat{\beta}_j) v_j, & \hat{\beta}_j \neq 0, \\ |\mathbf{X}_j^T (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}) / n| \leq \lambda + |v_j|, & \hat{\beta}_j = 0, \end{cases}$$

as solutions of (2.7) can be constructed with separable penalties $\text{Pen}(\mathbf{b}) = \sum_{j=1}^p \rho_j(b_j; \lambda)$ with a common penalty level λ and potentially different concavity satisfying $\bar{\kappa}(\rho_j) \leq \kappa_*$. Compared with (2.6), (2.7) is more explicit. However, unfortunately, they do not cover solutions with sorted penalties.

We study solutions of (2.7) by comparing them with an oracle coefficient vector β^o . We assume that for a certain sparse subset $\mathcal{S}$ of $\{1, \dots, p\}$, $\lambda_* > 0$ and $\eta \in [0, 1)$, the oracle β^o satisfies the following:

$$(2.8) \quad \text{supp}(\beta^o) \subseteq \mathcal{S}, \quad \|X^T(y - X^T\beta^o)/n\|_\infty < \eta\lambda_*.$$

We may take $\beta^o \in \mathbb{R}^p$ as the true coefficient vector β^* with $\mathcal{S} = \text{supp}(\beta^*)$, or the oracle LSE $\hat{\beta}^o$ given by

$$(2.9) \quad \hat{\beta}_S^o = (X_S^T X_S)^{-1} X_S^T y, \quad \hat{\beta}_{S^c}^o = \mathbf{0}.$$

When $\epsilon = y - X\beta^* \sim N(\mathbf{0}, V)$ with $\phi_{\max}(V) \vee \max_{j \leq p} \mathbf{x}_j^T V \mathbf{x}_j / n \leq \sigma^2$,

$$(2.10) \quad y - X\beta^o \sim N(\mathbf{0}, V^o) \quad \text{with} \quad \phi_{\max}(V^o) \vee \max_{j \leq p} \mathbf{x}_j^T V^o \mathbf{x}_j / n \leq \sigma^2$$

for such β^o , so that (2.8) holds with at least probability $1 - \sqrt{2/(\pi \log p)}$ for

$$\lambda_* = (\sigma/\eta) \sqrt{(2/n) \log p}.$$

Given an oracle β^o satisfying (2.8), we consider solutions of (2.7) with penalties satisfying the following conditions:

$$(2.11) \quad \lambda \geq \lambda_*, \quad \|\mathbf{v}_S\|_\infty / \lambda \leq (1 - \eta)\xi - (1 + \eta),$$

for some λ_* and η satisfying (2.8) and $\xi > 0$. Let $\mathcal{B}(\lambda_*, \kappa_*)$ be the set of all local solutions satisfying (2.7) with penalty level and approximation errors satisfying (2.11),

$$(2.12) \quad \mathcal{B}(\lambda_*, \kappa_*) = \{\hat{\beta} : (2.7) \text{ holds with } \lambda \text{ and } \mathbf{v} \text{ satisfying (2.11)}\}.$$

Given a penalty $\text{Pen}(\mathbf{b}) = \sum_{j=1}^p \rho_j(b_j; \lambda)$ satisfying $\bar{\kappa}(\rho_j) \leq \kappa_*$ and (2.11), $\mathcal{B}(\lambda_*, \kappa_*)$ contains all local minimizers of the penalized loss (1.2). Our theory is applicable to the subclass

$$(2.13) \quad \mathcal{B}_0(\lambda_*, \kappa_*) = \{\hat{\beta} : \hat{\beta} \text{ and } \mathbf{0} \text{ are connected in } \mathcal{B}(\lambda_*, \kappa_*)\}.$$

Here, $\hat{\beta}$ and $\mathbf{0}$ are connected if there exist $\hat{\beta}^{(k)} \in \mathcal{B}(\lambda_*, \kappa_*)$, $k = 1, \dots, k^*$ with penalty levels $\lambda^{(k)}$ such that for the $a_0 > 0$ in Proposition 2 below,

$$(2.14) \quad \hat{\beta}^{(0)} = \mathbf{0}, \quad \hat{\beta}^{(k^*)} = \hat{\beta}, \quad \|\hat{\beta}^{(k)} - \hat{\beta}^{(k-1)}\|_1 \leq a_0 \lambda^{(k)}.$$

This condition will be further relaxed in Section 3.

By definition, $\mathcal{B}_0(\lambda_*, \kappa_*)$ contains the set of all local solutions (2.7) computable by path following algorithms starting from the origin, with constraints on the penalty levels, concavities and approximation errors. This is a large class of statistical solutions as it includes all local solutions connected to the origin regardless of the specific algorithms used to compute the solution, and different types of penalties can be used in a single solution path. For example, the Lasso estimator belongs

to the class as it is connected to the origin through the LARS algorithm [10, 19, 20]. The SCAD and MCP solutions with $\lambda \geq \lambda_*$ and $\bar{\kappa} \leq \kappa_*$ belong to the class if they are computed by the PLUS algorithm [33] or by a continuous path following algorithm from the Lasso solution. More important, many iterative algorithms generating approximate solutions of Lasso, SCAD or MCP also belong to the class, for example, [1, 12, 31]. As $\hat{\beta} = \mathbf{0}$ is the sparsest solution, $\mathcal{B}_0(\lambda_*, \kappa_*)$ can be viewed as the sparse branch of the solution space $\mathcal{B}(\lambda_*, \kappa_*)$.

We prove that all solutions in (2.13) belong to the following cone:

$$(2.15) \quad \mathcal{C}(\mathcal{S}; \xi) = \{\mathbf{u} : \|\mathbf{u}_{\mathcal{S}^c}\|_1 \leq \xi \|\mathbf{u}_{\mathcal{S}}\|_1\},$$

when the RE below is no smaller than κ_* in (2.7).

2.2. Restricted eigenvalue condition. The RE condition, proposed in [5], is arguably the weakest available on the design to guarantee rate minimax performance in prediction and coefficient estimation for the Lasso. The RE for the ℓ_2 estimation loss can be defined as follows. For $\mathcal{S} \subset \{1, \dots, p\}$ and $\xi > 0$,

$$(2.16) \quad \text{RE}_2(\mathcal{S}; \xi) = \inf \left\{ \frac{(\mathbf{u}^T \bar{\Sigma} \mathbf{u})^{1/2}}{\|\mathbf{u}\|_2} : \mathbf{u} \in \mathcal{C}(\mathcal{S}; \xi) \right\}$$

with $\inf \emptyset = \phi_{\min}(\bar{\Sigma})$ for $\mathcal{S} = \emptyset$ and $\mathcal{C}(\mathcal{S}; \xi)$ as in (2.15). The RE condition refers to the property that $\text{RE}_2(\mathcal{S}; \xi)$ is no smaller than a certain positive constant for all design matrices under consideration. For prediction and ℓ_1 estimation, it suffices to impose a somewhat weaker compatibility condition [29]. The compatibility coefficient, also called ℓ_1 -RE [29], is defined as

$$(2.17) \quad \text{RE}_1(\mathcal{S}; \xi) = \inf \left\{ \frac{(\mathbf{u}^T \bar{\Sigma} \mathbf{u})^{1/2}}{\|\mathbf{u}_{\mathcal{S}}\|_1 / |\mathcal{S}|^{1/2}} : \mathbf{u} \in \mathcal{C}(\mathcal{S}; \xi) \right\}.$$

In addition to the RE above, we define a relaxed cone invertibility factor (RCIF) for prediction as

$$(2.18) \quad \text{RCIF}_{\text{pred}}(\mathcal{S}; \eta, \mathbf{w}) = \inf \left\{ \frac{\|\bar{\Sigma} \mathbf{u}\|_{\infty}^2 |\mathcal{S}|}{\mathbf{u}^T \bar{\Sigma} \mathbf{u}} : \|\mathbf{u}_{\mathcal{S}^c}\|_1 < -\frac{\mathbf{w}_{\mathcal{S}}^T \mathbf{u}_{\mathcal{S}}}{1 - \eta} \right\},$$

with $\eta \in [0, 1)$ and a vector $\mathbf{w} \in \mathbb{R}^p$, and a RCIF for the ℓ_q estimation as

$$(2.19) \quad \text{RCIF}_{\text{est}, q}(\mathcal{S}; \eta, \mathbf{w}) = \inf \left\{ \frac{\|\bar{\Sigma} \mathbf{u}\|_{\infty} |\mathcal{S}|^{1/q}}{\|\mathbf{u}\|_q} : \|\mathbf{u}_{\mathcal{S}^c}\|_1 < -\frac{\mathbf{w}_{\mathcal{S}}^T \mathbf{u}_{\mathcal{S}}}{1 - \eta} \right\}.$$

The RCIF is a relaxation of the cone invertibility coefficient [32] for which the constraint $\|\mathbf{u}_{\mathcal{S}^c}\|_1 < \xi \|\mathbf{u}_{\mathcal{S}}\|_1$ is imposed.

The choices of ξ , η and $\mathbf{w}$ depend on the problem under consideration, but in our analysis of (2.6), we may let η and ξ satisfy (2.8) and (2.11) and

$$(2.20) \quad \begin{aligned} \mathbf{w} &= \{\text{Pen}(\beta^o) + \mathbf{v} - \mathbf{X}^T(\mathbf{y} - \mathbf{X}\beta^o)/n\}/\lambda \\ &= (\dot{\rho}_j(\beta_j^o; \lambda) + v_j - \mathbf{X}_j^T(\mathbf{y} - \mathbf{X}\beta^o)/n, j = 1, \dots, p)^T, \end{aligned}$$

where $\text{Pen}(\cdot)$ can be any penalty function of the form (2.1) satisfying $\lambda_j = \lambda$ and $\kappa(\rho_j) \leq \kappa_*$ and (2.11). It follows that $\|\mathbf{w}_S\|_\infty \leq (1 - \eta)\xi$, so that the minimization in (2.18) and (2.19) is taken over a smaller cone than (2.15). In the presence of partial signal strength, $\|\mathbf{w}_S\|_2$ can be much smaller than $|\mathcal{S}|^{1/2}(1 - \eta)\xi$. Moreover, for selection consistency, it is feasible to have $\mathbf{w}_S = \mathbf{0}$ under a beta-min condition when $\mathbf{v} = \mathbf{0}$. In the following subsection, we use an RE condition to prove cone membership of the estimation error of the concave PLSE and the RCIF to bound the prediction and coefficient estimation errors. The following proposition, which follows from the analysis in Section 3.2 of [32], shows that the RCIF may provide sharper bounds than the RE does.

PROPOSITION 1. *If $\|\mathbf{w}_S\|_\infty \leq (1 - \eta)\xi$, then*

$$\begin{aligned} \text{RCIF}_{\text{pred}}(\mathcal{S}; \eta, \mathbf{w}) &\geq \text{RE}_1^2(\mathcal{S}; \xi)/(1 + \xi)^2, \\ (2.21) \quad \text{RCIF}_{\text{est},1}(\mathcal{S}; \eta, \mathbf{w}) &\geq \text{RE}_1^2(\mathcal{S}; \xi)/(1 + \xi)^2, \\ \text{RCIF}_{\text{est},2}(\mathcal{S}; \eta, \mathbf{w}) &\geq \text{RE}_1(\mathcal{S}; \xi) \text{RE}_2(\mathcal{S}; \xi)/(1 + \xi). \end{aligned}$$

2.3. *Error bounds.* Let $\mathcal{B}(\lambda_*, \kappa_*)$ be as in (2.12), $\mathcal{B}_0(\lambda_*, \kappa_*)$ as in (2.13), $\mathcal{C}(\mathcal{S}; \xi)$ be as in (2.15). We define a set of “good solutions” for PLSE with separable penalties as

$$(2.22) \quad \mathcal{B}_0^*(\lambda_*, \kappa_*) = \mathcal{B}_0(\lambda_*, \kappa_*) \cup \{\hat{\boldsymbol{\beta}} \in \mathcal{B}(\lambda_*, \kappa_*) : \hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o \in \mathcal{C}(\mathcal{S}; \xi)\}.$$

Here, under an RE condition on the design matrix, we provide prediction and coefficient estimation error bounds for all good solutions and prove the sign consistency for exact solutions under a “beta-min” condition.

THEOREM 3. *Suppose (2.8) holds for certain $\boldsymbol{\beta}^o \in \mathbb{R}^p$ and $\text{RE}_2^2(\mathcal{S}; \xi) \geq \kappa_*$. Let $\hat{\boldsymbol{\beta}} \in \mathcal{B}_0^*(\lambda_*, \kappa_*)$ be a solution of (2.7) with penalty level λ . Then*

$$(2.23) \quad \|\mathbf{X}\hat{\boldsymbol{\beta}} - \mathbf{X}\boldsymbol{\beta}^o\|_2^2/n \leq \min \left\{ \frac{\bar{\lambda}^2 |\mathcal{S}|}{\text{RCIF}_{\text{pred}}(\mathcal{S}; \eta, \mathbf{w})}, \frac{(\bar{\lambda} + \eta\xi\lambda_*)^2 |\mathcal{S}|}{\text{RE}_1^2(\mathcal{S}; \xi)} \right\},$$

with $\bar{\lambda} = \lambda + \eta\lambda_* + \|\mathbf{v}\|_\infty \leq \lambda - \eta\xi\lambda_*$ and the $\mathbf{w}$ in (2.20), and

$$\begin{aligned} \|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o\|_1 &\leq \min \left\{ \frac{\bar{\lambda} |\mathcal{S}|}{\text{RCIF}_{\text{est},1}(\mathcal{S}; \eta, \mathbf{w})}, \frac{(\bar{\lambda} + \eta\xi\lambda_*)(1 + \xi) |\mathcal{S}|}{\text{RE}_1^2(\mathcal{S}; \xi)} \right\}, \\ (2.24) \quad \|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o\|_2 &\leq \min \left\{ \frac{\bar{\lambda} |\mathcal{S}|^{1/2}}{\text{RCIF}_{\text{est},2}(\mathcal{S}; \eta, \mathbf{w})}, \frac{(\bar{\lambda} + \eta\xi\lambda_*) |\mathcal{S}|^{1/2}}{\text{RE}_1(\mathcal{S}; \xi) \text{RE}_2(\mathcal{S}; \xi)} \right\}, \\ \|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o\|_q &\leq \frac{\bar{\lambda} |\mathcal{S}|^{1/q}}{\text{RCIF}_{\text{est},q}(\mathcal{S}; \eta, \mathbf{w})}, \quad \forall q \geq 1. \end{aligned}$$

Compared with Theorem 3, existing results on statistical solutions of concave PLSE [12, 31, 33] cover only separable penalties of form (2.2), require stronger conditions on the design (such as upper sparse eigenvalue) and offer less explicit error bounds. While the error bounds for concave penalty in Theorem 3 match existing ones for the Lasso [5, 29] and other methods [7, 9] up to a constant factor, they also hold when $\text{RE}_1(\mathcal{S}; \xi)$ and $\text{RE}_2(\mathcal{S}; \xi)$ in (2.23) and (2.24) are replaced by their larger version with the constraint $\|\mathbf{u}_{\mathcal{S}^c}\|_1 \leq \xi \|\mathbf{u}_{\mathcal{S}}\|_1$ replaced by the more stringent $(1 - \eta)\|\mathbf{u}_{\mathcal{S}^c}\|_1 \leq -\mathbf{w}_{\mathcal{S}}^T \mathbf{u}_{\mathcal{S}}$ in their definition, as $-\mathbf{w}_{\mathcal{S}}^T \mathbf{u}_{\mathcal{S}}$ could be much smaller than $(1 - \eta)\xi \|\mathbf{u}_{\mathcal{S}}\|_1$ when a significant proportion of the signals are strong. We describe the benefit of concave PLSE over the Lasso in such scenarios in the following two theorems.

THEOREM 4. Suppose (2.8) holds for an oracle solution $\boldsymbol{\beta}^o$ of (2.6) with $\mathbf{v} = 0$ and some $\rho_j(\cdot)$. Assume that $\text{RE}_2^2(\mathcal{S}; \xi) \geq \kappa_* \geq \max_j \bar{\kappa}(\rho_j)$. Let $\hat{\boldsymbol{\beta}} \in \mathcal{B}_0^*(\lambda_*, \kappa_*)$ be a solution of (2.6) with $\mathbf{v} = 0$ and the same $\rho_j(\cdot)$. Then

$$(2.25) \quad \hat{\boldsymbol{\beta}}_{\mathcal{S}^c} = \mathbf{0} \quad \text{and} \quad \text{sgn}(\hat{\boldsymbol{\beta}}_j) \text{sgn}(\boldsymbol{\beta}_j^o) \geq 0 \quad \forall j \in \mathcal{S}.$$

If in addition $\max_{j \in \mathcal{S}} \bar{\kappa}(0; \rho_j, \lambda) < \phi_{\min}(\bar{\boldsymbol{\Sigma}}_{\mathcal{S}, \mathcal{S}})$, for example, $\bar{\boldsymbol{\Sigma}}_{\mathcal{S}, \mathcal{S}^c} \neq \mathbf{0}$ a priori, then

$$(2.26) \quad \text{sgn}(\hat{\boldsymbol{\beta}}) = \text{sgn}(\boldsymbol{\beta}^o).$$

Theorem 4 provides selection consistency of concave PLSE under the restricted eigenvalue condition, compared with the required irrepresentable condition for the Lasso [16, 28, 37]. This is also new as existing results require the stronger sparse Riesz condition [33] or other combination of lower and upper sparse eigenvalue conditions on the design [12, 31, 35] for selection consistency.

The proof of Theorem 4 unifies the analysis for prediction, estimation and variable selection, as the proof of selection consistency is simply done by inspecting the case of $\mathbf{w}_{\mathcal{S}} = 0$ in the prediction and estimation error bounds. For $\mathbf{v} = 0$, the $\boldsymbol{\beta}^o$ in (2.8) is a solution of (2.6) iff $\mathbf{w}_{\mathcal{S}} = 0$, and this is the case for the oracle LSE $\boldsymbol{\beta}^o = \hat{\boldsymbol{\beta}}^o$ in (2.9) under the beta-min condition $\min_{j \in \mathcal{S}} |\hat{\beta}_j^o| \geq \gamma\lambda$ when $\dot{\rho}_j(t; \lambda) = 0$ for all $|t| \geq \gamma\lambda$. Regarding the additional condition for the sign consistency, we note that $\text{RE}_2^2(\mathcal{S}; \xi) < \text{RE}_2^2(\mathcal{S}; 0) = \phi_{\min}(\bar{\boldsymbol{\Sigma}}_{\mathcal{S}, \mathcal{S}})$ when $\bar{\boldsymbol{\Sigma}}_{\mathcal{S}, \mathcal{S}^c} \neq \mathbf{0}$ and $\bar{\kappa}(\rho_j) \leq \kappa_* \leq \text{RE}_2^2(\mathcal{S}; \xi)$ by the RE condition.

THEOREM 5. Suppose (2.8) holds for $\boldsymbol{\beta}^o \in \mathbb{R}^p$ and $\kappa_* \leq \text{RE}_2^2(\mathcal{S}; \xi)$. Let $\hat{\boldsymbol{\beta}} \in \mathcal{B}_0^*(\lambda_*, \kappa_*)$ be a solution of (2.6) with penalty level λ and $\rho_j(\cdot)$ satisfying $\bar{\kappa}(\rho_j) \leq \kappa_*$ and $\max_j \bar{\kappa}(\hat{\beta}_j; \rho_j, \lambda) \leq (1 - 1/C_0) \text{RE}_2^2(\mathcal{S}; \xi)$. Then

$$(2.27) \quad \|\mathbf{X}\hat{\boldsymbol{\beta}} - \mathbf{X}\boldsymbol{\beta}^o\|_2^2/n \leq (C_0\lambda)^2 \sup_{\mathbf{u} \neq \mathbf{0}} \frac{[-\mathbf{w}_{\mathcal{S}}^T \mathbf{u}_{\mathcal{S}} - (1 - \eta)\|\mathbf{u}_{\mathcal{S}^c}\|_1]_+^2}{\mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u}}$$

with the $\mathbf{w}$ in (2.20), and for any seminorm $\|\cdot\|$ as a loss function

$$(2.28) \quad \|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o\| \leq C_0 \lambda \sup_{\mathbf{u} \neq \mathbf{0}} \frac{\|\mathbf{u}\|[-\mathbf{w}_S^T \mathbf{u}_S - (1-\eta)\|\mathbf{u}_{S^c}\|_1]}{\mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u}}.$$

COROLLARY 1. Let $\boldsymbol{\beta}^o = \hat{\boldsymbol{\beta}}^o$ be the oracle LSE in (2.9). Suppose (2.8) holds with high probability and other conditions of Theorem 5 hold with $C_0^2/\text{RE}_2^2(\mathcal{S}; \xi) = O(1)$ and $\lambda = (\sigma/\eta)\sqrt{(2/n)\log p}$. Let $s = |\mathcal{S}|$ and $s_1 = \|(\text{Pen}(\boldsymbol{\beta}^o) + \mathbf{v})_{\mathcal{S}}/\lambda\|_2$ with $\text{Pen}(\cdot) = \sum_j \rho_j(\cdot; \lambda)$. Then

$$\|X\hat{\boldsymbol{\beta}} - X\hat{\boldsymbol{\beta}}^o\|_2^2/n + \|\hat{\boldsymbol{\beta}} - \hat{\boldsymbol{\beta}}^o\|_2^2 + \|\hat{\boldsymbol{\beta}} - \hat{\boldsymbol{\beta}}^o\|_1^2/s = O_{\mathbb{P}}(\sigma^2/n)s_1 \log p,$$

implying $\hat{\boldsymbol{\beta}} = \hat{\boldsymbol{\beta}}^o$ when $s_1 = 0$, and for the true $\boldsymbol{\beta}^*$

$$(2.29) \quad \begin{aligned} & \|X\hat{\boldsymbol{\beta}} - X\boldsymbol{\beta}^*\|_2^2/n + \|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*\|_2^2 + \|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*\|_1^2/s \\ & = O_{\mathbb{P}}(\sigma^2/n)(s_1 \log p + s). \end{aligned}$$

For $\mathbf{v} = \mathbf{0}$ and $\rho_j(\cdot; \lambda)$ satisfying $\text{supp}(\dot{\rho}_j(t; \lambda)) \subseteq [-\gamma\lambda, \gamma\lambda]$,

$$(2.30) \quad s_1 \leq \#\{j \in \mathcal{S} : |\hat{\beta}_j^o| < \lambda\gamma\},$$

Theorem 5 and Corollary 1 demonstrate the benefits of concave PLSE, as $s_1 = s = |\mathcal{S}|$ in (2.29) for the Lasso but s_1 could be much smaller than s for concave penalties. For usual penalties with $\rho_1(\cdot) = \dots = \rho_p(\cdot)$, (2.30) holds with $\gamma = 1 + 1/\bar{\kappa}$ for the SCAD (2.4) and $\gamma = 1/\bar{\kappa}$ for the MCP (2.5).

For the exact Lasso solution with $\bar{\kappa}(\rho) = 0$, $C_0 = 1$ and $\|\mathbf{w}_S\|_{\infty} \leq 1 + \eta$, Theorem 5 yields the sharpest possible prediction and estimation error bounds based on the basic inequality $\mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u} + (1-\eta)\|\mathbf{u}_{S^c}\|_1 \leq (1+\eta)\|\mathbf{u}_S\|_1$ with $\mathbf{u} = (\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o)/\lambda$, as stated in the following corollary.

COROLLARY 2. Let $\hat{\boldsymbol{\beta}}$ be the Lasso estimator with penalty level $\lambda \geq \lambda_*$. If (2.8) holds for a coefficient vector $\boldsymbol{\beta}^o \in \mathbb{R}^p$, then

$$\frac{\|X\hat{\boldsymbol{\beta}} - X\boldsymbol{\beta}^o\|_2^2}{n(1+\eta)^2\lambda^2} \leq \sup_{\mathbf{u} \neq \mathbf{0}} \frac{\psi^2(\mathbf{u})}{\mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u}} = \max_{0 < t < 1} \frac{|\mathcal{S}|(1-t)^2}{\text{RE}_1^2(\mathcal{S}; t\xi)}$$

with $\psi(\mathbf{u}) = [\|\mathbf{u}_S\|_1 - \|\mathbf{u}_{S^c}\|_1/\xi]_+$ and $\xi = (1+\eta)/(1-\eta)$,

$$\frac{\|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o\|_2}{(1+\eta)\lambda} \leq \sup_{\mathbf{u} \neq \mathbf{0}} \frac{\|\mathbf{u}\|_2 \psi(\mathbf{u})}{\mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u}} = \max_{0 < t < 1} \frac{|\mathcal{S}|^{1/2}(1-t)}{\text{RE}_{1,2}^2(\mathcal{S}; t\xi)}$$

with $\text{RE}_{1,2}(\mathcal{S}; \xi) = \inf_{\|\mathbf{u}_{S^c}\| < \xi \|\mathbf{u}_S\|_1} \mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u} / (\|\mathbf{u}\|_2 \|\mathbf{u}_S\|_1 / |\mathcal{S}|^{1/2})$, and

$$\frac{\|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o\|_1}{(1+\eta)\lambda} \leq \sup_{\mathbf{u} \neq \mathbf{0}} \frac{\|\mathbf{u}\|_1 \psi(\mathbf{u})}{\mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u}} = \max_{0 < t < 1} \frac{|\mathcal{S}|^{1/2}(1+t\xi)(1-t)}{\text{RE}_1^2(\mathcal{S}; t\xi)}.$$

As Theorems 3–5 deal with the same estimator under the same RE conditions on the design, they give a unified treatment of the prediction, coefficient estimation and variable selection performance of the PLSE, including the ℓ_1 and concave penalties. For prediction and coefficient estimation, (2.23) and (2.24) match those of state-of-art for the Lasso in both the convergence rate and the regularity condition on the design, while (2.27), (2.28) and Corollary 1 demonstrate the advantages of concave penalization when s_1 is much smaller than s . Meanwhile, for selection consistency, Theorem 4 weakens existing conditions on the design to the same RE condition as required for ℓ_2 estimation with the Lasso. These RE-based results are significant as the existing theory for concave penalization, which requires substantially stronger conditions on the design, leaves a false impression that the Lasso has a technical advantage in prediction and parameter estimation by requiring much weaker conditions on the design than the concave PLSE.

The following lemma, which can be viewed as a basic inequality for analyzing concave PLSE, is the beginning point of our analysis.

LEMMA 1. *Let $\hat{\beta}$ be as in (2.7), β^o as in (2.8), $\mathbf{h} = \hat{\beta} - \beta^o$, $\mathbf{w}$ as in (2.20), and $\bar{\lambda} = \lambda + \eta\lambda_* + \|\mathbf{v}\|_\infty$ as in Theorem 3. Then*

$$(2.31) \quad \mathbf{h}^T \bar{\Sigma} \mathbf{h} \leq \min\{-\lambda \mathbf{h}^T \mathbf{w} + \kappa_* \|\mathbf{h}\|_2^2, \bar{\lambda} \|\mathbf{h}_S\|_1 + \eta\lambda_* \|\mathbf{h}_{S^c}\|_1\}.$$

Moreover, for a proper choice of $\dot{\text{Pen}}(\beta^o) \in \partial \text{Pen}(\beta^o)$ in (2.20),

$$(2.32) \quad \mathbf{h}^T \bar{\Sigma} \mathbf{h} + \{\lambda - \|X_{S^c}^T(\mathbf{y} - X\beta^o)/n\|_\infty\} \|\mathbf{h}_{S^c}\|_1 \leq -\lambda \mathbf{h}_S^T \mathbf{w}_S + \kappa_* \|\mathbf{h}\|_2^2.$$

Next, we prove that the solutions in $\mathcal{B}_0^*(\lambda_*, \kappa_*)$ in (2.22) and other approximate local solutions (2.7) in $\mathcal{B}(\lambda_*, \kappa_*)$ are separated by a gap of size $a_0\lambda_*$ in the ℓ_1 distance for some $a_0 > 0$. Consequently, $\hat{\beta} \in \mathcal{B}_0(\lambda_*, \kappa_*)$ implies $\hat{\beta} - \beta^o \in \mathcal{C}(\mathcal{S}; \xi)$ for the $\mathcal{B}_0(\lambda_*, \kappa_*)$ in (2.13) and $\mathcal{C}(\mathcal{S}; \xi)$ in (2.15).

PROPOSITION 2. *Let $\mathcal{B}(\lambda_*, \kappa_*)$ as in (2.12) and $\mathcal{C}(\mathcal{S}; \xi)$ as in (2.15). Suppose $\text{RE}_2^2(\mathcal{S}; \xi) \geq \kappa_*$. Let $a_1 = \eta - \|\mathbf{z}_{S^c}\|_\infty/\lambda_*$, $a_2 = a_1\xi/[2\max_j \bar{\kappa}(\tilde{\beta}_j; \rho_j, \tilde{\lambda}) \times (\xi + 1)]$ and $a_3 = a_1(1 - \eta)\xi/\{(1 - \eta)(\xi + 1) + a_1\}$. Let $\hat{\beta} \in \mathcal{B}(\lambda_*, \kappa_*)$ with penalty level λ , and $\tilde{\beta} \in \mathcal{B}(\lambda_*, \kappa_*)$ with $\tilde{\lambda}$. Then*

$$\hat{\beta} - \beta^o \in \mathcal{C}(\mathcal{S}; \xi) \quad \text{and} \quad \|\tilde{\beta} - \hat{\beta}\|_1 \leq \tilde{\lambda}a_0 \quad \text{imply} \quad \tilde{\beta} - \beta^o \in \mathcal{C}(\mathcal{S}; \xi),$$

with $a_0 = \min[a_2, a_2a_3/\{(1 - \eta)(\xi \vee 1)\}]$. Consequently,

$$\mathcal{B}_0^*(\lambda_*, \kappa_*) = \{\hat{\beta} \in \mathcal{B}(\lambda_*, \kappa_*) : \hat{\beta} - \beta^o \in \mathcal{C}(\mathcal{S}; \xi)\}.$$

It follows from Proposition 2 that our theoretical results are applicable to statistical choices of approximate local solutions (2.7) computable through a discrete path of solutions $\hat{\beta}^{(t)}$ satisfying $\|\hat{\beta}^{(t)} - \hat{\beta}^{(t-1)}\|_1 \leq a_0\lambda^{(t)}$ and beginning from $\mathbf{0}$ or the Lasso solution, as $\{\mathbf{0}\}$ and the Lasso solution both satisfy the condition $\hat{\beta} - \beta^o \in \mathcal{C}(\mathcal{S}; \xi)$.

3. Sorted and smaller penalties. We have studied in Section 2 exact solutions (2.5) and approximate solutions (2.6) and (2.7) for no smaller penalty level than λ_* in the event where λ_* is a strict upper bound of the supreme norm of $\mathbf{z} = \mathbf{X}^T(\mathbf{y} - \mathbf{X}\boldsymbol{\beta}^o)/n$ as in (2.8). Such penalty or threshold levels are commonly used in the literature to study regularized methods in high-dimensional regression, but this is conservative and may yield poor numerical results. Under the normality assumption (2.10), (2.8) requires $\lambda \geq \lambda_* = (\sigma/\eta)\sqrt{(2/n)\log p}$, but the rate optimal penalty level is the smaller $\lambda \geq \sigma\sqrt{(2/n)\log(p/s)}$ for prediction and coefficient estimation with $s = |\mathcal{S}|$. It is known that for $\log(p/s) \ll \log p$, rate optimal performance in prediction and coefficient estimation can be guaranteed by the Lasso with the nonadaptive smaller penalty level depending on s [3, 26] or the Slope to achieve adaptation to the smaller penalty level [3, 24]. However, it is unclear from the literature whether the same can be done with concave penalties to also take advantage of signal strength, and under what conditions on the design. Moreover, for computational considerations, it is desirable to relax the condition imposed in Section 2 on the approximation error. In this section, we consider approximate solutions for general sorted concave penalties and unsorted nonadaptive smaller penalties. This section is divided into four subsections to discuss respectively sorted penalties, collection of PLSE with sorted and smaller penalties under consideration, a relaxed RE condition and error bounds for prediction and coefficient estimation.

3.1. Sorted concave penalties. Given a sequence of sorted penalty levels $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p \geq 0$, the sorted ℓ_1 penalty [24] is defined as

$$(3.1) \quad \text{Pen}(\mathbf{b}) = \sum_{j=1}^p \lambda_j b_j^\#,$$

where $b_j^\#$ is the j th largest value among $|b_1|, \dots, |b_p|$.

Here, we extend the sorted penalty beyond ℓ_1 . Given a family of univariate penalty functions $\rho(t; \lambda)$ and a vector $\boldsymbol{\lambda} = (\lambda_1, \dots, \lambda_p)^T$ with nonincreasing nonnegative elements, we define the associated sorted penalty as

$$(3.2) \quad \text{Pen}(\mathbf{b}) = \rho_\#(\mathbf{b}; \boldsymbol{\lambda}) = \sum_{j=1}^p \rho(b_j^\#; \lambda_j).$$

Although (3.2) seems to be a superficial extension of (3.1), it brings upon potentially significant benefits and its properties are nontrivial. We say that the sorted penalty is concave if $\rho(t; \lambda)$ is concave in t in $[0, \infty)$. We will prove that under an RE condition, the sorted concave penalty inherits the benefits of both the concave and sorted penalties, namely bias reduction for strong signal components and adaptation of the penalty level to the unknown sparsity of $\boldsymbol{\beta}$.

Penalty level and concavity of univariate penalty functions are well understood as we briefly described below (2.2). However, for nonseparable penalties such as

sorted concave penalties, their penalty level and concavity need to be studied carefully. This is done in Section 5. Here, we provide the results for sorted concave penalties in the following proposition. For the sorted penalty (3.1), we define its subdifferential, denoted by $\partial \text{Pen}(\mathbf{b})$ or $\partial \rho_{\#}(\mathbf{b}; \boldsymbol{\lambda})$, as the set of all vectors $\mathbf{g}$ satisfying

$$(3.3) \quad \begin{cases} g_{k_j} = \dot{\rho}(b_{k_j}; \lambda_j), & |b_{k_1}| \geq \cdots \geq |b_{k_s}| > 0, \quad j \leq s_{\mathbf{b}}, \\ |g_{k_j}| \leq \lambda_j, & \{k_{s+1}, \dots, k_p\} = \mathcal{S}_{\mathbf{b}}^c, \quad j > s_{\mathbf{b}}, \end{cases}$$

where $\mathcal{S}_{\mathbf{b}} = \text{supp}(\mathbf{b})$ and $s_{\mathbf{b}} = |\mathcal{S}_{\mathbf{b}}|$. We denote such vectors $\mathbf{g}$ as $\dot{\text{Pen}}(\mathbf{b})$ or $\dot{\rho}_{\#}(\mathbf{b}; \boldsymbol{\lambda})$. As in the case separable penalties, $\text{Pen}(\mathbf{b})$ is taken as the favorable choice unless otherwise stated.

PROPOSITION 3. *Let $\text{Pen}(\mathbf{b}) = \rho_{\#}(\mathbf{b}; \boldsymbol{\lambda})$ be as in (3.2) with $\boldsymbol{\lambda} = \boldsymbol{\lambda}^{\#}$. Suppose $\rho(t; \lambda) = \int_0^{|t|} \dot{\rho}(x; \lambda) dx$. If $\dot{\rho}(x; \lambda)$ is continuous in $x > 0$, then*

$$(3.4) \quad \liminf_{t \rightarrow 0+} t^{-1} \{\text{Pen}(\mathbf{b} + t\mathbf{u}) - \text{Pen}(\mathbf{b})\} \geq \mathbf{g}^T \mathbf{u} \quad \forall \mathbf{u} \in \mathbb{R}^p, \mathbf{g} \in \partial \text{Pen}(\mathbf{b}).$$

If $\dot{\rho}(x; \lambda)$ is nondecreasing in λ almost everywhere in positive x , then

$$(3.5) \quad \mathbf{h}^T \{\dot{\rho}_{\#}(\mathbf{b}; \boldsymbol{\lambda}) - \dot{\rho}_{\#}(\mathbf{b} + \mathbf{h}; \boldsymbol{\lambda})\} \leq \bar{\kappa}(\rho) \|\mathbf{h}\|_2^2 \quad \forall \mathbf{b}, \mathbf{h}.$$

The monotonicity condition on $\dot{\rho}(x; \lambda)$ holds for the ℓ_1 , SCAD and MCP. In general, we may define the concavity of $\dot{\rho}_{\#}(\mathbf{b}; \boldsymbol{\lambda})$ at $\mathbf{b}$ as

$$(3.6) \quad \bar{\kappa}(\mathbf{b}) = \bar{\kappa}(\mathbf{b}; \boldsymbol{\lambda}) = \sup_{\mathbf{h} \neq 0} \{\mathbf{h}^T (\dot{\rho}_{\#}(\mathbf{b}; \boldsymbol{\lambda}) - \dot{\rho}_{\#}(\mathbf{b} + \mathbf{h}; \boldsymbol{\lambda})) / \|\mathbf{h}\|_2^2\}.$$

See Section 5. Thus, (3.5) asserts that sorting the penalty does not increase the maximum concavity of a penalty family. For the penalty level, (3.3) asserts that the maximum penalty level at each index $j \in \mathcal{S}_{\mathbf{b}}^c$ is λ_{s+1} . Although the penalty level does not reach λ_{s+1} simultaneously for all $j \in \mathcal{S}_{\mathbf{b}}^c$, we still take λ_{s+1} as the penalty level for the sorted penalty $\rho_{\#}(\mathbf{b}; \boldsymbol{\lambda})$. This is especially reasonable when λ_j decreases slowly in j . In Section 3.4, we show that this weaker version of the penalty level is adequate for Gaussian errors, provided a certain minimum penalty level condition in (3.13) below. More important, (3.3) shows sorted penalties automatically pick penalty level λ_{s+1} from the sequence $\{\lambda_j\}$ without requiring the knowledge of s .

A key element of the proof of Proposition 3 is to write (3.2) as

$$(3.7) \quad \rho_{\#}(\mathbf{b}; \boldsymbol{\lambda}) = \max \left\{ \sum_{j=1}^p \rho(b_j; \lambda_{k_j}) : (k_1, \dots, k_p)^T \in \text{perm}(p) \right\},$$

where $\text{perm}(p)$ is the set of all vectors generated by permuting $(1, \dots, p)^T$. For the Slope with $\rho(t; \lambda) = \lambda|b|$, (3.7) follows directly from the rearrangement inequality. The monotonicity of $\dot{\rho}(t; \lambda)$ in λ , imposed in Proposition 3, is actually also necessary for the equivalence of (3.7) and (3.2) for all $(\mathbf{b}, \boldsymbol{\lambda})$.

3.2. *PLSE with sorted and smaller penalties.* For general penalties $\text{Pen}(\cdot)$, we consider local approximate solutions satisfying

$$(3.8) \quad \mathbf{X}^T(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})/n = \dot{\text{Pen}}(\hat{\boldsymbol{\beta}}) + \mathbf{v},$$

including nonseparable sorted penalties and separable penalties (2.6). However, instead of imposing ℓ_∞ bounds on the approximation error as in (2.11), we impose in this section a more practical ℓ_∞ - ℓ_2 split bound as in (3.12) below. Computational algorithms for approximate solutions and statistical properties of the resulting estimators have been considered in [1, 12, 15, 17, 31] among others. However, these studies of approximate solutions all focus on separable penalties with penalty level (2.8) or higher.

Given a sorted penalty with $\lambda_1 \geq \lambda_2 \geq \cdots \geq \lambda_p \geq 0$, define a sorted norm

$$(3.9) \quad \|\mathbf{b}\|_{\#,s} = \sum_{j=1}^{p-s} (\lambda_{s+j}/\lambda_{s+1}) b_j^\# \quad \forall \mathbf{b} \in \mathbb{R}^{p-s}, s = |\mathcal{S}|,$$

a standardized dual norm induced by the set of $\mathbf{g}_{\mathcal{S}^c}$ given in (3.3). Here, $b_j^\#$ is the j th largest value among $|b_1|, \dots, |b_{p-s}|$. Note that $\|\mathbf{b}\|_{\#,s}$ reduces to $\|\mathbf{b}_{\mathcal{S}^c}\|_1$ when $\lambda_{s+1} = \cdots = \lambda_p$.

Define $\lambda = \lambda_{s+1}$ as the effective penalty level for sorted penalties. Let $\mathbf{z} = \mathbf{X}^T(\mathbf{y} - \mathbf{X}\boldsymbol{\beta}^o)/n$. For positive r and γ and $\{\mathbf{w}, \mathbf{v}\} \subset \mathbb{R}^p$, define

$$(3.10) \quad \Delta(r, \mathbf{w}, \mathbf{v}) = \sup_{\mathbf{u} \neq \mathbf{0}} \frac{\mathbf{u}_{\mathcal{S}^c}^T \mathbf{z}_{\mathcal{S}^c} / \lambda - \eta \|\mathbf{u}_{\mathcal{S}^c}\|_{\#,s} - \mathbf{u}_{\mathcal{S}}^T \mathbf{w}_{\mathcal{S}} - \mathbf{u}^T \mathbf{v} / \lambda}{r \max\{\|\mathbf{u}\|_2, \|\mathbf{X}\mathbf{u}\|_2 \sqrt{\gamma/n}\}}.$$

We assume that for a certain constant $\xi > 0$,

$$(3.11) \quad \Delta(r_1, \mathbf{w}, \mathbf{v}) < 1 \quad \text{for some } r_1 \leq (1 - \eta)\xi|\mathcal{S}|^{1/2}$$

for each approximate solutions (3.8) under consideration, where $\mathbf{v}$ is as in (3.8) and $\mathbf{w} = \{\dot{\text{Pen}}(\boldsymbol{\beta}^o) - \mathbf{z}\}/\lambda$ with a favorable $\dot{\text{Pen}}(\boldsymbol{\beta}^o)$ as in (3.3). While $(1 - \eta)\xi|\mathcal{S}|^{1/2}$ is imposed as an upper bound for r_1 for all approximate solutions under consideration, (3.11) also allows individual approximate solutions to have a r_1 of smaller order, for example, to have smaller $\mathbf{w}_{\mathcal{S}}$ and/or $\mathbf{v}$. The seemingly complicated (3.10), which summarize conditions in our analysis on the noise vector $\boldsymbol{\varepsilon}$, the penalty level λ and approximation error $\mathbf{v}$, is actually not hard to decipher. Consider $\mathbf{v}$ satisfying

$$(3.12) \quad \mathbf{u}^T \mathbf{v} \leq \eta_1 \lambda_{s+1} \|\mathbf{u}_{\mathcal{S}^c}\|_{\#,s} + r_2 \lambda_{s+1} \|\mathbf{u}\|_2 \quad \forall \mathbf{u} \in \mathbb{R}^p$$

with $\eta_1 \in (0, \eta)$ and $r_2 > 0$. Proposition 4 below shows that under the normality assumption (2.10) and the minimum penalty level condition

$$(3.13) \quad \lambda_j \geq \lambda_{*,j} = A_0 \sigma \sqrt{(2/n) \log(p/(\alpha j))}, \quad j = 1, \dots, p,$$

with $A_0 > 1 > \alpha$, (3.11) holds with high probability.

In the simpler case with unsorted smaller penalties, we impose the same condition (3.11). By Proposition 10 in [26] and some algebra, (3.11) also holds with high probability under the same normality assumption (2.10) and the minimum penalty level condition

$$(3.14) \quad \lambda \geq \lambda_* = (\eta - \eta_1)^{-1} \sigma L / n^{1/2}$$

with $\Phi(-L) \leq s/p$, for example, $L = \sqrt{2 \log(p/s)}$, when $\|\mathbf{w}_S\|_\infty \leq (1 - \eta)\xi'$ with $\xi' < \xi$, $s = |\mathcal{S}|$ and $r_2/s^{1/2}$ is sufficiently small.

Building upon (3.11) and the approximate KKT condition (3.8), we define a solution set as in (2.12) as follows:

$$(3.15) \quad \mathcal{B}(\lambda_*, \kappa_*, \gamma) = \{\hat{\boldsymbol{\beta}} : (3.8) \text{ and } (3.11) \text{ hold, } \bar{\kappa}(\hat{\boldsymbol{\beta}}) \leq \kappa_*\},$$

where $\bar{\kappa}(\mathbf{b})$ is the concavity of the sorted penalty as defined in (3.6). Here, λ_* is a minimum penalty level requirement implicit in (3.11), for example, (3.13) for sorted penalties and (3.14) for unsorted smaller penalties. As in Section 2, we define below a subclass of (3.15) to which our analysis applies.

We first relax (2.14), the condition that requires the solution to be connected to $\mathbf{0}$. Let

$$(3.16) \quad \|\mathbf{u}\|_{\#,*} = \|\mathbf{u}_{S^c}\|_{\#,*} + |\mathcal{S}|^{1/2} \max\{\|\mathbf{u}\|_2, \|\mathbf{X}\mathbf{u}\|_2 \sqrt{\gamma/n}\}.$$

We say that two solutions $\hat{\boldsymbol{\beta}}$ and $\tilde{\boldsymbol{\beta}}$ in $\mathcal{B}(\lambda_*, \kappa_*, \gamma)$ are connected by an a_0 -chain if there exist $\hat{\boldsymbol{\beta}}^{(k)} \in \mathcal{B}(\lambda_*, \kappa_*, \gamma)$ with sorted penalty levels $\{\lambda_1^{(k)}, \dots, \lambda_p^{(k)}\}$ such that for the $a_0 > 0$ in Proposition 5 below:

$$(3.17) \quad \hat{\boldsymbol{\beta}}^{(0)} = \tilde{\boldsymbol{\beta}}, \quad \hat{\boldsymbol{\beta}}^{(k^*)} = \hat{\boldsymbol{\beta}}, \quad \|\hat{\boldsymbol{\beta}}^{(k)} - \hat{\boldsymbol{\beta}}^{(k-1)}\|_{\#,*} \leq a_0 |\mathcal{S}| \lambda_{s+1}^{(k)},$$

$k = 1, \dots, k^*$. This condition holds if $\hat{\boldsymbol{\beta}}$ and $\tilde{\boldsymbol{\beta}}$ are connected through a continuous path in $\mathcal{B}(\lambda_*, \kappa_*, \gamma)$. Similar to (2.13), we define the sparse branch of the solution set $\mathcal{B}(\lambda_*, \kappa_*, \gamma)$ as

$$(3.18) \quad \mathcal{B}_0(\lambda_*, \kappa_*, \gamma) = \{\hat{\boldsymbol{\beta}} \in \mathcal{B}(\lambda_*, \kappa_*, \gamma) : (3.17) \text{ holds with } \tilde{\boldsymbol{\beta}} = \mathbf{0}\}.$$

3.3. RE condition for sorted and smaller penalties. As in Section 2, we shall prove that all solutions in the sparse branch (3.18) belong to the following cone:

$$(3.19) \quad \mathcal{C}_\#(\mathcal{S}; \xi, \gamma) = \{\mathbf{u} : \|\mathbf{u}_{S^c}\|_{\#,*} \leq \xi |\mathcal{S}|^{1/2} \max(\|\mathbf{u}\|_2, (\gamma \mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u})^{1/2})\},$$

when the RE below is no smaller than the κ_* in (3.15).

The RE for sorted penalties is defined as

$$(3.20) \quad \text{RE}_\#(\mathcal{S}; \xi, \gamma) = \inf \left\{ \frac{(\mathbf{u}^T \bar{\boldsymbol{\Sigma}} \mathbf{u})^{1/2}}{\|\mathbf{u}\|_2} : \mathbf{0} \neq \mathbf{u} \in \mathcal{C}_\#(\mathcal{S}; \xi, \gamma) \right\}.$$

As $\mathcal{C}_\#(S; \xi, \gamma)$ in (3.19) depends on the sorted $\boldsymbol{\lambda} = (\lambda_1, \dots, \lambda_p)^T$, the infimum in (3.20) is taken over all $\boldsymbol{\lambda}$ under consideration. We note that

$$\text{RE}_\#(S; \xi, \gamma) \geq \left[\inf \left\{ \frac{(\mathbf{u}^T \bar{\Sigma} \mathbf{u})^{1/2}}{\|\mathbf{u}\|_2} : \|\mathbf{u}_{\mathcal{S}^c}\|_{\#,s} < \xi |\mathcal{S}|^{1/2} \|\mathbf{u}\|_2 \right\} \right] \wedge \frac{1}{\gamma}.$$

When the cone is confined to $\lambda_1 = \lambda_p$, that is, $\|\mathbf{u}_{\mathcal{S}^c}\|_{\#,s} = \|\mathbf{u}_{\mathcal{S}^c}\|_1$, the RE condition on $\text{RE}_\#^2(S; \xi, \gamma)$ is equivalent to the restricted strong convexity condition [17] as $\|\mathbf{u}_{\mathcal{S}^c}\|_1 < \xi |\mathcal{S}|^{1/2} \|\mathbf{u}\|_2$ implies $\|\mathbf{u}\|_1 < (\xi + 1) |\mathcal{S}|^{1/2} \|\mathbf{u}\|_2$. Compared with (2.16), the RE in (3.20) is smaller due to the use of a larger cone. However, this is hard to avoid because the sorted penalty may not control the ℓ_∞ measure of the noise as in (2.8) and we do not wish to impose uniform bound on the approximation error $\mathbf{v}$.

3.4. Error bounds. Let $\mathcal{B}(\lambda_*, \kappa_*, \gamma)$ be as in (3.15), $\mathcal{B}_0(\lambda_*, \kappa_*, \gamma)$ as in (3.18), $\mathcal{C}_\#(S; \xi, \gamma)$ be as in (3.20). We define the good solution set as

$$(3.21) \quad \begin{aligned} & \mathcal{B}_0^*(\lambda_*, \kappa_*, \gamma) \\ &= \mathcal{B}_0(\lambda_*, \kappa_*, \gamma) \cup \{\hat{\boldsymbol{\beta}} \in \mathcal{B}(\lambda_*, \kappa_*, \gamma) : \hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o \in \mathcal{C}_\#(S; \xi, \gamma)\}. \end{aligned}$$

This is the set of approximate solutions (3.8) with the estimation error inside the cone or connected to the origin through a chain. Now we provide predication and estimation error bounds for the good solution under the RE condition.

THEOREM 6. Suppose $\text{RE}_\#^2(S; \xi, \gamma) \geq \kappa_* > 0$. Let $\hat{\boldsymbol{\beta}} \in \mathcal{B}_0^*(\lambda_*, \kappa_*, \gamma)$ satisfying $\bar{\kappa}(\hat{\boldsymbol{\beta}}) \leq (1 - 1/C_0) \text{RE}_\#^2(S; \xi, \gamma)$ and condition (3.11) with a certain $r_1 \leq (1 - \eta) \xi |\mathcal{S}|^{1/2}$ and λ for either a sorted ($\lambda = \lambda_{s+1}$) or separable penalty. Then, for any seminorm $\|\cdot\|$,

$$(3.22) \quad \|\mathbf{h}\| \leq C_0 \lambda \sup_{\mathbf{u} \neq 0} \frac{\|\mathbf{u}\| \{r_1 F(\mathbf{u}) - (1 - \eta) \|\mathbf{u}_{\mathcal{S}^c}\|_{\#,s}\}}{\mathbf{u}^T \bar{\Sigma} \mathbf{u}},$$

where $\mathbf{h} = \hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o$ and $F(\mathbf{u}) = \|\mathbf{u}\|_2 \vee (\gamma \mathbf{u}^T \bar{\Sigma} \mathbf{u})^{1/2}$. In particular,

$$\frac{\|\mathbf{h}_{\mathcal{S}}\|_1 + \|\mathbf{h}_{\mathcal{S}^c}\|_{\#,s}}{(1 + \xi) |\mathcal{S}|^{1/2}} \leq F(\mathbf{h}) \leq \frac{(\mathbf{h}^T \bar{\Sigma} \mathbf{h})^{1/2}}{\text{RE}_\#(S; \xi, \gamma)} \leq \frac{C_0 r_1 \lambda}{\text{RE}_\#^2(S; \xi, \gamma)}.$$

As an extension of Theorem 5, Theorem 6 provides prediction and ℓ_2 estimation error bounds in the same form along with a comparable sorted ℓ_1 error bound. It demonstrates the benefit of sorted concave penalization as $r_1^2 \asymp s_1 + s/\log^2(p/s) + r_2^2$ in standard settings as described in Theorem 7 and Corollary 3 below, where $s = |\mathcal{S}|$, s_1 can be viewed as the number of small nonzero coefficients, and r_2 is the ℓ_2 component of the approximation error as in (3.12).

Theorem 6 applies to approximate solutions of PLSE with sorted penalties that adaptively choose penalty level λ_{s+1} from assigned sorted sequence $\lambda_1 \geq \dots \geq \lambda_p$,

or unsorted smaller penalties ($\lambda_j = \lambda \ \forall j$). However, it does not guarantee the selection consistency of (3.8) as a false positive cannot be ruled out at the smaller penalty level or with general approximation errors. On the other hand, as properties of the Lasso at smaller penalty levels and the Slope have been studied in [3, 24, 26] among others, Theorem 6 can be viewed as an extension of their results to concave and/or sorted penalties.

It is also possible to derive error bounds in the case of $C_0 = \infty$ as in Theorem 3 if the ℓ_∞ norm in the definition of the RCIF is replaced by the norm $\|\cdot\|_{\#,*}$ in (3.16). We omit details for the sake of space.

We still need to prove that (3.11) holds with high probability under the normality assumption (2.10) and the minimum penalty level condition (3.13) and (3.14), respectively for sorted and fixed penalty levels. To this end, we need to define a few more quantities. For nonnegative integer s and $0 < \alpha < 1 < A$, define $p_{\alpha,A} = 2\alpha \sum_{k=0}^{\infty} \alpha^{(A-1)A^k}$, $q_{\alpha,A} = (1 - \sqrt{2p_{\alpha,A}})_+$, $x_1 = s/q_{\alpha,A}$, $L_x = \sqrt{2\log(p/(\alpha x))}$, and

$$(3.23) \quad \mu_{\#,s} = \left\{ \frac{2s(x_1/p)^{A^2-1}(2/q_{\alpha,A})^2}{A^2 L_{s+1}^2 (A^2 L_{x_1}^2 + 2)} \right\}^{1/2} I_{\{s=0\}}.$$

We assume here that $x_1 \leq p$. This is reasonable when p/s is large.

THEOREM 7. *Suppose the normality condition (2.10) on the noise and condition (3.12) on the approximation error with $\eta_1 \in (0, \eta)$ and $r_2 > 0$. Let $A > 1$. Suppose (3.13) holds for sorted penalties with $\alpha \in (0, 1/4)$ and $A_0 = A/(\eta - \eta_1)$, and that (3.14) holds for unsorted smaller penalties. Let $\beta^o = \hat{\beta}^o$ be the oracle LSE as in (2.9) and $\mathcal{B}_0^*(\lambda_*, \kappa_*, \gamma)$ be the solution set in (3.21). Let $F(\mathbf{u}) = \max\{\|\mathbf{u}\|_2, (\gamma \mathbf{u}^T \bar{\Sigma} \mathbf{u})^{1/2}\}$, $\eta_* = \{\sum_{j=1}^s \lambda_j^2 / (\lambda_{s+1}^2 (s \vee 1))\}^{1/2}$,*

$$\xi = \{(1 - 1/a)(1 - \eta)\}^{-1} [\{(\eta - \eta_1)^2 \mu_{\#,s}^2 + \eta_*^2\}^{1/2} + r_2/s^{1/2}]$$

with the $\mu_{\#,s}$ in (3.23) and $a \in (0, 1)$. Suppose $\kappa_ \leq (1 - 1/C^*) \text{RE}_{\#}^2(\mathcal{S}; \xi, \gamma)$. Then there exists an event Ω such that $\mathbb{P}\{\Omega\} \geq 1 - e^{-\xi_* s \log(p/(\alpha(s+1)))}$ with $\xi_* = (1 - \eta)^2 \xi^2 \gamma / \{a(\eta - \eta_1)\}^2$, and that in the event Ω*

$$(3.24) \quad \|\hat{\beta} - \hat{\beta}^o\| \leq (1 - \eta) C^* \lambda \sup_{\mathbf{u} \neq 0} \frac{\|\mathbf{u}\| \{ \xi s^{1/2} F(\mathbf{u}) - \|\mathbf{u}\|_{S^c} \|_{\#,s} \}}{\mathbf{u}^T \bar{\Sigma} \mathbf{u}}$$

for all approximate solutions $\hat{\beta} \in \mathcal{B}_0^(\lambda_*, \kappa_*, \gamma)$ and seminorms $\|\cdot\|$, where $\lambda = \lambda_{s+1}$ for sorted penalties. In addition, for all sorted or unsorted penalties satisfying $\|\text{Pen}(\beta^o)/\lambda\|_2 \leq s_1$,*

$$(3.25) \quad \|\hat{\beta} - \hat{\beta}^o\| \leq C^* \lambda \sup_{\mathbf{u} \neq 0} \frac{\|\mathbf{u}\| \{ r_1 F(\mathbf{u}) - (1 - \eta) \|\mathbf{u}\|_{S^c} \|_{\#,s} \}}{\mathbf{u}^T \bar{\Sigma} \mathbf{u}}$$

with $r_1 = (1 - 1/a)^{-1} [\{(\eta - \eta_1)^2 \mu_{\#,s}^2 + s_1^2\}^{1/2} + r_2]$ and at least probability $\mathbb{P}\{\Omega\} - e^{-\xi'_ r_1 \log(p/(\alpha(s+1)))}$ with $\xi'_* = \gamma / \{a(\eta - \eta_1)\}^2$.*

COROLLARY 3. Suppose $\lambda = \lambda_{s+1} \asymp \sigma \sqrt{(2/n) \log(p/s)}$ in (3.25) and $\bar{\kappa}(\hat{\beta}) \leq (1 - 1/C_0) \text{RE}_{\#}^2(\mathcal{S}; \xi, \gamma)$ with $C_0^2 / \text{RE}_{\#}^2(\mathcal{S}; \xi, \gamma) = O(1)$. Then

$$\begin{aligned} & \|X\hat{\beta} - X\hat{\beta}^o\|_2^2/n + \|\hat{\beta} - \hat{\beta}^o\|_2^2 + \|\hat{\beta} - \hat{\beta}^o\|_{\#,s}^2/s = O_{\mathbb{P}}(\sigma^2/n) r_1^2 \log(p/s) \\ & \text{with } r_1^2 = s_1 + s/\log^2(p/s) + r_2^2, \text{ where } s_1 = \|\dot{\text{Pen}}(\hat{\beta}^o)/\lambda\|_2^2, \text{ and} \\ & \|X\hat{\beta} - X\beta^*\|_2^2/n + \|\hat{\beta} - \beta^*\|_2^2 + \|\hat{\beta} - \beta^*\|_{\#,s}^2/s \\ & = O_{\mathbb{P}}(\sigma^2/n) \{(s_1 + r_2^2) \log(p/s) + s\}. \end{aligned} \quad (3.26)$$

Moreover, for sorted concave penalty (3.2) with $\text{supp}(\dot{\rho}(\cdot; \lambda)) \subseteq [-\gamma\lambda, \gamma\lambda]$,

$$s_1 \leq \#\{j \leq |\mathcal{S}| : (\beta^o)_j^{\#} \leq \gamma\lambda_j\}.$$

Corollary 3 extends Corollary 1 to smaller penalty levels $\lambda = \lambda_{s+1} \geq A_0 \sigma \sqrt{(2/n) \log(p/\alpha(s+1))}$ and sorted penalties. In the worst case scenario where $s_1 \asymp s$, the error bounds in Corollary 3 attain the minimax rate [3, 32].

Theorem 7 and Corollary 3 provide sufficient conditions to guarantee simultaneous adaptation of sorted concave PLSE: (a) picking level λ_{s+1} automatically from $\{\lambda_1, \dots, \lambda_p\}$, and (b) partially removing the bias of the Slope [3, 24] when $s_1 \ll s$, without requiring the knowledge of s or s_1 .

Theorem 7 is a direct consequence of Theorem 6 and the following proposition.

PROPOSITION 4. Suppose the normality assumption (2.10) holds. Let $\eta > \eta_1$ and $\{\lambda_j\}$ be as in (3.13) for all sorted penalties. Let $\{\sum_{j=1}^s \lambda_j^2 / (\lambda_{s+1}^2 (s \vee 1))\}^{1/2}$, $\mathbf{z} = \mathbf{X}^T(\mathbf{y} - \mathbf{X}\hat{\beta}^o)/n$, $\mathbf{w} = \{\dot{\text{Pen}}(\hat{\beta}^o) - \mathbf{z}\}/\lambda_{s+1}$, $\mu_{\#,s}$ be as in (3.23) with $q_{\alpha,A} > 0$ and $L = \sqrt{2 \log(p/(\alpha(s+1)))}$. Suppose $\text{supp}(\mathbf{z}) \subseteq \mathcal{S}^c$. Then

$$(3.27) \quad \mathbb{P}\{(3.12) \text{ and } \|\mathbf{w}_{\mathcal{S}}\|_2 \leq w \text{ imply } \Delta(r_1, \mathbf{w}, \mathbf{v}) < 1\} \geq \Phi\left(\frac{r_1 L \sqrt{\gamma}}{a(\eta - \eta_1)}\right)$$

when $(1 - 1/a)r_1 \geq \{(\eta - \eta_1)^2 \mu_{\#,s}^2 + w^2\}^{1/2} + r_2$ for $a > 0$. In particular, when $w = \eta_* s^{1/2}$, $\|\mathbf{w}_{\mathcal{S}}\|_2 \leq w$ holds automatically and

$$(3.28) \quad \mathbb{P}\{(3.12) \text{ implies (3.11)}\} \geq \Phi\left(\frac{(1 - \eta)\xi s^{1/2} L}{a(\eta - \eta_1) \gamma^{-1/2}}\right),$$

with $\xi \geq \{(1 - 1/a)(1 - \eta)\}^{-1} [(\eta - \eta_1)^2 \mu_{\#,s}^2/s + \eta_*^2]^{1/2} + r_2/s^{1/2}$. Moreover, when (3.8) is confined to penalties with the minimum fixed penalty level $\lambda \geq \lambda_*$ as in (3.14), (3.28) holds with the L in (3.14) and $\mu_{\#,s} = \sqrt{4s/(L^4 + 2L^2)}$.

The following lemma provides the basic inequality for analyzing PLSE with sorted and smaller penalties.

LEMMA 2. Let $\widehat{\boldsymbol{\beta}}$ be a solution of (3.8), $\mathbf{h} = \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o$, $\mathbf{z} = \mathbf{X}^T(\mathbf{y} - \mathbf{X}\boldsymbol{\beta}^o)/n$, $\mathbf{w} = \{\dot{\text{Pen}}(\boldsymbol{\beta}^o) - \mathbf{z}\}/\lambda$ with λ being the fixed penalty level for unsorted penalties or $\lambda = \lambda_{s+1}$ for sorted penalties. Then

$$(3.29) \quad \begin{aligned} & \mathbf{h}^T \overline{\boldsymbol{\Sigma}} \mathbf{h} + (1 - \eta)\lambda \|\mathbf{h}_{\mathcal{S}^c}\|_{\#,s} - \bar{\kappa}(\widehat{\boldsymbol{\beta}}) \|\mathbf{h}\|_2^2 \\ & \leq (\mathbf{h}_{\mathcal{S}^c}^T \mathbf{z}_{\mathcal{S}^c} - \eta\lambda \|\mathbf{h}_{\mathcal{S}^c}\|_{\#,s}) - \lambda \mathbf{h}_{\mathcal{S}}^T \mathbf{w}_{\mathcal{S}} - \mathbf{h}^T \mathbf{v}. \end{aligned}$$

Next, we provide conditions under which the error $\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o$ belongs to the cone $\mathcal{C}_{\#}(\mathcal{S}; \xi, \gamma)$ for the solutions in $\mathcal{B}_0^*(\lambda_*, \kappa_*, \gamma)$ in (3.18).

PROPOSITION 5. Let $\bar{\kappa}(\mathbf{b})$ be as in (3.6), $\mathcal{B}(\lambda_*, \kappa_*, \gamma)$ as in (3.15) and $\mathcal{C}_{\#}(\mathcal{S}; \xi, \gamma)$ as in (3.19). Suppose $\text{RE}_{\#}^2(\mathcal{S}; \xi, \gamma) \geq \kappa_*$ and

$$(3.30) \quad \Delta((1 - \eta - a_1)\xi|\mathcal{S}|^{1/2}, \mathbf{w}, \mathbf{v}) \leq 1, \quad 0 < a_1 < 1 - \eta,$$

as in (3.11) for all $\{\lambda, \mathbf{w}, \mathbf{v}\}$ associated with solutions in $\mathcal{B}(\lambda_*, \kappa_*, \gamma)$. Let $a_2 = a_1\xi^2/\{2\bar{\kappa}(\widehat{\boldsymbol{\beta}})\}$, $a_3 = a_1(1 - \eta)\xi/\{(1 - \eta - a_1)(\xi + 1) + a_1\}$ and $a_0 = \min[a_2, a_2a_3/\{(1 - \eta)(\xi \vee 1)\}]$. Then

$$\mathcal{B}_0^*(\lambda_*, \kappa_*, \gamma) = \{\widehat{\boldsymbol{\beta}} \in \mathcal{B}(\lambda_*, \kappa_*, \gamma) : \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^o \in \mathcal{C}_{\#}(\mathcal{S}; \xi, \gamma)\}.$$

4. Local convex approximation. In this section, we develop a local convex approximation (LCA) algorithm for general penalized optimization problem (1.2), especially for sorted penalties. The LCA is a majorization-minimization (MM) algorithm, and is closely related to and in fact very much inspired by the LQA [11] and LLA [14, 36, 38]. We describe the LCA algorithm and how it can be applied to sorted penalizations in the first subsection, and prove the error bounds for estimation and prediction in the second subsection.

4.1. *The LCA algorithm.* Consider a general minimization problem

$$(4.1) \quad L(\mathbf{b}) + \text{Pen}(\mathbf{b}), \quad \mathbf{b} \in \mathbb{R}^p,$$

with $L(\mathbf{b})$ a differentiable loss function and $\text{Pen}(\mathbf{b})$ a multivariate penalty function. Suppose for a certain continuously differentiable convex $\text{Pen}_{-}(\mathbf{b})$,

$$(4.2) \quad \text{Pen}_{+}(\mathbf{b}) = \text{Pen}(\mathbf{b}) + \text{Pen}_{-}(\mathbf{b})$$

is convex. The LCA algorithm can be written as

$$(4.3) \quad \mathbf{b}^{(\text{new})} = \arg \min_{\mathbf{b}} \{L(\mathbf{b}) + \text{Pen}_{+}(\mathbf{b}) - \mathbf{b}^T \dot{\text{Pen}}_{-}(\mathbf{b}^{(\text{old})})\}.$$

This is clearly a MM-algorithm: Because

$$\text{Pen}^{(\text{new})}(\mathbf{b}) = \text{Pen}_{+}(\mathbf{b}) - \text{Pen}_{-}(\mathbf{b}^{(\text{old})}) - (\mathbf{b} - \mathbf{b}^{(\text{old})})^T \dot{\text{Pen}}_{-}(\mathbf{b}^{(\text{old})})$$

is a convex majorization of $\text{Pen}(\mathbf{b})$ with $\text{Pen}^{(\text{new})}(\mathbf{b}^{(\text{old})}) = \text{Pen}(\mathbf{b}^{(\text{old})})$,

$$\begin{aligned} L(\mathbf{b}^{(\text{new})}) + \text{Pen}(\mathbf{b}^{(\text{new})}) &\leq L(\mathbf{b}^{(\text{new})}) + \text{Pen}^{(\text{new})}(\mathbf{b}^{(\text{new})}) \\ (4.4) \qquad \qquad \qquad &\leq L(\mathbf{b}^{(\text{old})}) + \text{Pen}^{(\text{new})}(\mathbf{b}^{(\text{old})}) \\ &= L(\mathbf{b}^{(\text{old})}) + \text{Pen}(\mathbf{b}^{(\text{old})}). \end{aligned}$$

Let $\rho_{\#}(\mathbf{b}; \lambda)$ be the sorted concave penalty in (3.2) with a penalty family $\rho(t; \lambda)$ and a vector of sorted penalty levels $\lambda = (\lambda_1, \dots, \lambda_p)^T$. Suppose $\dot{\rho}(x; \lambda) = (\partial/\partial x)\rho(x; \lambda)$ is nondecreasing in λ almost everywhere in positive x , so that Proposition 3 applies. Suppose for a certain continuously differentiable convex function $\rho_{-}(t)$,

$$(4.5) \qquad \rho_{+}(t; \lambda_j) = \rho(t; \lambda_j) + \rho_{-}(t) \quad \text{is convex in } t \text{ for } j = 1, \dots, p.$$

By (3.7), $\rho_{+, \#}(\mathbf{b}; \lambda) = \rho_{\#}(\mathbf{b}; \lambda) + \rho_{-}(\mathbf{b})$, the sorted penalty with $\rho_{+}(t; \lambda)$, is convex in $\mathbf{b}$, so that the LCA algorithm for $\rho_{\#}(\mathbf{b}; \lambda)$ can be written as

$$(4.6) \qquad \mathbf{b}^{(\text{new})} = \arg \min_{\mathbf{b}} \{L(\mathbf{b}) + \rho_{+, \#}(\mathbf{b}; \lambda) - \mathbf{b}^T \dot{\rho}_{-}(\mathbf{b}^{(\text{old})})\},$$

where $\dot{\rho}_{-}(\mathbf{b})$ is the gradient of $\rho_{-}(\mathbf{b}) = \sum_{j=1}^p \rho_{-}(b_j)$. The simplest version of LCA takes $\rho_{-}(t) = t^2 \bar{\kappa}(\rho)/2$ with the maximum concavity defined in (2.4), but this is not necessary as (4.5) is only required to hold for the given λ .

Figure 1 demonstrates that for $p = 1$, the LCA with $\rho_{-}(t) = t^2 \bar{\kappa}(\rho)/2$ also majorizes the LLA with $\text{Pen}^{(\text{new})}(b) = \rho(|b^{(\text{old})}|; \lambda) + \dot{\rho}(|b^{(\text{old})}|; \lambda)(|b| - |b^{(\text{old})}|)$. With $\rho_{-}(t) = \lambda|t| - \rho(t; \lambda)$ in (4.5), the LCA is identical to an unfolded LLA with $\text{Pen}^{(\text{new})}(b) = \lambda|b| + \{\dot{\rho}(b^{(\text{old})}; \lambda) - \lambda \text{sgn}(b^{(\text{old})})\}(b - b^{(\text{old})})$. The situation is the

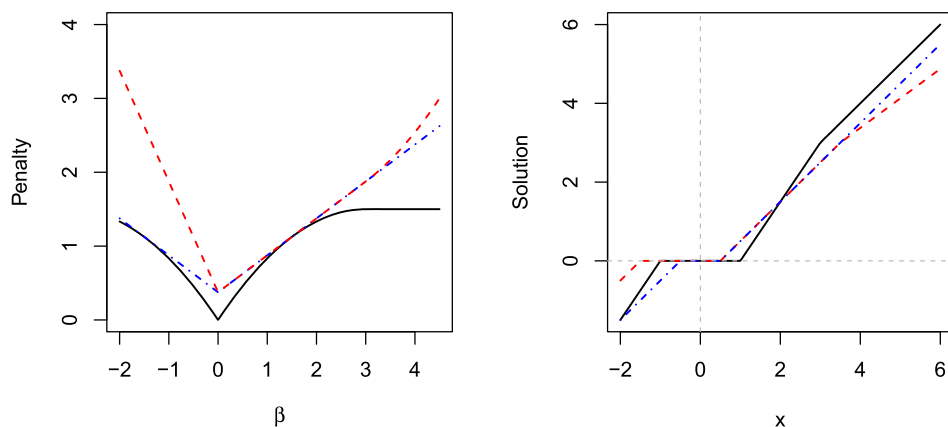


FIG. 1. Local convex approximation (red dashed), local linear approximation (blue mixed) and original penalty (black solid) for MCP with $\lambda = 1$ and $\bar{\kappa} = 1/3$ at $b^{(\text{old})} = 1.5$. Left: penalty function and its approximations; Right: $\arg \min_b \{(x - b)^2/2 + \text{Pen}(b)\}$.

Algorithm 1 FISTA for LCA

Initialization: $\mathbf{x}^1 = \mathbf{b}^0 = \mathbf{b}^{(\text{old})}$, $t_1 = 1$
 Iteration: $\mathbf{b}^k = \text{prox}(\mathbf{x}^k - t_* \nabla L(\mathbf{x}^k) + t_* \dot{\rho}_-(\mathbf{b}^{(\text{old})}); t_* \rho_{+,\#}(\cdot; \boldsymbol{\lambda}))$
 $t_{k+1} = \{1 + (1 + 4t_k^2)^{1/2}\}/2$
 $\mathbf{x}^{k+1} = \mathbf{b}^k + \{(t_k - 1)/t_{k+1}\}(\mathbf{b}^k - \mathbf{b}^{k-1})$

same for separable penalties, that is, $\lambda_1 = \lambda_p$. However, the LLA is not feasible for truly sorted concave penalties with $\lambda_1 > \lambda_p$. As the LCA also majorized the LLA, it imposes larger penalty on solutions with larger step size compared with the LLA, but this has little effect in our theoretical analysis in Sections 4.2.

As in (3.8), we may write approximate solutions of (4.3) and (4.6) as

$$(4.7) \quad \mathbf{0} = \dot{L}(\mathbf{b}^{(\text{new})}) + \dot{\text{Pen}}_+(\mathbf{b}^{(\text{new})}) - \dot{\text{Pen}}_-(\mathbf{b}^{(\text{old})}) + \mathbf{v}.$$

Such approximate solutions can be viewed as output of iterative algorithms such as the proximal gradient algorithms [2, 18, 21], which approximate $L(\mathbf{b})$ by $L(\mathbf{x}) + (\mathbf{b} - \mathbf{x})^T \nabla L(\mathbf{x}) + \|\mathbf{b} - \mathbf{x}\|_2^2/(2t)$ around $\mathbf{x}$. For example, for (4.6) with sorted concave penalty, the FISTA [2] can be written as Algorithm 1 where $\rho_{+,\#}(\mathbf{b}; \boldsymbol{\lambda}) = \sum_{j=1}^p \rho_+(b_j^\#; \lambda_j)$, t_* is the reciprocal of a Lipschitz constant for ∇L or determined in the iteration by backtracking, and

$$(4.8) \quad \text{prox}(\mathbf{x}; \text{Pen}) = \arg \min_{\mathbf{b}} \{\|\mathbf{b} - \mathbf{x}\|_2^2/2 + \text{Pen}(\mathbf{b})\}$$

is the proximal mapping for convex Pen, for example, $\text{Pen}(\mathbf{b}) = t\rho_{+,\#}(\mathbf{b}; \boldsymbol{\lambda})$.

For sorted penalties $\rho_{\#}(\mathbf{b}; \boldsymbol{\lambda})$, the proximal mapping is not separable but still preserves the sign and ordering in absolute value of the input. Thus, after removing the sign and sorting the input and output simultaneously, it can be solved with the isotonic proximal mapping,

$$(4.9) \quad \text{iso.prox}(\mathbf{x}; \text{Pen}) = \arg \min_{\mathbf{b}} \{\|\mathbf{b} - \mathbf{x}\|_2^2/2 + \text{Pen}(\mathbf{b}) : b_j \downarrow \text{ in } j\},$$

with $\text{Pen}(\mathbf{b}) = t\rho_{+}(\mathbf{b}; \boldsymbol{\lambda})$. Moreover, similar to the Slope [6], (4.9) can be computed by Algorithm 2 as a pool adjacent violators algorithm.

For the MCP, $\rho_{+}(t; \lambda) = \rho(t; \lambda) + \bar{\kappa}t^2/2 = \int_0^{|t|} \{\lambda \vee (\bar{\kappa}x)\} dx$, the univariate

$$\text{prox}(x; t\rho_{+}(\cdot; \lambda)) = \text{sgn}(x) \min\{(|x| - t\lambda)_{+}, |x|/(1 + t\bar{\kappa})\},$$

is a combination of the soft threshold and shrinkage estimators. Figure 2 plots this univariate proximal mapping for a specific (λ, t) . The computational cost of Algorithm 2 for the MCP is no greater than $O(p \log p)$, the same as sorting its input, because the while loop is to run no more than p times and for each merge the cost of the search for the cutoff is $O(\log p)$. We formally state the above discussion in the following proposition.

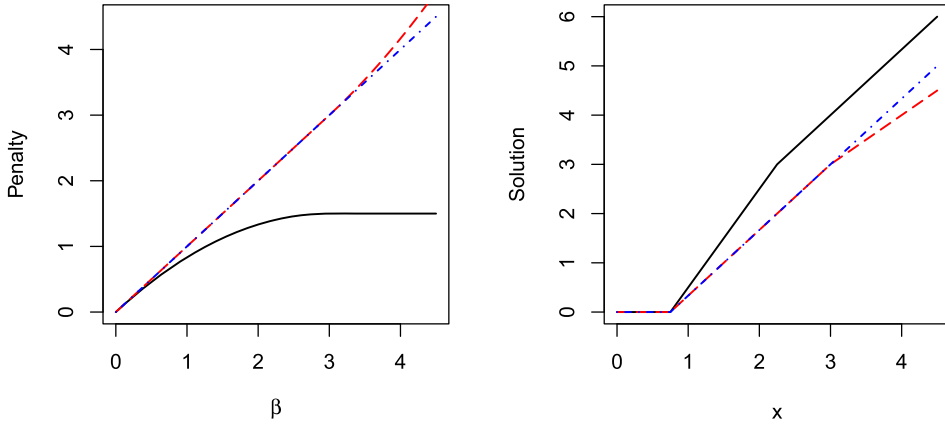


FIG. 2. $\rho_+(\cdot; \lambda) = \rho(t; \lambda) + \bar{\kappa}t^2/2$ for LCA (red dashed), Lasso (blue mixed) and MCP (black solid) with $\lambda = 1$ and $\bar{\kappa} = 1/3$. Left: penalties; Right: proximal mappings.

PROPOSITION 6. Let λ be a vector with components $\lambda_1 \geq \dots \geq \lambda_p \geq 0$.

(i) Let $\rho_{\#}(\mathbf{b}; \lambda) = \sum_{j=1}^p \rho(b_j^{\#}; \lambda_j)$ as in (3.2), and $\mathbf{b} = \text{prox}(\mathbf{x}; \rho_{\#}(\cdot; \lambda))$. Then $\text{sgn}(b_j) = \text{sgn}(x_j)$, $|x_j| \geq |x_k|$ implies $|b_j| \geq |b_k|$, and

$$(4.10) \quad (b_1^{\#}, \dots, b_p^{\#})^T = \text{iso.prox}(\mathbf{x}^{\#}; \rho(\cdot; \lambda)).$$

(ii) $\text{iso.prox}(\mathbf{x}; t\rho_+(\cdot; \lambda))$ is solved by Algorithm 2 when $x_1 \geq \dots \geq x_p \geq 0$.

(iii) For the MCP, the computational cost of Algorithm 2 is $O(p \log p)$ in the number of operations, and the cost of each iteration of Algorithm 1 is a sum of $O(p \log p)$ and the cost of computing the gradient $\nabla L(\mathbf{x}^k)$.

In linear regression, the cost of computing $\nabla L(\mathbf{x}^k)$ exactly is of the order np , and the cost of uniformly approximating $\nabla L(\mathbf{x}^k)$ in stochastic gradient descent is of no smaller order than $p \log p$. As $\log p \lesssim n$ is required for the RE condition to

Algorithm 2 $\text{iso.prox}(\mathbf{x}; t\rho_+(\cdot; \lambda))$ with monotone input

Input: $\lambda \downarrow, \mathbf{x} \downarrow$

Compute $b_j = \arg \min_b \{(x_j - b)^2/2 + t\rho_+(b; \lambda_j)\}$

While $\mathbf{b}$ is not nonincreasing **do**

Identify blocks of violators of the monotonicity constraint,

$$b_{j'-1} > b_{j'} \leq b_{j'+1} \leq \dots \leq b_{j''} > b_{j''+1}, b_{j'} < b_{j''}$$

Replace b_j , $j' \leq j \leq j''$, with common value b as follows:

For the MCP, $\rho_+(b; \lambda) = \int_0^{|b|} \{\lambda \vee (\bar{\kappa}x)\} dx$, b solves

$$\sum_{j=j'}^{j''} \{b - x_j + t \max(\lambda_j, \bar{\kappa}b)\} = 0$$

Else $b = \arg \min_{t \geq 0} \sum_{j=j'}^{j''} \{(x_j - t)^2/2 + t\rho_+(t; \lambda_j)\}$

hold uniformly with $|\mathcal{S}| = 2$, the cost of each iteration in Algorithm 1 for sorted MCP is dominated by that of computing or approximating the gradient, required for all penalties.

The computation of the isotonic proximal mapping for SCAD can be carried out in a similar but more complicated fashion, as the region for shrinkage is broken by an interval for soft thresholding.

4.2. Error bounds of the LCA. Let $\rho_{\#}(\mathbf{b}; \lambda)$ be the sorted concave penalty in (3.7), $\rho_{+,\#}(\mathbf{b}; \lambda)$ also given by (3.7) with $\rho(t; \lambda)$ replaced by the $\rho_{+}(t; \lambda)$ in (4.5), and $\rho_{-}(\mathbf{b}) = \sum_{j=1}^p \rho_{-}(b_j)$. For sorted concave PLSE, we write the approximate solution (4.7) of an LCA step as

$$(4.11) \quad \mathbf{X}^T(\mathbf{y} - \mathbf{X}\mathbf{b}^{(\text{new})})/n = \dot{\rho}_{+,\#}(\mathbf{b}^{(\text{new})}; \lambda^{(\text{new})}) - \dot{\rho}_{-}(\mathbf{b}^{(\text{old})}) + \mathbf{v}^{(\text{new})},$$

as $L(\mathbf{b}) = \|\mathbf{y} - \mathbf{X}\mathbf{b}\|_2^2/(2n)$, where $\dot{\rho}_{+,\#}(\mathbf{b}; \lambda)$ and $\dot{\rho}_{-}(\mathbf{b})$ can be any members of the respective subdifferentials as characterized in (3.3) and Proposition 3. We study in this subsection statistical properties of estimates generated by iterative applications of (4.11).

We first provide a basic inequality for (4.11) as in Lemmas 1 and 2.

LEMMA 3. Let $\mathbf{b}^{(\text{new})}$ be as in (4.11), $\mathbf{w} = (\dot{\rho}_{\#}(\boldsymbol{\beta}^o; \lambda) - \mathbf{z})/\lambda_{s+1}^{(\text{new})}$, $\lambda = \lambda_{s+1}^{(\text{new})}$, $\lambda = \lambda^{(\text{new})}$, $\mathbf{h} = \mathbf{b}^{(\text{new})} - \boldsymbol{\beta}^o$ and $\mathbf{v}_{\text{carry}} = \dot{\rho}_{-}(\mathbf{b}^{(\text{old})}) - \dot{\rho}_{-}(\boldsymbol{\beta}^o)$ be the carryover error in the gradient. Then

$$(4.12) \quad \begin{aligned} & \mathbf{h}^T \bar{\Sigma} \mathbf{h} + (1 - \eta)\lambda \|\mathbf{h}_{\mathcal{S}^c}\|_{\#,s} \\ & \leq \mathbf{h}^T \bar{\Sigma} \mathbf{h} + (1 - \eta)\lambda \|\mathbf{h}_{\mathcal{S}^c}\|_{\#,s} + D_{+}(\mathbf{b}^{(\text{new})}, \boldsymbol{\beta}^o) \\ & = (\mathbf{h}_{\mathcal{S}^c}^T \mathbf{z}_{\mathcal{S}^c} - \eta\lambda \|\mathbf{h}_{\mathcal{S}^c}\|_{\#,s}) - \lambda \mathbf{h}_{\mathcal{S}}^T \mathbf{w}_{\mathcal{S}} - \mathbf{h}^T (\mathbf{v}^{(\text{new})} - \mathbf{v}_{\text{carry}}) \end{aligned}$$

when the favorable $\dot{\rho}_{\#}(\boldsymbol{\beta}^o; \lambda)$ is taken from the subdifferential (3.3), where $D_{+}(\mathbf{b}, \boldsymbol{\beta}) = (\mathbf{b} - \boldsymbol{\beta})^T \{\dot{\rho}_{+,\#}(\mathbf{b}; \lambda) - \dot{\rho}_{+,\#}(\boldsymbol{\beta})\}$ is the symmetric Bregman divergence for $\rho_{+,\#}(\mathbf{b}; \lambda)$.

Lemma 3 asserts that the basic inequality (4.6) holds for the LCA with $\bar{\kappa}(\hat{\boldsymbol{\beta}}) = 0$ and an extra carryover term $\mathbf{h}^T \mathbf{v}_{\text{carry}}$. Therefore, when

$$(4.13) \quad \Delta((1 - \eta)\xi|\mathcal{S}|^{1/2}, \mathbf{w}, \mathbf{v}^{(\text{new})} - \mathbf{v}_{\text{carry}}) < 1$$

as in (3.11), Theorem 6 is applicable to the LCA with $\mathbf{h} = \mathbf{b}^{(\text{new})} - \boldsymbol{\beta}^o$. We can always take $\rho_{-}(t)$ with $|(\partial/\partial t)^2 \rho_{-}(t)| \leq \bar{\kappa}(\rho)$, so that

$$(4.14) \quad \|\mathbf{v}_{\text{carry}}\|_2 \leq \|\dot{\rho}_{-}(\mathbf{b}^{(\text{old})}) - \dot{\rho}_{-}(\boldsymbol{\beta}^o)\|_2 \leq \bar{\kappa}(\rho) \|\mathbf{b}^{(\text{old})} - \boldsymbol{\beta}^o\|_2.$$

Next, we summarize the above discussion to show that Theorem 6 can be iteratively applied to the LCA. Let

$$(4.15) \quad \mathbf{b}^{(t)} \leftarrow \text{LCA}(\mathbf{b}^{(t-1)}, \rho^{(t)}, \lambda^{(t)}, \mathbf{v}^{(t)}),$$

where $\mathbf{b}^{(\text{new})} \leftarrow \text{LCA}(\mathbf{b}^{(\text{old})}, \rho, \boldsymbol{\lambda}^{(\text{new})}, \mathbf{v}^{(\text{new})})$ is the one-step approximate LCA (4.11) with a sorted concave penalty (3.7) generated from $\rho = \rho(t; \boldsymbol{\lambda})$.

THEOREM 8. Let $\mathbf{b}^{(t)}$ be as in (4.15) with penalty level $\boldsymbol{\lambda}^{(t)}$ and concavity $\kappa(\rho^{(t)}) \leq \kappa_0$. Let $\lambda^{(t)} = \lambda_{s+1}^{(t)}$ and $\mathbf{w}^{(t)} = (\dot{\rho}_{\#}^{(t)}(\boldsymbol{\beta}^o; \boldsymbol{\lambda}^{(t)}) - \mathbf{z})/\lambda^{(t)}$. Suppose

$$(4.16) \quad \Delta(r_1^{(t)}, \mathbf{w}^{(t)}, \mathbf{v}^{(t)}) \leq 1, \quad t = 1, \dots, t_{\text{fin}},$$

with the $\Delta(r, \mathbf{w}, \mathbf{v})$ in (3.10) and certain $r_1^{(t)} > 0$. Let $\mathbf{h}^{(t)} = \mathbf{b}^{(t)} - \boldsymbol{\beta}^o$ and $v_0 = \kappa_0 \|\mathbf{h}^{(0)}\|_2$ be the initial carryover error. Suppose the RE condition

$$\text{RE}_{\#}(\mathcal{S}; \xi, \gamma) \geq \kappa_0 \{\lambda^{(t)}/\lambda^{(t+1)}\} \{r_1^{(t)}\lambda^{(1)}/v_0 + 1\}$$

with $(1 - \eta)\xi s^{1/2} > v_0/\lambda^{(1)} + r_1^{(t)}$, $t = 1, \dots, t_{\text{fin}}$. Then, for $t \leq t_{\text{fin}}$

$$(4.17) \quad F(\mathbf{h}^{(t)}) \leq \frac{r_1^{(t)}\lambda^{(t)}}{\text{RE}_{\#}^2(\mathcal{S}; \xi, \gamma)} + \theta_0 \|\mathbf{h}^{(t-1)}\|_2 \leq \sum_{k=1}^t \frac{\theta_0^{t-k} r_1^{(k)} \lambda^{(k)}}{\text{RE}_{\#}^2(\mathcal{S}; \xi, \gamma)} + \theta_0^t \|\mathbf{h}^{(0)}\|_2$$

with $F(\mathbf{u}) = \max\{\|\mathbf{u}\|_2, (\gamma \mathbf{u}^T \bar{\Sigma} \mathbf{u})^{1/2}\}$ and $\theta_0 = \kappa_0/\text{RE}_{\#}^2(\mathcal{S}; \xi, \gamma)$.

We note that Theorem 8 is also applicable to the separable penalty (2.2) as $\lambda_1^{(t)} = \dots = \lambda_p^{(t)} = \lambda^{(t)}$ is allowed here.

To find an approximate solution of PLSE with sorted penalty $\rho_{\#}(\mathbf{b}; \boldsymbol{\lambda}_*)$, we may implement (4.15) with a fixed penalty family $\rho(x; \boldsymbol{\lambda})$ and $\boldsymbol{\lambda}^{(t)} \rightarrow \boldsymbol{\lambda}_*$,

$$\mathbf{X}^T(\mathbf{y} - \mathbf{X}\mathbf{b}^{(t)})/n = \dot{\rho}_{+, \#}(\mathbf{b}^{(t)}; \boldsymbol{\lambda}^{(t)}) - \dot{\rho}_{-}(\mathbf{b}^{(t-1)}) + \mathbf{v}^{(t)},$$

$t = 1, \dots, t_{\text{fin}}$, where $\boldsymbol{\lambda}_* = (\lambda_{*,1}, \dots, \lambda_{*,p})^T$ is a vector of target penalty levels, sorted or fixed. For example, we may take

$$(4.18) \quad \begin{aligned} \lambda_j^{(t)} &= \lambda_{*,j} \vee (\theta^{t-1} \lambda^{(1)}), & \lambda^{(1)} &= \|\mathbf{X}^T \mathbf{y}/n\|_{\infty}, \\ \mathbf{b}^{(0)} &= \mathbf{0}, & \theta &\in (0, 1). \end{aligned}$$

Alternatively, we may move gradually from the Lasso to $\rho_{\#}(\mathbf{b}; \boldsymbol{\lambda}_*)$, with

$$(4.19) \quad \begin{cases} \lambda_j^{(t)} = \max(\lambda_{*,j}, \theta^t \lambda_{*,1}), & j \leq p, \\ \rho^{(t)}(\mathbf{b}; \boldsymbol{\lambda}^{(t)}) = \theta^t \lambda_{*,1} \|\mathbf{b}\|_1 + (1 - \theta^t) \rho_{\#}(\mathbf{b}; \boldsymbol{\lambda}^{(t)}), \end{cases}$$

for a suitable $\theta < 1$. We recall that the prediction and squared ℓ_2 estimation error bounds are of the order $(\sigma^2/n)s \log p$ for the Lasso and $(\sigma^2/n)\{s + s_1 \log(p/s)\}$ for sorted concave penalties, where s_1 can be understood as the number of small nonzero coefficients.

COROLLARY 4. Let $\boldsymbol{\lambda}_* = (\lambda_{*,1}, \dots, \lambda_{*,p})^T$ be given by (3.13). Suppose the conditions of Corollary 3 hold with $r_2^2 \lesssim s_1$. Suppose the conditions of Theorem 8

hold for the penalty sequences (4.18) and (4.19). If the LCA (4.15) is implemented with the penalty sequence (4.18), then the right-hand side of (4.17) is of no greater order than (3.26) when

$$(4.20) \quad t_{\text{fin}} \geq 3 \left\lceil \frac{\log(\lambda^{(1)}/\lambda_{*,p})}{\log(1/\theta)} \vee \frac{\log(n\|\boldsymbol{\beta}^*\|_2^2/\{\sigma^2(s + s_1 \log(p/s))\})}{2\log(1/\theta_0)} \right\rceil.$$

If the LCA (4.15) is implemented with the penalty sequence (4.19), then the right-hand side of (4.17) is of no greater order than (3.26) when

$$t_{\text{fin}} \geq 3 \left\lceil \frac{\log((\log p)/\log(p/s))}{2\log(1/\theta)} \vee \frac{\log((s \log p)/(s + s_1 \log(p/s)))}{2\log(1/\theta_0)} \right\rceil.$$

The cost of computing a sorted concave PLSE by the LCA is determined by the number of required LCA steps, for example, t_{fin} in Corollary 4, the number of proximal-gradient iterations in each LCA step, for example, as in Algorithm 1, and the cost per proximal-gradient iteration. By Proposition 3, the cost per proximal-gradient iteration for sorted MCP is dominated by the cost of computing the gradient for the squared loss, which does depend on the choice of penalty. By a variation of the analysis in [31], the number of required proximal-gradient iterations in each LCA step is bounded by $C_1 \log(C_2 \sqrt{s})$ for some constants C_1 and C_2 . Thus, in view of (4.20), the cost of computing the sorted PLSE with the MCP is of the same order as that of computing the Lasso and other PLSE as in [31].

5. Discussion. In Sections 2 and 3, we have studied separable and sorted penalties in (2.1) and (3.2), respectively. However, our analysis does not require the penalty to have these specific forms as long as the penalty level and concavity of the penalty are controlled within proper levels. Such a more general version of our theory can be found in the first version of this paper on arXiv.

5.1. Subdifferential, penalty level and concavity. As in the cases of separable and sorted penalties, we may define the subdifferential $\partial \text{Pen}(\mathbf{b})$ of a penalty $\text{Pen}(\mathbf{b})$ as the set of all vectors $\mathbf{g}$ satisfying (3.4). As $\mathbf{g}^T \mathbf{u}$ is continuous in $\mathbf{g}$, $\partial \text{Pen}(\mathbf{b})$ is always a closed convex set.

Suppose the loss function $L(\mathbf{b})$ in (4.1) is differentiable with derivative $\dot{L}(\mathbf{b})$. It follows immediately from the definition of the subdifferential that

$$\liminf_{t \rightarrow 0+} \frac{1}{t} [\{L(\hat{\boldsymbol{\beta}} + t\mathbf{u}) + \text{Pen}(\hat{\boldsymbol{\beta}} + t\mathbf{u})\} - \{L(\hat{\boldsymbol{\beta}}) + \text{Pen}(\hat{\boldsymbol{\beta}})\}] \geq 0$$

for all $\mathbf{u} \in \mathbb{R}^p$ iff $-\dot{L}(\hat{\boldsymbol{\beta}}) \in \partial \text{Pen}(\hat{\boldsymbol{\beta}})$. This includes all local minimizers. Let $\dot{\text{Pen}}(\mathbf{b})$ denote a member of $\partial \text{Pen}(\mathbf{b})$. We say that $\hat{\boldsymbol{\beta}}$ is a local solution for minimizing (4.1) iff the following estimating equation is feasible:

$$(5.1) \quad -\dot{L}(\hat{\boldsymbol{\beta}}) = \dot{\text{Pen}}(\hat{\boldsymbol{\beta}}).$$

For simplicity, we define the penalty level of $\text{Pen}(\cdot)$ at a point $\mathbf{b} \in \mathbb{R}^p$ as

$$(5.2) \quad \lambda(\mathbf{b}) = \sup\{\lambda : \{\mathbf{g}_{S_b^c} : \mathbf{g} \in \partial \text{Pen}(\mathbf{b})\} \supseteq [-\lambda, \lambda]^{|S_b^c|}\}.$$

This definition is designed to achieve sparsity for solutions of (5.1). Although $\lambda(\mathbf{b})$ is a function of $\mathbf{b}$ in general, it depends solely on $\text{Pen}(\cdot)$ for many commonly used penalty functions. For sorted penalties, $\lambda(\mathbf{b}) = \lambda_{s+1}$ is a somewhat weaker penalty level as we discussed in Section 3.1.

We define the concavity of $\text{Pen}(\cdot)$ at $\mathbf{b}$, relative to a target $\tilde{\mathbf{b}}$, as

$$(5.3) \quad \bar{\kappa}(\mathbf{b}) = \bar{\kappa}(\mathbf{b}, \tilde{\mathbf{b}}) = \sup\{(\tilde{\mathbf{b}} - \mathbf{b})^T (\dot{\text{Pen}}(\mathbf{b}) - \dot{\text{Pen}}(\tilde{\mathbf{b}})) / \|\mathbf{b} - \tilde{\mathbf{b}}\|_2^2\}$$

with the convention $0/0 = 0$, where the supremum is taken over all choices $\dot{\text{Pen}}(\mathbf{b}) \in \partial \text{Pen}(\mathbf{b})$ and $\dot{\text{Pen}}(\tilde{\mathbf{b}}) \in \partial \text{Pen}(\tilde{\mathbf{b}})$. We use $\bar{\kappa} = \bar{\kappa}(\text{Pen}) = \sup_{\mathbf{b}, \tilde{\mathbf{b}}} \bar{\kappa}(\mathbf{b}, \tilde{\mathbf{b}})$ to denote the maximum concavity of $\text{Pen}(\cdot)$. This definition of concavity includes (2.3) for separable penalties and (3.5) for sorted penalties. However, we may also consider a relaxed concavity, given $s \geq \|\tilde{\mathbf{b}}\|_0$ and $\xi > 0$, as

$$(5.4) \quad \bar{\kappa}_{1,2}(\mathbf{b}; \xi) = \inf\{\bar{\kappa}_2(\mathbf{b}) + (1 + \xi)^2 s \bar{\kappa}_1(\mathbf{b})\},$$

where infimum is taken over all nonnegative $\bar{\kappa}_1(\mathbf{b})$ and $\bar{\kappa}_2(\mathbf{b})$ satisfying

$$(5.5) \quad (\mathbf{b} - \tilde{\mathbf{b}})^T (\dot{\text{Pen}}(\tilde{\mathbf{b}}) - \dot{\text{Pen}}(\mathbf{b})) \leq \bar{\kappa}_1(\mathbf{b}) \|\mathbf{h}\|_1^2 + \bar{\kappa}_2(\mathbf{b}) \|\mathbf{h}\|_2^2$$

for all choices of $\dot{\text{Pen}}(\mathbf{b})$ and $\dot{\text{Pen}}(\tilde{\mathbf{b}})$. This notion of concavity is more relaxed than the ℓ_2 one in (5.3) because $\bar{\kappa}_{1,2}(\mathbf{b}; \xi) \leq \bar{\kappa}(\mathbf{b}, \tilde{\mathbf{b}})$ always holds due to the option of picking $\bar{\kappa}_1(\mathbf{b}) = 0$.

5.2. Multivariate mixed penalties. For $\boldsymbol{\lambda} = (\lambda_1, \dots, \lambda_p)^T \in [0, \infty)^p$, let $\rho(\mathbf{b}; \boldsymbol{\lambda}) = \sum_{j=1}^p \rho(b_j; \lambda_j)$ be a separable penalty function with different penalty levels for different coefficients b_j . A multivariate mixed penalty can be constructed by mixing penalties $\rho(\mathbf{b}; \boldsymbol{\lambda})$ with a probability measure $\nu(d\boldsymbol{\lambda})$ and a real r_n as follows:

$$(5.6) \quad \rho_v(\mathbf{b}) = -r_n^{-1} \log \int \exp\{-r_n \rho(\mathbf{b}; \boldsymbol{\lambda})\} \nu(d\boldsymbol{\lambda}),$$

with the convention $\rho_v(\mathbf{b}) = \int \rho(\mathbf{b}; \boldsymbol{\lambda}) \nu(d\boldsymbol{\lambda})$ for $r_n = 0$. In particular, for $\rho(t; \lambda) = \lambda|t|$, $r_n = n$ and mixing distributions ν giving i.i.d. two-point components λ_j , (5.6) gives the spike-and-slab Lasso penalty as in [22].

By definition, the subdifferential of (5.6) can be written as

$$(5.7) \quad \partial \rho_v(\mathbf{b}) = \left\{ \frac{\int \dot{\rho}(\mathbf{b}; \boldsymbol{\lambda}) \exp\{-r_n \rho(\mathbf{b}; \boldsymbol{\lambda})\} \nu(d\boldsymbol{\lambda})}{\int \exp\{-r_n \rho(\mathbf{b}; \boldsymbol{\lambda})\} \nu(d\boldsymbol{\lambda})} : \dot{\rho}(\mathbf{b}; \boldsymbol{\lambda}) \in \partial \rho(\mathbf{b}; \boldsymbol{\lambda}) \right\}$$

with $\partial \rho(\mathbf{b}; \boldsymbol{\lambda})$ being the set of all vectors $\dot{\rho}(\mathbf{b}; \boldsymbol{\lambda}) = (\dot{\rho}(b_1; \lambda_1), \dots, \dot{\rho}(b_p; \lambda_p))^T$. If we treat $\exp\{-r_n \rho(t; \boldsymbol{\lambda})\} \nu(d\boldsymbol{\lambda}) / \int \exp\{-r_n \rho(t; \mathbf{x})\} \nu(d\mathbf{x})$ as conditional density of $\boldsymbol{\lambda}$ under a joint probability $\mathbb{P}_v$, we may write (5.7) as

$$(5.8) \quad \partial \rho_v(\mathbf{b}) = \{\mathbb{E}_v[\dot{\rho}(\mathbf{b}; \boldsymbol{\lambda}) | \mathbf{b}] : \dot{\rho}_j(\mathbf{b}; \boldsymbol{\lambda}) = (\dot{\rho}(b_1; \lambda_1), \dots, \dot{\rho}(b_p; \lambda_p))^T\}.$$

PROPOSITION 7. Let $\rho_v(\mathbf{b})$ be a mixed penalty in (5.6). Let $\mathcal{S}_b = \text{supp}(\mathbf{b})$ with $s_b = |\mathcal{S}_b| < p$. Then the concavity of $\rho_v(\mathbf{b})$ satisfies

$$(5.9) \quad \bar{\kappa}(\mathbf{b}) \leq \bar{\kappa}(\rho) + \sup_{\mathbf{u} \in [\tilde{\mathbf{b}}, \mathbf{b}]} \phi_{\max}(r_n \text{Cov}_v(\dot{\rho}(\mathbf{u}; \boldsymbol{\lambda}), \dot{\rho}(\mathbf{u}; \boldsymbol{\lambda})|\mathbf{u})),$$

where $[\tilde{\mathbf{b}}, \mathbf{b}] = \{\mathbf{u} = t\mathbf{b} + (1-t)\tilde{\mathbf{b}} : 0 \leq t \leq 1\}$, and (5.5) holds with

$$\bar{\kappa}_2(\mathbf{b}) \leq \bar{\kappa}(\rho), \quad \bar{\kappa}_1(\mathbf{b}) \leq (r_n \vee 0) \sup_{\mathbf{u}} \max_{1 \leq j \leq p} \text{Var}_v(\dot{\rho}(u_j; \lambda_j)|\mathbf{u}).$$

If the components of $\boldsymbol{\lambda}$ are independent given θ , then (5.5) holds with

$$\bar{\kappa}_2(\mathbf{b}) \leq \bar{\kappa}(\rho) + (r_n \vee 0) \sup_{\mathbf{u}} \max_{1 \leq j \leq p} \mathbb{E}_v[\text{Var}(\dot{\rho}(u_j; \lambda_j)|\mathbf{u}, \theta)|\mathbf{u}],$$

$$\bar{\kappa}_1(\mathbf{b}) \leq (r_n \vee 0) \sup_{\mathbf{u}} \max_{1 \leq j \leq p} \text{Var}_v(\mathbb{E}_v[\dot{\rho}(u_j; \lambda_j)|\mathbf{u}, \theta])|\mathbf{u}.$$

If in addition $\boldsymbol{\lambda}$ is exchangeable under $v(d\boldsymbol{\lambda})$, the penalty level of (5.6) is

$$(5.10) \quad \lambda(\mathbf{b}) = \mathbb{E}[\lambda_j|\mathbf{b}] \quad \forall j \notin \mathcal{S}_b.$$

Interestingly, (5.9) indicates that mixing $\rho(\mathbf{b}; \boldsymbol{\lambda})$ with $r_n < 0$ makes the penalty more convex.

For the nonseparable spike-and-slab Lasso [22], the prior is hierarchical where $\beta_j|\boldsymbol{\lambda} \sim (r_n \lambda_j/2)e^{-|t|r_n \lambda_j}$ are independent, $\lambda_j|\theta$ are i.i.d. with $\pi(\lambda_j = \lambda'|\theta) = \theta = 1 - \pi(\lambda_j = \lambda''|\theta)$ for some given constants λ' and λ'' , and $\theta \sim \pi(d\theta)$. As $\pi(\mathbf{0}|\theta) = \{(r_n/2)(\theta\lambda' + (1-\theta)\lambda'')\}^p$ and $\pi(\mathbf{0}) = \int \pi(\mathbf{0}|\theta)\pi(d\theta)$, the penalty can be written as

$$\rho_v(\mathbf{b}) = \frac{-1}{r_n} \log \frac{\pi(\mathbf{b})}{\pi(\mathbf{0})} = \frac{-1}{r_n} \log \int \exp \left\{ -r_n \sum_{j=1}^p \lambda_j |b_j| \right\} v(d\boldsymbol{\lambda}),$$

where $\lambda_j \in \{\lambda', \lambda''\}$ are i.i.d. given θ with $v(\lambda_j = \lambda'|\theta) = \theta\lambda'/\{\theta\lambda' + (1-\theta)\lambda''\}$ and $v(d\theta) = \pi(\mathbf{0}|\theta)\pi(d\theta)/\pi(\mathbf{0})$. The penalty level is given by

$$\lambda(\mathbf{b}) = \frac{\int \lambda_\theta \exp\{-r_n \sum_{j \in \mathcal{S}_b} \rho_\theta(b_j)\} v(d\theta)}{\int \exp\{-r_n \sum_{j \in \mathcal{S}_b} \rho_\theta(b_j)\} v(d\theta)},$$

with $\rho_\theta(t) = \{\theta\lambda'e^{-r_n\lambda'|t|} + (1-\theta)\lambda''e^{-r_n\lambda''|t|}\}/\{\theta\lambda' + (1-\theta)\lambda''\}$ and $\lambda_\theta = \{\theta(\lambda')^2 + (1-\theta)(\lambda'')^2\}/\{\theta\lambda' + (1-\theta)\lambda''\}$, and the relaxed concavity in (5.4)–(5.5) are bounded by $\bar{\kappa}_2(\mathbf{b}) = 0$, $\bar{\kappa}_1(\mathbf{b}) \leq r_n(\lambda' - \lambda'')^2/4$.

SUPPLEMENTARY MATERIAL

Supplement to “Sorted concave penalized regression” (DOI: [10.1214/18-AOS1759SUPP](https://doi.org/10.1214/18-AOS1759SUPP); .pdf). The Supplementary Material contains detailed proofs for Lemmas 1–3, Propositions 2–7, Theorems 3–6, 8 and Corollary 4. We omit proofs of Theorems 1, 2 and 7 and Proposition 1 as explained above or below their statements.

REFERENCES

- [1] AGARWAL, A., NEGAHBAN, S. and WAINWRIGHT, M. J. (2012). Fast global convergence of gradient methods for high-dimensional statistical recovery. *Ann. Statist.* **40** 2452–2482. [MR3097609](#)
- [2] BECK, A. and TBOULLE, M. (2009). A fast iterative shrinkage-thresholding algorithm for linear inverse problems. *SIAM J. Imaging Sci.* **2** 183–202. [MR2486527](#)
- [3] BELLEC, P. C., LECUÉ, G. and TSYBAKOV, A. B. (2018). Slope meets Lasso: Improved oracle bounds and optimality. *Ann. Statist.* **46** 3603–3642. [MR3852663](#)
- [4] BELLONI, A. and CHERNOZHUKOV, V. (2013). Least squares after model selection in high-dimensional sparse models. *Bernoulli* **19** 521–547. [MR3037163](#)
- [5] BICKEL, P. J., RITOV, Y. and TSYBAKOV, A. B. (2009). Simultaneous analysis of lasso and Dantzig selector. *Ann. Statist.* **37** 1705–1732. [MR2533469](#)
- [6] BOGDAN, M., VAN DEN BERG, E., SABATTI, C., SU, W. and CANDÈS, E. J. (2015). SLOPE—Adaptive variable selection via convex optimization. *Ann. Appl. Stat.* **9** 1103–1140. [MR3418717](#)
- [7] CANDÈS, E. and TAO, T. (2007). The Dantzig selector: Statistical estimation when p is much larger than n . *Ann. Statist.* **35** 2313–2351. [MR2382644](#)
- [8] CANDÈS, E. J. and TAO, T. (2005). Decoding by linear programming. *IEEE Trans. Inform. Theory* **51** 4203–4215. [MR2243152](#)
- [9] DALALYAN, A. and TSYBAKOV, A. B. (2008). Aggregation by exponential weighting, sharp pac-Bayesian bounds and sparsity. *Mach. Learn.* **72** 39–61.
- [10] EFRON, B., HASTIE, T., JOHNSTONE, I. and TIBSHIRANI, R. (2004). Least angle regression. *Ann. Statist.* **32** 407–499. With discussion, and a rejoinder by the authors. [MR2060166](#)
- [11] FAN, J. and LI, R. (2001). Variable selection via nonconcave penalized likelihood and its oracle properties. *J. Amer. Statist. Assoc.* **96** 1348–1360. [MR1946581](#)
- [12] FAN, J., LIU, H., SUN, Q. and ZHANG, T. (2018). I-LAMM for sparse learning: Simultaneous control of algorithmic complexity and statistical error. *Ann. Statist.* **46** 814–841. [MR3782385](#)
- [13] FENG, L. and ZHANG, C.-H. (2019). Supplement to “Sorted concave penalized regression.” DOI:10.1214/18-AOS1759SUPP.
- [14] HUANG, J. and ZHANG, C.-H. (2012). Estimation and selection via absolute penalized convex minimization and its multistage adaptive applications. *J. Mach. Learn. Res.* **13** 1839–1864. [MR2956344](#)
- [15] LOH, P.-L. and WAINWRIGHT, M. J. (2015). Regularized M -estimators with nonconvexity: Statistical and algorithmic theory for local optima. *J. Mach. Learn. Res.* **16** 559–616. [MR3335800](#)
- [16] MEINSHAUSEN, N. and BÜHLMANN, P. (2006). High-dimensional graphs and variable selection with the lasso. *Ann. Statist.* **34** 1436–1462. [MR2278363](#)
- [17] NEGAHBAN, S. N., RAVIKUMAR, P., WAINWRIGHT, M. J. and YU, B. (2012). A unified framework for high-dimensional analysis of M -estimators with decomposable regularizers. *Statist. Sci.* **27** 538–557. [MR3025133](#)
- [18] NESTEROV, Y. (2007). Gradient methods for minimizing composite functions. *Math. Program.* **140** 125–161. [MR3071865](#)
- [19] OSBORNE, M. R., PRESNELL, B. and TURLACH, B. A. (2000). A new approach to variable selection in least squares problems. *IMA J. Numer. Anal.* **20** 389–403. [MR1773265](#)
- [20] OSBORNE, M. R., PRESNELL, B. and TURLACH, B. A. (2000). On the LASSO and its dual. *J. Comput. Graph. Statist.* **9** 319–337. [MR1822089](#)
- [21] PARIKH, N. and BOYD, S. (2013). Proximal algorithms. In *Foundations and Trends in Optimization*.

- [22] ROČKOVÁ, V. and GEORGE, E. I. (2018). The spike-and-slab LASSO. *J. Amer. Statist. Assoc.* **113** 431–444. [MR3803476](#)
- [23] RUDELSON, M. and ZHOU, S. (2013). Reconstruction from anisotropic random measurements. *IEEE Trans. Inform. Theory* **59** 3434–3447. [MR3061256](#)
- [24] SU, W. and CANDÈS, E. (2016). SLOPE is adaptive to unknown sparsity and asymptotically minimax. *Ann. Statist.* **44** 1038–1068. [MR3485953](#)
- [25] SUN, T. and ZHANG, C.-H. (2012). Scaled sparse linear regression. *Biometrika* **99** 879–898. [MR2999166](#)
- [26] SUN, T. and ZHANG, C.-H. (2013). Sparse matrix inversion with scaled lasso. *J. Mach. Learn. Res.* **14** 3385–3418. [MR3144466](#)
- [27] TIBSHIRANI, R. (1996). Regression shrinkage and selection via the lasso. *J. Roy. Statist. Soc. Ser. B* **58** 267–288. [MR1379242](#)
- [28] TROPP, J. A. (2006). Just relax: Convex programming methods for identifying sparse signals in noise. *IEEE Trans. Inform. Theory* **52** 1030–1051. [MR2238069](#)
- [29] VAN DE GEER, S. A. and BÜHLMANN, P. (2009). On the conditions used to prove oracle results for the Lasso. *Electron. J. Stat.* **3** 1360–1392. [MR2576316](#)
- [30] WAINWRIGHT, M. J. (2009). Sharp thresholds for high-dimensional and noisy sparsity recovery using ℓ_1 -constrained quadratic programming (Lasso). *IEEE Trans. Inform. Theory* **55** 2183–2202. [MR2729873](#)
- [31] WANG, Z., LIU, H. and ZHANG, T. (2014). Optimal computational and statistical rates of convergence for sparse nonconvex learning problems. *Ann. Statist.* **42** 2164–2201. [MR3269977](#)
- [32] YE, F. and ZHANG, C.-H. (2010). Rate minimaxity of the Lasso and Dantzig selector for the ℓ_q loss in ℓ_r balls. *J. Mach. Learn. Res.* **11** 3519–3540. [MR2756192](#)
- [33] ZHANG, C.-H. (2010). Nearly unbiased variable selection under minimax concave penalty. *Ann. Statist.* **38** 894–942. [MR2604701](#)
- [34] ZHANG, C.-H. and HUANG, J. (2008). The sparsity and bias of the LASSO selection in high-dimensional linear regression. *Ann. Statist.* **36** 1567–1594. [MR2435448](#)
- [35] ZHANG, C.-H. and ZHANG, T. (2012). A general theory of concave regularization for high-dimensional sparse estimation problems. *Statist. Sci.* **27** 576–593. [MR3025135](#)
- [36] ZHANG, T. (2010). Analysis of multi-stage convex relaxation for sparse regularization. *J. Mach. Learn. Res.* **11** 1081–1107. [MR2629825](#)
- [37] ZHAO, P. and YU, B. (2006). On model selection consistency of Lasso. *J. Mach. Learn. Res.* **7** 2541–2563. [MR2274449](#)
- [38] ZOU, H. and LI, R. (2008). One-step sparse estimates in nonconcave penalized likelihood models. *Ann. Statist.* **36** 1509–1533. [MR2435443](#)

SCHOOL OF DATA SCIENCE
CITY UNIVERSITY OF HONG KONG
83 TAT CHEE AVENUE, KOWLOON TONG
HONG KONG
E-MAIL: lfengstat@gmail.com

DEPARTMENT OF STATISTICS AND BIOSTATISTICS
RUTGERS UNIVERSITY
PISCATAWAY, NEW JERSEY 08854
USA
E-MAIL: czhang@stat.rutgers.edu