

SLOPE for Sparse Linear Regression: Asymptotics and Optimal Regularization

Hong Hu and Yue M. Lu

Abstract

In sparse linear regression, the SLOPE estimator generalizes LASSO by penalizing different coordinates of the estimate according to their magnitudes. In this paper, we present a precise performance characterization of SLOPE in the asymptotic regime where the number of unknown parameters grows in proportion to the number of observations. Our asymptotic characterization enables us to derive the fundamental limits of SLOPE in both estimation and variable selection settings. We also provide a computationally feasible way to optimally design the regularizing sequences such that the fundamental limits are reached. In both settings, we show that the optimal design problem can be formulated as certain infinite-dimensional convex optimization problems, which have efficient and accurate finite-dimensional approximations. Numerical simulations verify all our asymptotic predictions. They demonstrate the superiority of our optimal regularizing sequences over other designs used in the existing literature.

I. INTRODUCTION

A. Motivation and Problem Setup

In sparse linear regression, we seek to estimate a sparse vector $\beta \in \mathbb{R}^p$ from

$$\mathbf{y} = \mathbf{A}\beta + \mathbf{w}, \quad (1)$$

where $\mathbf{A} \in \mathbb{R}^{n \times p}$ is the design matrix and $\mathbf{w}$ denotes the observation noise. In this paper, we study the *sorted ℓ_1 penalization estimator* (SLOPE) [2] (see also [3], [4]), a new paradigm for sparse linear regression. Given a non-decreasing regularization sequence $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_p]^\top$ with $0 \leq \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_p$, SLOPE estimates β by solving the following optimization problem

$$\hat{\beta} \in \underset{\mathbf{b}}{\operatorname{argmin}} \frac{1}{2} \|\mathbf{y} - \mathbf{A}\mathbf{b}\|_2^2 + \sum_{i=1}^p \lambda_i |b|_{(i)}, \quad (2)$$

where $|b|_{(1)} \leq |b|_{(2)} \leq \dots \leq |b|_{(p)}$ is a reordering of the absolute values $|b_1|, |b_2|, \dots, |b_p|$ in increasing order. In [2], the regularization term $J_\lambda(\mathbf{b}) \stackrel{\text{def}}{=} \sum_{i=1}^p \lambda_i |b|_{(i)}$ is referred to as the “sorted ℓ_1 norm” of $\mathbf{b}$. The same regularizer

H. Hu and Y. M. Lu are with the John A. Paulson School of Engineering and Applied Sciences, Harvard University, Cambridge, MA 02138, USA. (e-mails: honghu@g.harvard.edu and yuelu@seas.harvard.edu). This work was supported by the Harvard FAS Dean’s Fund for Promising Scholarship, and by the US National Science Foundation under grants CCF-1718698 and CCF-1910410. The preliminary version of this work has been presented at the 2019 Signal Processing with Adaptive Sparse Structured Representations (SPARS) workshop and the 2019 IEEE International Symposium on Information Theory (ISIT) [1].

was independently developed in a different line of work [3]–[6], where the motivation is to promote group selection in the presence of correlated covariates.

The classical LASSO estimator is a special case of SLOPE. It corresponds to using a constant regularization sequence, *i.e.*, $\lambda_1 = \lambda_2 = \dots = \lambda_p = \lambda$. However, with more general λ -sequences, SLOPE has the flexibility to penalize different coordinates of the estimate according to their magnitudes. This adaptivity endows SLOPE with some nice *statistical* properties that are not possessed by LASSO. For example, it is shown in [7], [8] that SLOPE achieves the minimax ℓ_2 estimation rate with high probability. When applied in variable selection problem, SLOPE is shown to control the false discovery rate (FDR) for orthogonal design matrices [2], which is not the case for LASSO. In addition, the new regularizer $J_\lambda(\mathbf{b})$ is still a norm [2], [4]. Thus, the optimization problem associated with SLOPE remains convex, and it can be efficiently solved by using *e.g.*, proximal gradient descent [2], [4].

Although the flexible regularization of SLOPE creates the hope of potential performance enhancement, to fully realize SLOPE's potential, we have to carefully design the regularizing sequence λ . Note that this is equivalent to specifying the empirical distribution of λ . Popular choices in the previous works include delta distribution (*i.e.*, LASSO), uniform distribution [5], chi-distribution [9], etc. These regularization schemes are mostly devised based on statistical insights gained from simpler models and they are indeed superior than LASSO in several applications. However, the success of these regularizing sequences provide no quantitative answer to the following two questions:

- 1) What is the fundamental limit of SLOPE?
- 2) How to optimally design λ to reach the fundamental limit?

The aforementioned studies on analyzing SLOPE provide very limited information for us to address the above two questions, since in these works, the SLOPE's performance is characterized in an order-wise manner, which contains loose constants. What we need is an *exact* performance characterization of SLOPE estimator, which is still absent in the existing literature. On the other hand, however, exact asymptotic analysis has been carried out for LASSO [10], [11] and several other regularized regression techniques [12]–[15], under certain statistical assumptions on the sensing matrix $\mathbf{A}$. One key feature of all these results is that the performance in the originally high-dimensional model can be well-captured by some low dimensional problems, which are much easier to handle. The technical hurdle that has precluded a similar treatment for SLOPE is that unlike all the regularizer considered in these works, the SLOPE norm $J_\lambda(\mathbf{x})$ is *non-separable*: it cannot be written as a sum of component-wise functions, *i.e.*, $J_\lambda(\mathbf{x}) \neq \sum_{i=1}^p J_i(x_i)$. This makes a similar low-dimensional reduction more challenging.

B. Main Contributions

In this paper, we answer the questions raised above. Our main contributions are listed as follows:

1) *Asymptotic separability*: As mentioned above, the main obstacle in analyzing SLOPE asymptotics is the non-separability of SLOPE regularizer $J_\lambda(\mathbf{b}) = \sum_{i=1}^p \lambda_i |b|_{(i)}$. We overcome this challenge by showing that the proximal operator of $J_\lambda(\mathbf{b})$ is *asymptotically separable*. To be more concrete, we first give a technically light overview of this result. The *proximal operator* of $J_\lambda(\mathbf{b})$ is defined as:

$$\text{Prox}_\lambda(\mathbf{y}) \stackrel{\text{def}}{=} \underset{\mathbf{x}}{\text{argmin}} \frac{1}{2} \|\mathbf{y} - \mathbf{x}\|_2^2 + J_\lambda(\mathbf{x}) \quad (3)$$

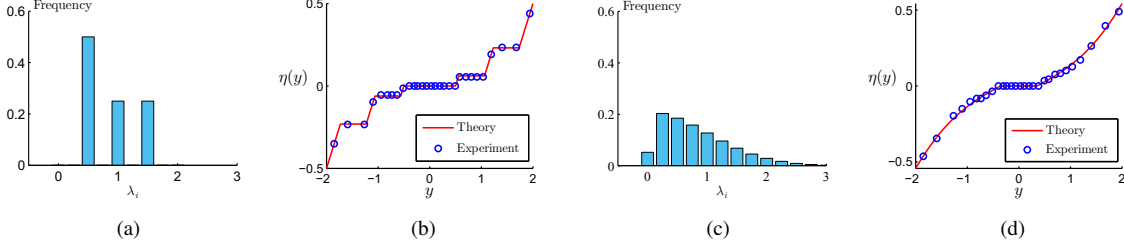


Figure 1: (a) and (c): The histograms of two different λ -sequences. (b) and (d): Sample points of $(y_i, [\text{Prox}_\lambda(\mathbf{y})]_i)$ (the blue dots) compared against the limiting scalar functions $\eta(y)$ (the red curves). In this experiment, $p = 1024$ and $y_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, 1)$. For better visualization, we randomly sample 3% of all $(y_i, [\text{Prox}_\lambda(\mathbf{y})]_i)$.

In the case of LASSO, where we choose $\lambda_1 = \lambda_2 = \dots = \lambda_p = \lambda$, characterizing $\text{Prox}_\lambda(\mathbf{y})$ is easy, since the optimization in (3) is equivalent to p scalar problems: $\sum_{i=1}^p \min_{x_i} \frac{1}{2}(y_i - x_i)^2 + \lambda|x_i|$. Correspondingly, the proximal operator is *separable*: $[\text{Prox}_\lambda(\mathbf{y})]_i = \text{sign}(y_i) \max(|y_i| - \lambda, 0)$. In other words, the i th element of $\text{Prox}_\lambda(\mathbf{y})$ is solely determined by y_i . However, this separability property does not hold for a general regularizing sequence. When p is finite, $[\text{Prox}_\lambda(\mathbf{y})]_i$ depends not only on y_i but also on other elements of $\mathbf{y}$. As one of the core results in this paper, we show that if the empirical distributions of $\mathbf{y}$ and λ converge as $p \rightarrow \infty$, then

$$\frac{1}{p} \|\text{Prox}_\lambda(\mathbf{y}) - \eta(\mathbf{y})\|^2 \rightarrow 0,$$

where η is a limiting scalar function that is uniquely determined by the limiting empirical measures of $\mathbf{y}$ and λ (for the exact form, see Proposition 1). This result is illustrated in Fig. 1, where we compare the actual proximal operator $\text{Prox}_\lambda(\mathbf{y})$ and the limiting scalar function $\eta(y)$, for two different λ -sequences shown in Fig. 1a and Fig. 1c. It can be seen that under a moderate dimension, the proximal operator $\text{Prox}_{\tau\lambda}(\mathbf{y})$ can already be very accurately approximated by $\eta(\mathbf{y})$.

2) *Exact characterization*: The asymptotic separability allows us to obtain the exact characterization of SLOPE's performance in the linear asymptotic regime: $n, p \rightarrow \infty$ and $n/p \rightarrow \delta$, under the assumption that sensing matrix $\mathbf{A}$ is generated from i.i.d. Gaussian. On a high level, our main results show that the joint empirical distribution of $\{(\hat{\beta}_i, \beta_i)\}_{i=1}^p$ converges to a well-defined limiting measure (the precise description can be found in Theorem 1). Note that the performance metrics of interests such as mean square error (MSE), type-I error, power are all functional of the empirical measure $\{(\hat{\beta}_i, \beta_i)\}_{i=1}^p$. Therefore, this makes it possible us to compute the high-dimensional limits of all these quantities. Compared with the probabilistic bounds derived in previous work, our results are asymptotically *exact*.

3) *Fundamental limits and optimal regularization*: The exact asymptotic characterization finally enables us to derive the fundamental limits of SLOPE in both estimation and variable selection tasks: (1) the minimum MSE that can be achieved by SLOPE; and (2) the highest possible power achievable under any given level of Type-I error. Moreover, we show that in both cases, the optimal λ sequence can be obtained by solving certain infinite-dimensional convex optimization problems, which have efficient and accurate finite-dimensional approximations. It

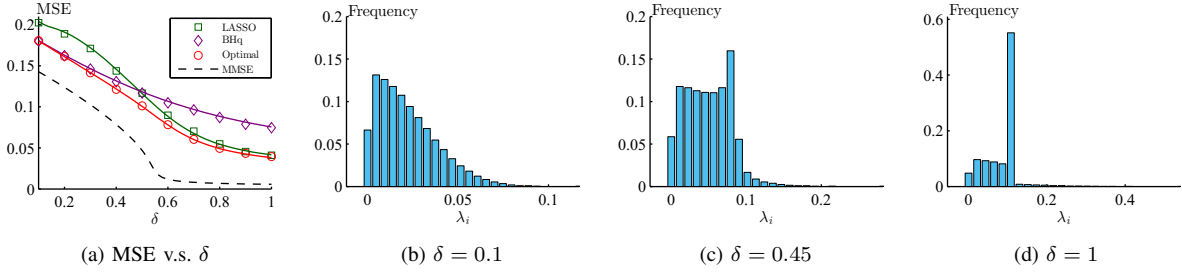


Figure 2: (a): Theoretical predictions (solid lines) v.s. empirical results. Here, β_i are i.i.d. Bernoulli random variables with $\mathbb{P}(\beta_i = 1) = 0.2$ and $w_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, 0.04)$. In our simulation, we choose $p = 2048$ and the empirical results are averaged over 20 independent trials. (b)-(d): Empirical distributions of optimal regularizing sequences under 3 different sampling ratios.

is worth mentioning that a caveat of our optimal design is that it requires knowing the limiting empirical measure of β (e.g., the sparsity level and the distribution of its nonzero coefficients). For this reason, our results are *oracle optimal*. However, it provides the first step towards optimal sequence designs under more realistic setting, where no or only limited information about β is available.

An illustration of asymptotic characterization and optimality results stated above are presented in Fig. 2. We consider three different regularizing sequences: LASSO, BHq sequence proposed in [9] and the optimal sequence given by Proposition 4 below. In Fig. 2a, we plot the empirical MSEs and compare them with the theoretical results. We can see they match well under all settings. Moreover, all the recorded MSE values are lower bounded by the fundamental limits predicted by our theory (red curve in the figure) and they can be achieved by the optimally designed sequences (red circles in the figure). For comparison, we also enclose the curve of minimum mean square error (MMSE) of linear Gaussian model, which was derived in [16], [17]. Finally, to help the readers get a sense of what the optimal regularizing sequences look like, in Fig. 2b-2d we plot their empirical distributions under 3 different sampling ratios δ . Interestingly, we can find they exhibit very different distributions as we change δ .

C. Related Work

1) *Exact asymptotic characterization*: There has been a growing body of works studying the exact asymptotics in high-dimensional statistical problems under random design assumptions. A partial list of these works includes [10], [18]–[30]. One distinct feature of these type of results is that they provide sharp performance guarantee that does not contain loose constants. From a technical viewpoint, these works are built on powerful tools including statistical physics [31], [32], approximate message passing (AMP) [19], [20], Gaussian width or statistical dimensions [21], [25], leave-one-out analysis [13], [24], Gordon’s Gaussian comparison lemma [33], etc. Our main asymptotic characterization is proved based on convex min-max Gaussian theorem (CGMT) [12], [26], [34], which is a tight version of Gordon’s comparison lemma in the convex setting. This framework was developed through a series of works [12], [26], [29], [34] and have now been successfully applied in a variety of problems such as binary

detection [14], regularized M-estimation [26], [29], phase retrieval [28], [35] and high-dimensional classification [36]–[38].

2) *Optimal M estimation in high dimensions*: The optimality part of this work falls within the line of research pursuing the optimal M-estimation in high-dimensional regression. The general form of M-estimator is as follows:

$$\hat{\beta} \in \underset{\mathbf{b}}{\operatorname{argmin}} \sum_{i=1}^n \ell(y_i, \mathbf{a}_i^\top \mathbf{b}) + r(\mathbf{b}) \quad (4)$$

and the question is what is the optimal statistical performance achievable by (4) and how to optimally design the loss function ℓ and the regularizer r . The exact asymptotic characterizations open up the possibility of obtaining a precise answer to the above question. This line of research is initiated by the papers [39] and [27], where the authors study the fundamental limits of the unregularized M-estimator (*i.e.*, the case when $r = 0$) in the linear model. In particular, a computational feasible recipe is provided in [39] for constructing the optimal loss function ℓ that minimizes the estimation errors. Similar types of results are also recently established for the binary models [40]. When a regularizer is included, the optimal performance of (4) in the linear model is studied in [41] and recently extended to binary model for the special case of quadratic regularization [42], [43]. In the meantime, a series of papers study the optimal ℓ_q -norm regularized least square regression [15], [44], [45]. In some limiting regimes, explicit answers are provided regarding the optimal choice of q . Note that all the aforementioned works consider the separable regularizer: $r(\mathbf{b}) = \sum_{i=1} r_i(b_i)$, while SLOPE regularizer considered in this paper is not separable.

Closely related with current work is the paper by Celentano and Montanari [46]. One of their main results is on the optimal estimation performance achievable by quadratic loss regularized by any lower semi-continuous, proper, convex and symmetric¹ function. It is not hard to check that SLOPE norm belongs to this family of functions. In fact, the optimality results presented in their paper and ours share a very similar form. We will elaborate more on this in Sec. IV-A.

3) *Three Parallel works*: Finally, we mention three parallel works that also study the limiting behavior of SLOPE under the same asymptotic setting.

- 1) From an algorithmic perspective, [47] consider solving the SLOPE minimization problem (2) using the AMP algorithm. By relating the stationary point of AMP iterations to SLOPE estimator, they also establish the same characterization (as shown in Theorem 1 below). In the proof, they also utilize the asymptotic separability property proved in Proposition 1.
- 2) The CGMT framework is also applied in [48] to obtain the limiting mean square errors (MSE) of SLOPE, together with a finite-sample concentration bound. The authors quantitatively compares the MSEs of different regularizing sequences in some limiting regimes. In particular, it is shown that in the high SNR regimes, LASSO regularization is optimal. A major difference from our work is that they do not exploit the asymptotic separability of SLOPE and the optimal performance in the general regime is not addressed.

¹Symmetric means $r(\mathbf{b})$ is permutation invariant to coordinates of $\mathbf{b}$.

- 3) In [49], the asymptotic separability properties is further extended to all lsc, proper, convex and symmetric regularizers using an elegant lifting and embedding idea. A finite-sample concentration bound is also given. Using the general asymptotic separability results, the author proves a conjecture in [46]: the MSE lower bound achievable by *non-separable* convex symmetric regularizers will be the same if we are restricted to the *separable* convex regularizers. However, the performance of variable selection is not addressed.

D. Notations

For a vector $\mathbf{x} \in \mathbb{R}^p$ and a scalar function $f(\cdot) : \mathbb{R} \rightarrow \mathbb{R}$, $f(\mathbf{x})$ means $f(\cdot)$ is applied to vector $\mathbf{x}$ coordinate-wise. $\|\mathbf{x}\|$ denotes the ℓ_2 norm, x_i (or $[\mathbf{x}]_i$) denotes the i th coordinate of $\mathbf{x}$ and $|x|_{(i)}$ (or $|\mathbf{x}|_{(i)}$) denotes the i th largest coordinate of $|\mathbf{x}|$. The Euclidean ball in $\mathbb{R}^p$ centered on $\mathbf{a}$ with radius $r \geq 0$ is denoted as: $\mathcal{B}_r(\mathbf{a}) := \{\mathbf{v} : \|\mathbf{v} - \mathbf{a}\| \leq r\}$ and $\mathcal{B}_r \stackrel{\text{def}}{=} \mathcal{B}_r(\mathbf{0})$. Also we define $\mathcal{B}_r^o(\mathbf{a}) \stackrel{\text{def}}{=} \{\mathbf{v} : \|\mathbf{v} - \mathbf{a}\| \geq r\}$.

For a probability measure μ , we denote $\text{Supp}(\mu)$ as its support. For random variables X, Y , we denote $\mu_{X,Y}$ and μ_X, μ_Y as their joint and marginal measures and F_X, F_Y as the corresponding (marginal) cumulative distribution function (CDF). The quantile function of random variable X is denoted as $F_X^{-1}(p)$, where $F_X^{-1}(p) \stackrel{\text{def}}{=} \inf\{x \in \mathbb{R} : F_X(x) \geq p\}$. Specifically, we use Φ and Φ^{-1} to denote the CDF and quantile function of standard Gaussian. For vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^p$, we denote $\mu_{\mathbf{x},\mathbf{y}}$ and $\mu_{\mathbf{x}}, \mu_{\mathbf{y}}$ as their joint and marginal empirical measures and $F_{\mathbf{x}}, F_{\mathbf{y}}$ as the corresponding (marginal) empirical CDF. Also we denote the empirical quantile function of $\mathbf{x}$ as $F_{\mathbf{x}}^{-1}$.

We denote $\mathcal{P}_q(\mathbb{R}^k)$, for some $q \geq 1$ and $k \in \mathbb{Z}^+$, as the space of all probability measures on $\mathbb{R}^k$ with bounded moments of order q , i.e., for any $\mu \in \mathcal{P}_q(\mathbb{R}^k)$, it holds that $\mathbb{E}_{\mu}(\|\mathbf{X}\|^q) < \infty$. For two measures $\mu, \nu \in \mathcal{P}_q(\mathbb{R}^k)$, their Wasserstein- q distance is defined as:

$$W_q(\mu, \nu) \stackrel{\text{def}}{=} \left(\inf_{\pi \in \Pi(\mu, \nu)} \mathbb{E} \|\mathbf{X} - \mathbf{Y}\|_2^q \right)^{1/q},$$

where $(X, Y) \sim \pi$ and $\Pi(\mu, \nu)$ is the set of all couplings of μ and ν .

E. Asymptotic Setting

There are four main objects in the description of our model and algorithm: (1) the unknown vector β ; (2) the design matrix $\mathbf{A}$; (3) the noise vector $\mathbf{w}$; and (4) the regularizing sequence λ . Since we study the asymptotic limit (with $p \rightarrow \infty$), we will consider a sequence of instances $\{\beta^{(p)}, \mathbf{A}^{(p)}, \mathbf{w}^{(p)}, \lambda^{(p)}\}_{p \in \mathbb{N}}$ with increasing dimensions p , where $\beta^{(p)}, \lambda^{(p)} \in \mathbb{R}^p$, $\mathbf{A}^{(p)} \in \mathbb{R}^{n \times p}$ and $\mathbf{w}^{(p)} \in \mathbb{R}^n$. A sequence of vectors $\{\mathbf{x}^{(p)}\}_{p \in \mathbb{Z}}$ (or $\{\mathbf{x}^{(p)}, \mathbf{y}^{(p)}\}_{p \in \mathbb{Z}}$), with p indexing the growing dimensions, is called a *converging sequence*, if its empirical measure $\mu_{\mathbf{x}^{(p)}}$ (or $\mu_{\mathbf{x}^{(p)}, \mathbf{y}^{(p)}}$) converges in Wasserstein-2 distance to a probability measure μ_X (or $\mu_{X,Y}$) as $p \rightarrow \infty$. For notational brevity, we will omit the superscript “ (p) ” when it is clear from the context.

F. Paper Outline

The rest of the paper is organized as follows. In Sec. II, we first prove the asymptotic separability of the proximal operator associated with $J_{\lambda}(\mathbf{x})$. This property allows us to derive our asymptotic characterization of SLOPE in Sec. III. Based on this analysis, we derive the fundamental limit and present the optimal design of the regularizing

sequence in Sec. IV. Numerical simulations are provided to verify our asymptotic characterizations. They also demonstrate the superiority of our optimal regularization over LASSO and BHq sequence in [7]. In Sec. V, we provide the proof of all our main results. We conclude the paper in Sec. VI and discuss some possible directions for future work.

II. PROXIMAL PROBLEM AND ASYMPTOTIC SEPARABILITY

We start by studying the following proximal problem:

$$\mathcal{M}_{\lambda}(\mathbf{y}; \tau) \stackrel{\text{def}}{=} \min_{\mathbf{x}} \frac{1}{2\tau} \|\mathbf{y} - \mathbf{x}\|_2^2 + J_{\lambda}(\mathbf{x}), \quad (5)$$

where $\tau > 0$, $\mathbf{y} \in \mathbb{R}^p$ and $J_{\lambda}(\mathbf{x}) = \sum_{i=1}^p \lambda_i |x|_{(i)}$, with $0 \leq \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_p$. $\mathcal{M}_{\lambda}(\mathbf{y}; \tau)$ in (5) is known as the *Moreau envelope* of $J_{\lambda}(\mathbf{x})$ evaluated at $\mathbf{y}$ and τ is the smoothing parameter. The unique minimizer of (5) is the *proximal operator* associated with $J_{\lambda}(\mathbf{x})$ under parameter τ . From (5), we know the proximal operator of $J_{\lambda}(\mathbf{x})$ is fully determined by $\mathbf{y}$ and $\tau\lambda$, so we simply denote it as $\text{Prox}_{\tau\lambda}(\mathbf{y})$. It turns out that the asymptotics of the original problem (2) is closely related to (5). Thus, as a preliminary step, we will first analyze its limiting properties.

To state our result, we introduce the following functional optimization problem. For $\mu_Y, \mu_{\Lambda} \in \mathcal{P}_2(\mathbb{R})$, with $\mathbb{P}(\Lambda \geq 0) = 1$, define

$$\mathcal{M}_{\mu_{\Lambda}}(\mu_Y; \tau) \stackrel{\text{def}}{=} \min_{g \in \mathcal{I}} \frac{1}{2\tau} \mathbb{E}_{\mu_Y} [Y - g(Y)]^2 + \int_0^1 F_{\Lambda}^{-1}(u) F_{|g(Y)|}^{-1}(u) du, \quad (6)$$

where

$$\mathcal{I} \stackrel{\text{def}}{=} \{g(y) \mid g(y) \text{ is odd, non-decreasing and 1-Lipschitz}\}. \quad (7)$$

Also we denote $\eta(\cdot; \mu_Y, \mu_{\tau\Lambda})$ as the optimal solution of (6). Comparing (6) with (5), we can intuitively interpret $\mathcal{M}_{\mu_{\Lambda}}(\mu_Y; \tau)$ and $\eta(\cdot; \mu_Y, \mu_{\tau\Lambda})$ as the functional-form Moreau envelope and proximal operator.

We are now ready to state our main result on the asymptotics of the proximal problem (5).

Proposition 1: Let $\{\mathbf{y}\}_{p \in \mathbb{N}}$ and $\{\lambda\}_{p \in \mathbb{N}}$ be two converging sequences, with limiting measures μ_Y and μ_{Λ} satisfying $\mathbb{P}(\Lambda \geq 0) = 1$. It holds that for any $\tau > 0$,

$$\frac{1}{p} \mathcal{M}_{\lambda}(\mathbf{y}; \tau) \rightarrow \mathcal{M}_{\mu_{\Lambda}}(\mu_Y; \tau) \quad (8)$$

and

$$\frac{1}{p} \|\text{Prox}_{\tau\lambda}(\mathbf{y}) - \eta(\mathbf{y}; \mu_Y, \mu_{\tau\Lambda})\|^2 \rightarrow 0, \quad (9)$$

where $\mathcal{M}_{\mu_{\Lambda}}(\mu_Y; \tau)$ and $\eta(\cdot; \mu_Y, \mu_{\tau\Lambda})$ are the optimal value and the unique (up to a set of measure zero with respect to μ_Y) optimal solution of (6).

The proof of Proposition 1 will be provided in Appendix V-A. We will also see that the limiting characterization of $\mathcal{M}_{\lambda}(\mathbf{y}; \tau)$ in (6) and the asymptotic separability of $\text{Prox}_{\tau\lambda}(\cdot)$ in (9) greatly facilitates our asymptotic analysis and the optimal design of λ , since this allows us to reduce the original high-dimensional problem to an equivalent one-dimensional problem, as in the LASSO case. Indeed, $\eta(\cdot; \mu_Y, \mu_{\tau\Lambda})$ in (9) is exactly the limiting scalar function

$\eta(\cdot)$ shown earlier in Fig. 1. We will still sometimes adopt the lighter notation $\eta(\cdot)$, when doing so causes no confusion.

Note that (6) is involved with an infinite-dimensional optimization, which typically permits no simple analytical solutions. To gain more intuition, before moving on, let us consider two examples where closed-form solutions do exist.

Example 1 (LASSO): The LASSO case corresponds to $\mathbb{P}(\Lambda = \lambda) = 1$. When $\tau = 1$, optimization in (6) then reduces to

$$\min_{g \in \mathcal{I}} \frac{1}{2} \mathbb{E} \left[\underbrace{|Y| - \lambda - g(|Y|)}_{:= f(|Y|)} \right]^2 + \text{constant}. \quad (10)$$

Note that the function $f(y) = y - \lambda$ in (10) is non-decreasing and 1-Lipschitz on $\mathbb{R}_{\geq 0}$ and $f(0) \leq 0$. It is not hard to show in this case, the optimal solution of (10) equals to

$$\begin{aligned} \eta(y; \mu_Y, \mu_\Lambda) &= \text{sign}(y) \max(f(|y|), 0) \\ &= \text{sign}(y) \max(|y| - \lambda, 0), \end{aligned}$$

which is exactly the soft-thresholding function.

Example 2 (BHq [9]): The BHq regularization corresponds to $\Lambda \sim \Phi^{-1}(1 - \frac{q}{2} + \frac{q}{2}U)$, where $q \in (0, 1]$ and U is uniformly distributed over $[0, 1]$. Then we have $F_\Lambda^{-1}(u) = \Phi^{-1}(1 - \frac{q}{2} + \frac{q}{2}u)$. Further, we consider $Y \sim \mathcal{N}(0, 1)$. It holds that $F_{|Y|}(y) = 2\Phi(y) - 1$ and $F_{|g(Y)|}^{-1}(F_{|Y|}(y)) = g(y)$, for $y \geq 0$. Therefore,

$$\int_0^1 F_\Lambda^{-1}(u) F_{|g(Y)|}^{-1}(u) du = \int_0^\infty \underbrace{\Phi^{-1}(1 - q + q \cdot \Phi(y))}_{:= \lambda(y)} g(y) dF_{|Y|}(y), \quad (11)$$

where we apply a change of variable $u = F_{|Y|}(y)$. In this case, (6) becomes

$$\begin{aligned} &\min_{g \in \mathcal{I}} \frac{1}{2} \mathbb{E} [|Y| - g(|Y|)]^2 + \mathbb{E} [\lambda(|Y|)g(|Y|)] \\ &= \min_{g \in \mathcal{I}} \frac{1}{2} \mathbb{E} [|Y| - \lambda(|Y|) - g(|Y|)]^2 + \text{constant}. \end{aligned}$$

On the other hand, by direct differentiation of $\lambda(y)$ in (11), we can get $\lambda'(y) = \frac{q\phi(y)}{\phi(\lambda(y))}$, where ϕ is the density function of standard Gaussian. It is not hard to verify $\lambda'(y) \in (0, 1]$ when $y \geq 0$. Therefore, $y \mapsto y - \lambda(y)$ is non-decreasing and 1-Lipschitz on $\mathbb{R}_{\geq 0}$. On the other hand, $\lambda(0) = \Phi^{-1}(1 - \frac{q}{2}) \geq 0$. Then following the same argument in Example 1, we get $\eta(y; \mu_Y, \mu_\Lambda) = \text{sign}(y) \max(|y| - \lambda(|y|), 0)$.

Remark 1: More generally, we can show when Y has a density supported on $\mathbb{R}$ and $y \mapsto y - F_\Lambda^{-1}(F_{|Y|}(y))$ is non-decreasing and 1-Lipschitz on $\mathbb{R}_{\geq 0}$, then $\eta(y; \mu_Y, \mu_\Lambda) = \text{sign}(y) \max(|y| - F_\Lambda^{-1}(F_{|Y|}(|y|)), 0)$. In some sense, $F_\Lambda^{-1}(F_{|Y|}(|y|))$ can be viewed as the equivalent regularization function. This equivalent regularization is adaptive to y . As a comparison, the regularization is a constant λ in the LASSO case.

III. ASYMPTOTIC CHARACTERIZATION OF SLOPE

Based on the asymptotic separability properties established in the last section, we are now ready to tackle the original optimization problem (2). We are going to obtain the precise characterizations of SLOPE in both estimation and variable selection problems.

A. Technical Assumptions

Our results are proved under the following assumptions:

- (A.1) The number of observations grows in proportion to p : $n^{(p)}/p \rightarrow \delta \in (0, \infty)$.
- (A.2) The number of nonzero elements in $\beta^{(p)}$ grows in proportion to p : $r_0^{(p)}/p \rightarrow \rho \in [0, 1]$.
- (A.3) The elements of $\mathbf{A}^{(p)}$ are i.i.d. Gaussian distribution: $A_{ij}^{(p)} \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \frac{1}{n})$.
- (A.4) $\{\beta^{(p)}\}_{p \in \mathbb{N}}$, $\{\mathbf{w}^{(p)}\}_{p \in \mathbb{N}}$ and $\{\boldsymbol{\lambda}^{(p)}\}_{p \in \mathbb{N}}$ are converging sequences. The limiting measures are denoted by μ_B , μ_W and μ_Λ , respectively. In addition, $\mathbb{P}(B \neq 0) = \rho$, $\sigma_w^2 = \mathbb{E}[W^2] > 0$ and $\mathbb{P}(\Lambda \neq 0) > 0$ when $\delta \leq 1$, where the probability $\mathbb{P}(\cdot)$ and the expectations $\mathbb{E}[\cdot]$ are all computed with respect to the limiting measures.

B. Asymptotic Performance of Estimation

The main goal of this section is to derive the limiting MSE of SLOPE: $\lim_{p \rightarrow \infty} \frac{1}{p} \|\hat{\beta} - \beta\|^2$. As in [10], we are going to prove a more general result, which characterizes the joint empirical measure of $(\hat{\beta}, \beta)$ through its action on *pseudo-Lipschitz* functions.

Definition 1 (Pseudo-Lipschitz function): A function $\psi : \mathbb{R}^2 \rightarrow \mathbb{R}$ is called pseudo-Lipschitz if $|\psi(\mathbf{x}) - \psi(\mathbf{y})| \leq L(1 + \|\mathbf{x}\| + \|\mathbf{y}\|)\|\mathbf{x} - \mathbf{y}\|$ for all $\mathbf{x}, \mathbf{y} \in \mathbb{R}^2$, where L is a positive constant.

To compute the limiting MSE, we just need to let $\psi(\mathbf{x}) = (x_1 - x_2)^2$, which is a pseudo-Lipschitz function by the above definition. The general theorem is as follows, whose proof is deferred to Sec. V-B.

Theorem 1: Assume (A.1) – (A.4) hold. For any pseudo-Lipschitz function ψ , we have

$$\frac{1}{p} \sum_{i=1}^p \psi(\hat{\beta}_i, \beta_i) \xrightarrow{\mathbb{P}} \mathbb{E}[\psi(\eta(Y_*; \mu_{Y_*}, \mu_{\tau_*\Lambda}), B)], \quad (12)$$

where $Y_* = B + \sigma_* H$ with $B \sim \mu_B$, $H \sim \mathcal{N}(0, 1)$ independent and η is the limiting scalar function defined in Proposition 1. In the above, the scalar pair (σ_*, τ_*) is the unique solution of the following equations:

$$\sigma^2 = \sigma_w^2 + \frac{1}{\delta} \mathbb{E}[(\eta(B + \sigma H; \mu_{B+\sigma H}, \mu_{\tau\Lambda}) - B)^2] \quad (13)$$

$$1 = \tau \left[1 - \frac{1}{\delta} \mathbb{E} \eta'(B + \sigma H; \mu_{B+\sigma H}, \mu_{\tau\Lambda}) \right]. \quad (14)$$

Theorem 1 essentially says that the joint empirical measure of $(\hat{\beta}^{(p)}, \beta^{(p)})$ converges to the law of $(\eta(Y_*; \mu_{Y_*}, \mu_{\tau_*\Lambda}), B)$. This means that although the original problem (2) is high-dimensional, its asymptotic performance can be succinctly captured by merely two scalar random variables. From (12) and (13), we know the limiting MSE equals to

$$\lim_{p \rightarrow \infty} \frac{1}{p} \|\hat{\beta} - \beta\|^2 = \delta(\sigma_*^2 - \sigma_w^2). \quad (15)$$

Readers familiar with the asymptotic analysis of LASSO will recognize that the forms of (13) and (14) look identical to the results of LASSO obtained in [10], [50]. Indeed, the proof of Theorem 1 directly applies the framework of analyzing LASSO asymptotics using convex Gaussian min-max theorem (CGMT) [26], [29], [50]. In a nutshell, the CGMT framework builds a connection between the asymptotics of the original high-dimensional problem (2) and the optimal solution of the following two-dimensional minimax problem:

$$\min_{\sigma \geq \sigma_w} \max_{\theta \geq 0} \frac{\theta}{2} \left(\frac{\sigma_w^2}{\sigma} + \sigma \right) - \frac{\theta^2}{2} + \frac{1}{\delta} \underbrace{\left[\lim_{p \rightarrow \infty} \frac{1}{p} \mathcal{M}_\lambda(\beta + \sigma \mathbf{h}; \frac{\sigma}{\theta}) - \frac{\theta \sigma}{2} \right]}_{:= \mathcal{F}(\sigma, \theta)}, \quad (16)$$

where β is the true signal vector in (1), $\mathbf{h} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$ and $\mathcal{M}_\lambda(\cdot; \cdot)$ is the Moreau envelope defined in (5). In fact, equation (13) and (14) corresponds to the first-order optimality condition of (16). Proposition 1 enables us to justify and explicitly compute the limit in (16), as well as the first-order derivatives $\frac{\partial \mathcal{F}(\sigma, \theta)}{\partial \sigma}$ and $\frac{\partial \mathcal{F}(\sigma, \theta)}{\partial \theta}$, which are crucial in obtaining the optimal point of (16).

C. Asymptotic Performance of Variable Selection

Next we study the asymptotic performance of SLOPE, when it is used as a variable selection methodology. Under this setting, the goal is to accurately select all the non-zero coordinates of β . Based on SLOPE estimate, we select the non-zero coordinates of estimate $\hat{\beta}$. Ideally, we hope that the selected set includes the non-zero coordinates of β , while do not contain zero coordinates of β . The usual performance metrics for this task include Type-I error, power, false discovery rate (FDR), etc. Most of these performance metrics can be expressed as a function of the sparsity level $r_0^{(p)}$ and the following quantities

$$R_0^{(p)} = \frac{1}{p} \sum_{i=1}^p \mathbb{I}_{\hat{\beta}_i=0}, \quad V^{(p)} = \frac{1}{p} \sum_{i=1}^p \mathbb{I}_{\hat{\beta}_i \neq 0, \beta_i=0}, \quad (17)$$

where $R_0^{(p)}$ and $V^{(p)}$ are the proportions of discoveries and false discoveries. In the following, we will adopt Type-I error and power as our performance metrics, which can be written as

$$\text{Type-I error} = \frac{V^{(p)}}{\max\{1 - r_0^{(p)}, 1/p\}}, \quad \text{Power} = \frac{1 - R_0^{(p)} - V^{(p)}}{\max\{r_0^{(p)}, 1/p\}}. \quad (18)$$

In order to study the asymptotics of these testing statistics, we need to obtain the limits of $R_0^{(p)}$ and $V^{(p)}$ in (17).

Note that the test functions involved in (17) ($\mathbb{I}_{x=0}$ and $\mathbb{I}_{x \neq 0, y=0}$) are discontinuous, so we can not directly apply (12) in Theorem 1 to compute $\lim_{p \rightarrow \infty} R_0^{(p)}$ and $\lim_{p \rightarrow \infty} V^{(p)}$. Further justifications are needed to obtain companion results for the testing-related statistics in (17). Before delving into technical descriptions, we first show that counter examples do exist where the quantities in (17) fail to converge, while the assumptions in Theorem 1 are still satisfied. This is different from the LASSO case, where the prediction (12) is shown to be still correct for the above non-smooth indicator functions [9].

Example 3 (A counter example): Consider μ_B being a spike-and-slab distribution: $\mu_B = 0.5 \cdot \delta_0 + 0.5 \cdot \mathcal{N}(1, 0.5^2)$ and $\{\beta_i\}_{i \in [p]}$ are i.i.d. generated from μ_B . Let (σ_*, τ_*) be the solution of (13)-(14) in the LASSO case, where $\mathbb{P}(\Lambda = 1) = 1$. Then we construct the following class of distribution of Λ , parameterized by $\vartheta \in [0, 1]$:

$$\Lambda_\vartheta = \begin{cases} \vartheta & |Y_*| < \vartheta \tau_*, \\ \frac{|Y_*|}{\tau_*} & \vartheta \tau_* \leq |Y_*| < \tau_*, \\ 1 & |Y_*| \geq \tau_*, \end{cases} \quad (19)$$

where $Y_* = B + \sigma_* H$. Here ϑ is a tuning parameter and $\vartheta = 1$ corresponding to the LASSO regularization. In Fig. 3, we plot the empirical $R_0^{(p)}$ and MSEs under different values of ϑ . It can be seen from Fig. 3b that for different values of ϑ , the empirical MSEs all concentrate around the predicted values from Theorem 1, when $\Lambda = 1$. On the contrary, from Fig. 3a we can find when $\vartheta < 1$, $R_0^{(p)}$ does not converge to $\mathbb{P}(\eta(Y_*) = 0)$, which is the limit

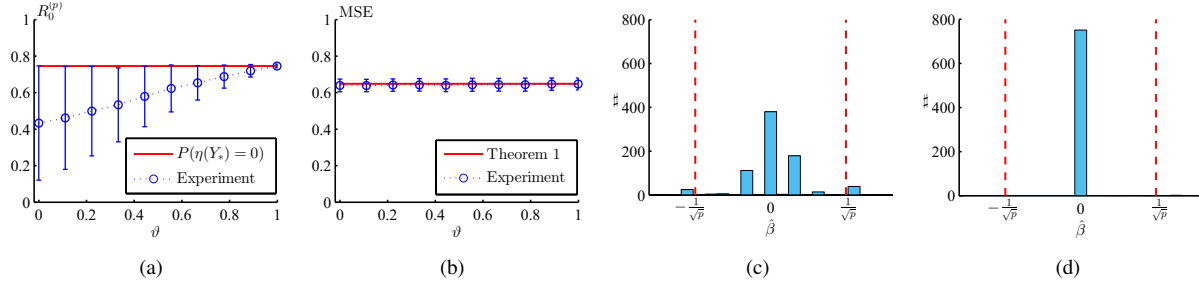


Figure 3: Counter example when $R_0^{(p)} \not\rightarrow \mathbb{P}(\eta(Y_*) = 0)$. In the experiment, $p = 1024$. The regularizing sequence λ is generated by reordering p i.i.d. samples of Λ_ϑ defined in (19). (a) and (b): empirical $R_0^{(p)}$ and MSE v.s. theoretical predictions based on Theorem 1, under different values of ϑ . The error bars in (a) and (b) are plotted using 1000 independent runs. (c) and (d): the histograms of $\hat{\beta}$ near zero when $\vartheta = 0$ and $\vartheta = 1$. When $\vartheta = 0$, it can be observed that two clumps of pseudo-zero entries appear within the tiny interval $[-\frac{1}{\sqrt{p}}, \frac{1}{\sqrt{p}}]$, while when $\vartheta = 1$, there is no pseudo-zero cluster.

indicated by Theorem 1. Moreover, as ϑ becomes smaller, the SLOPE estimator becomes less conservative and the variances of $R_0^{(p)}$ become increasingly notable. Also we can see $R_0^{(p)}$ does converge to $\mathbb{P}(\eta(Y_*) = 0)$, when $\vartheta = 1$.

We will explain the logic behind the construction of Λ_ϑ in Remark 2 below. The counter example above suggests that some additional constraints are needed, so that the testing statistics in (17) have well-defined limits and (12) can be used to compute these limits. It turns out that we just need one more condition to guarantee their convergence.

Proposition 2: Under the same settings as Theorem 1, define $q_0^* \stackrel{\text{def}}{=} \mathbb{P}(\eta(Y_*) = 0)$. If the following condition holds:

$$(R.1) \quad q_0^* = 0 \text{ or for any } t \in [0, q_0^*), \int_t^{q_0^*} F_{|Y_*|}^{-1}(u) du < \int_t^{q_0^*} F_{\tau_* \Lambda}^{-1}(u) du,$$

then we have

$$R_0^{(p)} \xrightarrow{\mathbb{P}} \mathbb{P}(\eta(Y_*) = 0) \text{ and } V^{(p)} \xrightarrow{\mathbb{P}} \mathbb{P}(\eta(Y_*) \neq 0, B = 0), \quad (20)$$

where $R_0^{(p)}$ and $V^{(p)}$ are defined in (17).

The proof of Proposition 2 will be provided in Appendix V-C, along with some explanations for condition (R.1) (see Remark 10).

Remark 2: In fact, Λ_ϑ in (19) is constructed so that condition (R.1) is violated for all $\vartheta < 1$. One can easily check that under the setting of Example 3, we have $q_0^* = F_{|Y_*|}(\tau_*) > 0$. From (19) we can get $F_{\tau_* \Lambda_\vartheta}^{-1}(u) = F_{|Y_*|}^{-1}(u)$, for all $u \in [F_{|Y_*|}(\vartheta \tau_*), F_{|Y_*|}(\tau_*)]$. Also due to the fact that Y_* is supported on $\mathbb{R}$, we have $F_{|Y_*|}(\vartheta \tau_*) < F_{|Y_*|}(\tau_*)$, when $\vartheta < 1$. Therefore, $\int_t^{q_0^*} F_{|Y_*|}^{-1}(u) du = \int_t^{q_0^*} F_{\tau_* \Lambda_\vartheta}^{-1}(u) du$ for any $t \in [F_{|Y_*|}(\vartheta \tau_*), q_0^*)$, where we have used $q_0^* = F_{|Y_*|}(\tau_*)$. This violates condition (R.1). On the other hand, we can also check when $\vartheta = 1$, i.e., in the LASSO case, condition (R.1) is satisfied. Indeed, in this case $\Lambda_\vartheta = 1$ and $F_{\tau_* \Lambda_\vartheta}^{-1}(u) = F_{\tau_*}^{-1}(u) = \tau_*$ for any $u \in [0, q_0^*]$. Besides, since $q_0^* = F_{|Y_*|}(\tau_*) > 0$ and Y_* is supported on $\mathbb{R}$, we get $F_{|Y_*|}^{-1}(u) < \tau_*$ for any $u \in [0, q_0^*)$. Therefore, $\int_t^{q_0^*} F_{|Y_*|}^{-1}(u) du < \int_t^{q_0^*} F_{\tau_* \Lambda_\vartheta}^{-1}(u) du$ for any $t \in [0, q_0^*)$.

Remark 3: In Example 3, a superficial reason for $R_0^{(p)} \not\rightarrow \mathbb{P}(\eta(Y_*) = 0)$ when $\vartheta < 1$ is that λ generated from such Λ will lead to many pseudo-zero entries in $\hat{\beta}$, i.e., entries that are very closed to 0, but not strictly 0. This is illustrated in Fig. 3c and 3d. In practice, the pseudo-zero effects can be mitigated by employing post-screening to $\hat{\beta}$. This is done by first specifying a threshold $h > 0$ and then setting all the entries in $\hat{\beta}$ with $|\hat{\beta}_i| < h$ to be zero. However, this creates a new problem of choosing the appropriate h . Our claim is that this problem can be completely avoided by adding an extra constraint on the regularizing sequence. Moreover, as will be clarified in Sec. IV-B, this additional constraint will not harm the diversity of our design choices.

Based on Proposition 2, we can now compute the limiting Type-I error and power of SLOPE.

Corollary 1: When $\mathbb{P}(B = 0) \in (0, 1)$, we have

$$\lim_{p \rightarrow \infty} \text{Type-I error} = \mathbb{P}(|\sigma_* H| \geq y_{\text{th}}^*) \quad (21)$$

and

$$\lim_{p \rightarrow \infty} \text{Power} = \mathbb{P}(|B + \sigma_* H| \geq y_{\text{th}}^* \mid B \neq 0), \quad (22)$$

where $y_{\text{th}}^* = \sup_{y \geq 0} \{y \mid \eta(y; \mu_{Y_*}, \mu_{\tau_* \Lambda}) = 0\}$.

The proof of Corollary 1 directly follows from (18), (20) and the assumption that $r_0^{(p)}/p \rightarrow \mathbb{P}(B \neq 0)$. Formulas (21) and (22) will be useful in Sec. IV-B, where we analyze the optimal performance of SLOPE for variable selection.

Remark 4: In Corollary 1, we require that $\mathbb{P}(B = 0) \in (0, 1)$. This means asymptotically, the proportions of zero and non-zero entries of β are both non-vanishing. We need this assumption on the distribution of B , because otherwise the limiting formula of Type-I error and power will involve with $\frac{0}{0}$ term, when we apply (18). This is beyond the scope of asymptotic setting considered in this paper.

IV. FUNDAMENTAL LIMITS AND OPTIMAL REGULARIZATION

Armed with the asymptotic characterizations in Theorem 1 and Proposition 2, we are now ready to analyze the optimal performance of SLOPE in both estimation and variable selection setting.

A. Estimation with Minimum MSE

We first turn to the problem of finding the minimum MSE achievable by SLOPE estimator and the corresponding optimal regularization. In the current asymptotic setting, this can be formulated as follows:

$$\inf_{\mu_\Lambda \in \mathcal{P}_\Lambda} \lim_{p \rightarrow \infty} \frac{1}{p} \|\hat{\beta} - \beta\|_2^2 \quad (23)$$

where $\mathcal{P}_\Lambda \stackrel{\text{def}}{=} \{\mu_\Lambda \mid \mu_\Lambda \in \mathcal{P}_2(\mathbb{R}) \text{ and } \mathbb{P}(\Lambda \neq 0) > 0, \text{ when } \delta \leq 1\}$ is the admissible set of μ_Λ , under which the asymptotic characterization in Theorem 1 is valid. By (15), solving (23) is equivalent to solving

$$\inf_{\mu_\Lambda \in \mathcal{P}_\Lambda} \sigma_*. \quad (24)$$

In the current context, σ_* should be understood as a function of μ_Λ , but for notational simplicity, we will drop its dependency on μ_Λ , when doing so causes no confusion.

Note that σ_* is determined by μ_Λ implicitly through the nonlinear fixed point equation (13)-(14), so a direct optimization over μ_Λ as in (24) is not viable. To proceed, a key observation from (13)-(14) is that the influence of μ_Λ is exerted only through the limiting scalar function η . In light of this, (24) can be alternatively solved via the following two-step scheme:

Step 1. Search over all *realizable* η such that there exists $\sigma, \tau > 0$ satisfying

$$\sigma^2 = \sigma_w^2 + \frac{1}{\delta} \mathbb{E}[(\eta(B + \sigma H) - B)^2] \quad (25)$$

$$1 = \tau \left(1 - \frac{1}{\delta} \mathbb{E}[\eta'(B + \sigma H)]\right) \quad (26)$$

and find optimal η^* that yields the minimum feasible σ . Denote the corresponding solution of (25)-(26) as (σ^*, τ^*) .

Step 2. Find corresponding μ_Λ such that $\eta(y; \mu_{B+\sigma^*H}, \mu_{\tau^*\Lambda}) = \eta^*(y)$.

Note that in Step 1, η is treated as an optimization variable that do not depend on other parameters, which greatly simplifies the original formulation (24). However, to implement this scheme, we still need to guarantee two things. First, the *realizable* set of η (as required in Step 1) needs to be decided. Second, for any realizable η , the corresponding Λ can be efficiently computed. These are both addressed in the following result.

Proposition 3: For a probability measure $\mu_Y \in \mathcal{P}_2(\mathbb{R})$, define

$$\mathcal{M}_{\mu_Y} \stackrel{\text{def}}{=} \{\eta(\cdot; \mu_Y, \mu_\Lambda) \mid \mu_\Lambda \in \mathcal{P}_2(\mathbb{R})\}, \quad (27)$$

where $\eta(\cdot; \mu_Y, \mu_\Lambda)$ is the limiting scalar function in Proposition 1. Then for any $\mu_Y \in \mathcal{P}_2(\mathbb{R})$, we have $\mathcal{M}_{\mu_Y} = \mathcal{I}$. Correspondingly, for any $f(y) \in \mathcal{I}$, we can take $\Lambda \sim |Y| - f(|Y|)$, with $Y \sim \mu_Y$, so that $\eta(y; \mu_Y, \mu_\Lambda) = f(y)$.

The proof of Proposition 3 will be presented in Appendix L. It is the key ingredient in proving our optimality results. It shows that, with different choices of μ_Λ , one can reach any non-decreasing and odd function that is Lipschitz continuous with constant 1. Clearly, the soft-thresholding functions associated with LASSO belongs to $\mathcal{M}_{\mu_Y}$, but the set $\mathcal{M}_{\mu_Y}$ is much richer. This is how SLOPE generalizes LASSO: it allows for more degrees of freedom in the regularization.

Based on Proposition 3, we are now ready to show the two-step scheme sketched above indeed yield a computationally feasible procedure to obtain the minimum MSE and the optimal μ_Λ . Before that, we first introduce the following function:

$$\begin{aligned} \mathcal{L}(\sigma) &\stackrel{\text{def}}{=} \inf_{f \in \mathcal{I}} \mathbb{E}[f(B + \sigma H) - B]^2 \\ &\text{s.t. } \delta^{-1} \mathbb{E}[f'(B + \sigma H)] \leq 1. \end{aligned} \quad (28)$$

We will see for any $\sigma > 0$, problem (28) is convex and there exists a unique optimal solution. Given $\mathcal{L}(\sigma)$, we then introduce the following equation on σ :

$$\mathcal{L}(\sigma) = \delta(\sigma^2 - \sigma_w^2). \quad (29)$$

As is shown in Proposition 4 below, the minimum limiting MSE is closely related with the minimum solution of equation (29).

Proposition 4: Under the same setting as Theorem 1, we have

- (a) For any $\sigma > 0$, problem (28) is convex and there exists a unique optimal solution $f_\sigma \in \mathcal{I}$.
- (b) $\mathcal{L}(\sigma)$ defined in (28) is continuous on $\mathbb{R}_{>0}$ and equation (29) always has a solution. The minimum solution $\sigma_0 \in [\sigma_w, \sqrt{\sigma_w^2 + \delta^{-1} \mathbb{E} B^2}]$.
- (c) Define $Y_0 := B + \sigma_0 H$ and $\tau_0 := [1 - \delta^{-1} \mathbb{E} f'_{\sigma_0}(Y_0)]^{-1}$. It always holds that

$$\lim_{p \rightarrow \infty} \frac{1}{p} \|\hat{\beta} - \beta\|_2^2 \geq \delta(\sigma_0^2 - \sigma_w^2). \quad (30)$$

Moreover, if $\delta^{-1} \mathbb{E}[f'_{\sigma_0}(Y_0)] < 1$, the equality in (30) can be attained, when μ_Λ is the law of

$$\frac{1}{\tau_0} [|Y_0| - f_{\sigma_0}(|Y_0|)]. \quad (31)$$

The proof of Proposition 4 is deferred to Appendix V-D. To solve the infinite-dimensional optimization problem (28) in practice, we can discretize over $\mathbb{R}$ and obtain a finite-dimensional approximation. Naturally, this finite-dimensional problem is still convex. In our simulation, we use an approximation with 2048 grids.

We have a couple of comments regarding Proposition 4 as follows.

Remark 5 (Interpretation of $\mathcal{L}(\sigma)$): Consider the optimization in (28):

$$\inf_{f \in \mathcal{I}} \mathbb{E}[f(B + \sigma H) - B]^2 \quad (32)$$

and we neglect the constraint $\delta^{-1} \mathbb{E}[f'(B + \sigma H)] \leq 1$ for a moment. From Proposition 1 and Proposition 3 we know minimization in (32) is equivalent to

$$\inf_{\mu_\Lambda \in \mathcal{P}_2(\mathbb{R})} \lim_{p \rightarrow \infty} \frac{1}{p} \|\text{Prox}_\lambda(\beta + \sigma \mathbf{h}) - \beta\|^2,$$

where $\mathbf{h} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$ and $\mu_\Lambda = \lim_{p \rightarrow \infty} \mu_\lambda$. In other words, we are estimating β from the noisy observation: $\mathbf{y} = \beta + \sigma \mathbf{h}$ using SLOPE and we want to find the optimal regularization (specified by its limiting distribution μ_Λ) such that the estimation error of β is minimized. Then $\mathcal{L}(\sigma)$ can be understood as the minimum MSE we can achieve, if we put an additional constraint on the average slope of limiting scalar function. On the other hand, if at the optimal solution f_σ , the constraint is inactive, i.e., $\delta^{-1} \mathbb{E}[f'_\sigma(B + \sigma H)] < 1$, then $\mathcal{L}(\sigma) = \inf_{f \in \mathcal{I}} \mathbb{E}[f(B + \sigma H) - B]^2$. This can be easily verified as follows. Assume there exists $f_\star \in \mathcal{I}$ such that $\mathbb{E}[f_\star(B + \sigma H) - B]^2 < \mathcal{L}(\sigma)$. Then consider the convex combination $f_t := t f_\star + (1 - t) f_\sigma$, for $t \in (0, 1)$. Clearly, $f_t \in \mathcal{I}$ and it is not hard to check for small enough t , $\delta^{-1} \mathbb{E}[f'_t(B + \sigma H)] \leq 1$. However, due to the convexity of objective function in (28),

$$\mathbb{E}[f_t(B + \sigma H) - B]^2 \leq t \underbrace{\mathbb{E}[f_\star(B + \sigma H) - B]^2}_{< \mathcal{L}(\sigma)} + (1 - t) \underbrace{\mathbb{E}[f_\sigma(B + \sigma H) - B]^2}_{= \mathcal{L}(\sigma)} < \mathcal{L}(\sigma),$$

which leads to a contradiction.

Remark 6 (Tightness of lower bound (30)): We require $\delta^{-1} \mathbb{E} f'_{\sigma_0}(B + \sigma_0 H) < 1$ so that the lower bound (30) is tight. The question is whether it is possible that $\delta^{-1} \mathbb{E} f'_{\sigma_0}(B + \sigma_0 H) = 1$. This will not happen when $\delta > 1$, since $f'_{\sigma_0} \leq 1$. When $\delta \leq 1$, we do not have a rigorous proof yet. Numerically, this never happens either. Here we provide an intuitive argument. Suppose for certain configurations of $(\delta, \rho, \sigma_w, \mu_B)$, we do have $\delta^{-1} \mathbb{E} f'_{\sigma_0}(B + \sigma_0 H) = 1$. Under this scenario, let us consider the following approximation of (28) and (29), parameterized by $\varepsilon > 0$:

$$\begin{aligned} \mathcal{L}_\varepsilon(\sigma) &\stackrel{\text{def}}{=} \inf_{f \in \mathcal{I}} \mathbb{E}[f(B + \sigma H) - B]^2 \\ &\text{s.t. } \delta^{-1} \mathbb{E}[f'(B + \sigma H)] \leq 1 - \varepsilon. \end{aligned} \quad (33)$$

and

$$\mathcal{L}_\varepsilon(\sigma) = \delta(\sigma^2 - \sigma_w^2). \quad (34)$$

Denote $\sigma_{0,\varepsilon}$ as the minimum solution of equation (34) and f_ε as the optimal solution of (33), when $\sigma = \sigma_{0,\varepsilon}$. If we take μ_Λ to be the law of

$$\frac{1}{\tau_{0,\varepsilon}}[|Y_{0,\varepsilon}| - f_\varepsilon(|Y_{0,\varepsilon}|)], \quad (35)$$

where $Y_{0,\varepsilon} = B + \sigma_{0,\varepsilon}H$ and $\tau_{0,\varepsilon} = [1 - \delta^{-1}\mathbb{E}f'_\varepsilon(Y_{0,\varepsilon})]^{-1} < \infty$. Then similar as Proposition 4, it is not hard to show $\lim_{p \rightarrow \infty} \frac{1}{p}\|\hat{\beta}_\varepsilon - \beta\|_2^2 = \delta(\sigma_{0,\varepsilon}^2 - \sigma_w^2)$, where $\hat{\beta}_\varepsilon$ denotes the corresponding estimator. Intuitively, we could also expect $\sigma_{0,\varepsilon} \rightarrow \sigma_0$ and $\tau_{0,\varepsilon} \rightarrow \infty$, as $\varepsilon \rightarrow 0$. This implies the MSE can be made arbitrarily close to the lower bound (30) using a sequence $\{\mu_{\Lambda,\varepsilon}\}_{\varepsilon>0}$ which converges to the probability mass at 0 as $\varepsilon \rightarrow 0$. Recall that we have assumed $\sigma_w > 0$, so this means the optimal regularization in a noisy overparameterized linear model should be vanishingly small, which is not likely the case.

Remark 7 (Comparison with [46], [49]): In [46], [49], the authors also analyze the problem of optimal estimation in the linear model (1) with i.i.d. Gaussian design. For the convenience of comparison, here we rephrase their results in our notations. The optimality they consider is with respect to the following class of estimator:

$$\left\{ \hat{\beta} \mid \hat{\beta} \in \operatorname{argmin}_{\mathbf{b}} \frac{1}{2}\|\mathbf{y} - \mathbf{A}\mathbf{b}\|_2^2 + r_p(\mathbf{b}), r_p \in \mathcal{C}_p \right\} \quad (36)$$

where

$$\mathcal{C}_p \stackrel{\text{def}}{=} \{r_p : \mathbb{R}^p \rightarrow \bar{\mathbb{R}} \mid r_p \text{ is lsc, proper, convex and symmetric}\}.$$

The optimal estimation within the class of $\mathcal{C}_p$ is formulated as:

$$\text{MSE}_{\text{cvx}} \stackrel{\text{def}}{=} \inf_{\forall p, r_p \in \mathcal{C}_p \cap \mathcal{W}_p} \liminf_{p \rightarrow \infty} \frac{1}{p} \|\hat{\beta} - \beta\|_2^2, \quad (37)$$

where $\mathcal{W}_p$ is some set that ensures $\hat{\beta}$ is unique². One of their main results states that under certain conditions, the minimum achievable limiting MSE defined in (37) satisfies: $\text{MSE}_{\text{cvx}} \geq \delta(\sigma_{\text{cvx}}^2 - \sigma_w^2)$, where

$$\sigma_{\text{cvx}}^2 = \sup\{\sigma^2 \mid \delta(\sigma^2 - \sigma_w^2) < \inf_{f \in \mathcal{J}} \mathbb{E}[f(B + \sigma H) - B]^2\}, \quad (38)$$

with $\mathcal{J} \stackrel{\text{def}}{=} \{f : \mathbb{R} \rightarrow \mathbb{R} \mid f \text{ is non-decreasing and 1-Lipschitz continuous}\}$. Comparing their results with ours, we can find the lower bounds in both settings follow the same type of characterization. Specifically, lying at the heart of this characterization is an optimization problem: $\inf_{f \in \mathcal{F}} \mathbb{E}[f(B + \sigma H) - B]^2$, which aims at finding the optimal estimator f of B under the noisy observation $B + \sigma H$. The only difference is on the feasible set $\mathcal{F}$: $\mathcal{F} = \mathcal{J}$ in (38), while $\mathcal{F} = \mathcal{I} \subset \mathcal{J}$ in (28). This agreement is not a coincidence, but related with the fact that the proximal operator of all functions in $\mathcal{C}_p$ is asymptotically separable as proved in [49]. In fact, f corresponds to the limiting proximal operator of the regularizer r_p . In our settings, r_p is chosen from the set of all possible sorted ℓ_1 norms (denoted by $\mathcal{S}_p$), while in their settings, it is chosen from the set $\mathcal{C}_p$. Correspondingly, $\mathcal{I}$ is the set of all limiting proximal operators associated with $\mathcal{S}_p$ and $\mathcal{J}$ is the one associated with $\mathcal{C}_p$. It is not hard to check $\mathcal{S}_p \subset \mathcal{C}_p$ and consequently, we have $\mathcal{I} \subset \mathcal{J}$.

²In fact, $\mathcal{W}_p$ corresponds to the tightness condition $\delta^{-1}\mathbb{E}[f'_{\sigma_0}(B + \sigma_0 H)] < 1$ in Proposition 4 (c).

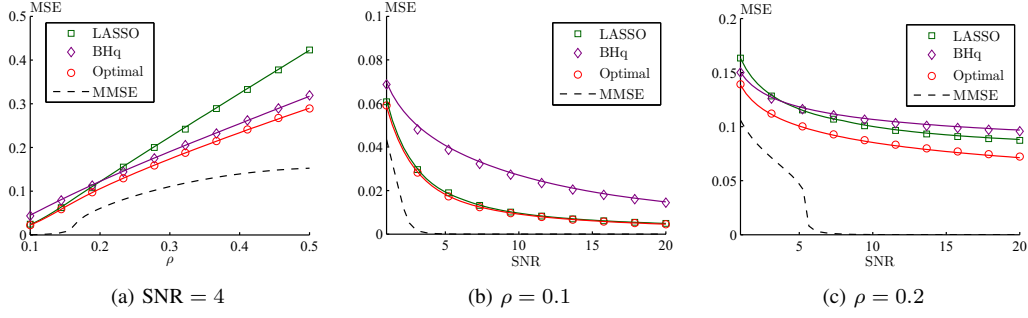


Figure 4: Comparison of MSEs obtained by three regularization sequences: LASSO, BHq and the oracle optimal design. Here, β_i are i.i.d. Bernoulli random variables, with $\mathbb{P}(\beta_i = 1) = \rho$ and $w_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \sigma_w^2)$, with $\sigma_w^2 = \rho/\text{SNR}$ and the BHq sequences are generated by reordering i.i.d. samples of: $\sigma_w \Phi^{-1}(1 - \frac{q}{2} + \frac{q}{2}U)$, where $q \in (0, 1]$ and U follows the uniform distribution over $[0, 1]$. In our simulation, we fix $p = 2048$, $\delta = 0.5$ and the empirical results are averaged over 20 independent trials. The dash curves correspond to the information-theoretic limit obtained in [16], [17].

In Fig. 4, we compare the MSEs achieved by different regularizing sequences (LASSO, BHq and oracle optimal design), at different SNR and sparsity levels. Since we are concerned with oracle optimality, for fair comparison, we search through the parameters of the BHq and LASSO sequences (in particular, q for BHq and λ for LASSO) and report the minimum MSEs that can be achieved. The solid curves correspond to the theoretical MSEs predicted by Theorem 1 and Proposition 4. Note that the empirical MSEs match well with theoretical predictions³. It is also observed that under each setting, the MSEs of different regularizing sequences are all above the lower bound obtained in (30) (red curve in the figure). Also we can see this lower bound can be attained when the limiting empirical distribution of λ follows prescribed optimal distribution (31). We also have the following findings:

- 1) As can be seen from Fig. 4a and Fig. 4b, when ρ is small, LASSO performs well and the corresponding MSEs almost match the theoretical lower bound, across different values of SNR. However, its performance degrades faster than the other two sequences, as ρ grows. This is because LASSO's penalization is not adaptive to the underlying sparsity levels and it incurs higher bias under larger ρ [7].
- 2) From Fig. 4b and Fig. 4c, we can find that at low SNR regimes, the BHq sequence can lead to comparable performance as the optimal design. However, at higher SNR regimes, the optimal design notably outperforms the BHq sequence. To explain this phenomenon, we plot in Fig. 5 the empirical distributions of the λ -sequences associated with the optimal design and the BHq design, respectively. It turns out that, in the low SNR case, the optimal design and BHq have similar distributions, while at higher SNRs, the distribution of the optimal design is close to a mixture of a delta mass and uniform distribution.

³Here, the MSEs of LASSO and BHq are obtained by optimizing over the parameters λ and q , so strictly speaking, the theoretical curves are valid only if a stronger uniform convergence result holds. The uniform convergence for LASSO case is proved in [50], [51] and we conjecture that it also holds true for BHq sequences.

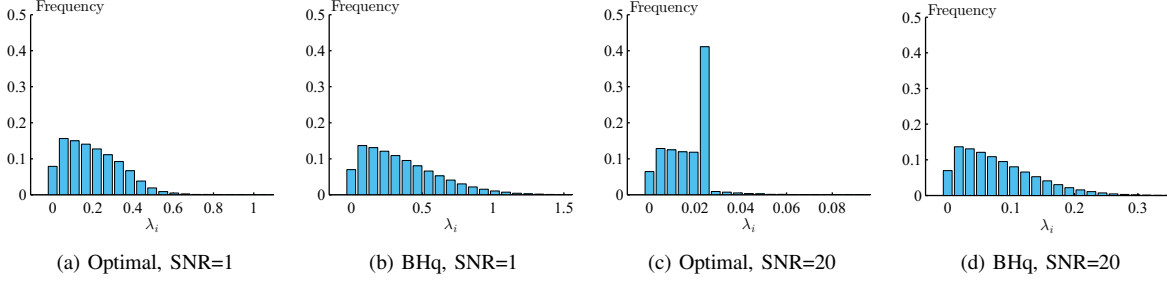


Figure 5: Comparison of empirical distributions of two regularizing sequences “BHq” and “Optimal” in Fig. 4c.

B. Variable Selection with Maximum Power

Next we consider using SLOPE for variable selection. Our goal is to find the optimal regularizing sequence to achieve the highest possible power, under a given level of type-I error α , formulated as:

$$\begin{aligned} \mathcal{P}(\alpha) &\stackrel{\text{def}}{=} \sup_{\Lambda \in \tilde{\mathcal{P}}_\Lambda} \lim_{p \rightarrow \infty} \text{Power} \\ \text{s.t. } &\lim_{p \rightarrow \infty} \text{Type-I error} \leq \alpha, \end{aligned} \quad (39)$$

where $\tilde{\mathcal{P}}_\Lambda \stackrel{\text{def}}{=} \{\Lambda \in \mathcal{P}_\Lambda : (\text{R.1}) \text{ is satisfied}\}$ is the admissible set of μ_Λ , with which the limits in (39) exist. In light of (21) and (22), if $\mathbb{P}(B = 0) \in (0, 1)$, optimization problem (39) is equivalent to:

$$\begin{aligned} \mathcal{P}(\alpha) &= \sup_{\Lambda \in \tilde{\mathcal{P}}_\Lambda} \mathbb{P}(|B + \sigma_* H| \geq y_{\text{th}}^* \mid B \neq 0) \\ \text{s.t. } &\mathbb{P}(|\sigma_* H| \geq y_{\text{th}}^*) \leq \alpha, \end{aligned} \quad (40)$$

where $y_{\text{th}}^* = \sup_{y \geq 0} \{y \mid \eta(y; \mu_{Y_*}, \mu_{\tau_* \Lambda}) = 0\}$. Comparing the admissible set $\tilde{\mathcal{P}}_\Lambda$ with $\mathcal{P}_\Lambda$ in (24), it can be seen the only difference is that here we need an additional condition (R.1) to ensure the limits of Type-I error and Power both exist (see Proposition 2).

To state our results, we first introduce the following function, which is the counterpart of (28).

$$\begin{aligned} \mathcal{L}_\alpha(\sigma) &\stackrel{\text{def}}{=} \inf_{f \in \mathcal{I} \cap \mathcal{F}_{\alpha, \sigma}} \mathbb{E}[f(B + \sigma H) - B]^2 \\ \text{s.t. } &\delta^{-1} \mathbb{E}[f'(B + \sigma H)] \leq 1 \end{aligned} \quad (41)$$

where $\alpha \in [0, 1]$ is a prescribed Type-I error level and $\mathcal{F}_{\alpha, \sigma} \stackrel{\text{def}}{=} \{f(y) : f(y) = 0 \text{ for } |y| \leq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma\}$. Similar as Proposition 4, we will see that the maximum power achievable by SLOPE under Type-I error level α is related with the following equation:

$$\mathcal{L}_\alpha(\sigma) = \delta(\sigma^2 - \sigma_w^2), \quad (42)$$

where $\mathcal{L}_\alpha(\sigma)$ is the function defined in (41).

We are now ready to state our main optimality results for variable selection.

Proposition 5: Under the same setting as Proposition 2, assume $\mathbb{P}(B = 0) \in (0, 1)$. Then we have

- (a) For any $\alpha \in [0, 1]$ and $\sigma > 0$, problem (41) is convex and there exists a unique optimal solution $f_{\alpha, \sigma} \in \mathcal{I}$.

- (b) For any $\alpha \in [0, 1]$, $\mathcal{L}_\alpha(\sigma)$ is continuous on $\mathbb{R}_{>0}$ and equation (42) always has a solution. The minimum solution $\sigma_{0,\alpha} \in [\sigma_w, \sqrt{\sigma_w^2 + \delta^{-1}\mathbb{E}B^2}]$.
- (c) Let $Y_{0,\alpha} := B + \sigma_{0,\alpha}H$ and $f_\alpha := f_{\alpha,\sigma_{0,\alpha}}$. If $\lim_{p \rightarrow \infty} \text{Type-I error} \leq \alpha$, then

$$\lim_{p \rightarrow \infty} \text{Power} \leq \mathbb{P}(|Y_{0,\alpha}| \geq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{0,\alpha} \mid B \neq 0). \quad (43)$$

Moreover, if $\delta^{-1}\mathbb{E}[f'_\alpha(Y_{0,\alpha})] < 1$ and $y_{0,\alpha} = \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{0,\alpha}$, the upper bound in (43) can be attained by $\mu_\Lambda = \mu_{\text{opt},\alpha}$, with $\mu_{\text{opt},\alpha}$ being the law of

$$\frac{1}{\tau_{0,\alpha}} \max \{y_{0,\alpha}, |Y_{0,\alpha}| - f_\alpha(|Y_{0,\alpha}|)\}. \quad (44)$$

Here, $y_{0,\alpha} = \sup_{y \geq 0} \{y \mid f_\alpha(y) = 0\}$ and $\tau_{0,\alpha} = [1 - \delta^{-1}\mathbb{E}f'_\alpha(Y_{0,\alpha})]^{-1}$.

The proof of Proposition 5, which is similar to that of Proposition 4, will be given in Sec. V-E. A key step is to show the realizable set of η in the variable selection setting is still equal to $\mathcal{I}$ (see Lemma 20 in Appendix N), although the admissible set of μ_Λ is replaced by $\tilde{\mathcal{P}}_\Lambda$, which is a subset of $\mathcal{P}_\Lambda$ in the estimation setting.

Remark 8: Comparing the results in Proposition 4 and Proposition 5, we can find that although at the beginning, we are dealing with two different problems (the objective of the first one is minimizing the MSE, while the other is on maximizing the power under a given Type-I error), we end up with two procedures of very similar natures. Both problems can finally be converted into a formulation involving finding the optimal estimation of β that can be achieved by SLOPE under the observation $\mathbf{y} = \beta + \sigma\mathbf{h}$, with $\mathbf{h} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. The only difference is that in the second problem, we need to enforce an additional restriction on the regularization sequence λ to ensure the Type-I error is below certain threshold α .

Remark 9 (Tightness of upper bound (43)): The tightness of the upper bound for power relies on the conditions: $\delta^{-1}\mathbb{E}[f'_\alpha(Y_{0,\alpha})] < 1$ and $y_{0,\alpha} = \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{0,\alpha}$. Numerically they hold under all the settings considered. We conjecture that within our assumptions, this condition always hold and the upper bound (43) is tight.

In Fig.6, we compare the variable selection performance achieved by the optimal regularization with that of LASSO and BHq sequences. We show both theoretical ROC curves and the empirical power under given Type-I error levels. Here each empirical (Type-I error, power) pair is generated by first fixing all the parameters (including the tuning parameters such as λ and q) and then averaging over 20 independent trials. It can be seen that the empirical results match well with the theoretical predictions (solid curves in the figures) and the optimal design of regularization dominates the other two regularizing sequences. We also have the following observations:

- 1) In all cases, the theoretical upper bounds on power (43) can be achieved by choosing μ_Λ to be the law of (44).
- 2) The performance of LASSO is closed to the fundamental limit at low sparsity and high SNR regimes, while its performance is significantly degraded as sparsity grows higher or SNR grows lower. In particular, we can find in such cases, the maximum power achievable by LASSO is less than 1. This phenomenon is also discussed in [2], [7], [11] and it is inherently connected with the so-called “noise-sensitivity” phase transition [52]. In comparison, the optimal and BHq sequences can both reach power 1, after Type-I errors are above certain thresholds.

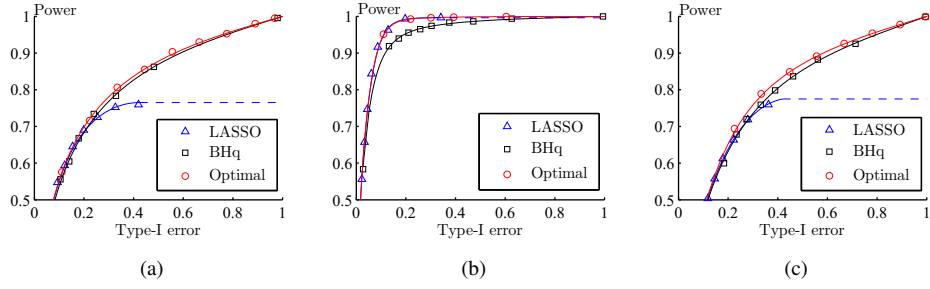


Figure 6: Theoretical predictions v.s. empirical results of testing performance using LASSO, BHq and oracle optimal sequences. Here, β_i are i.i.d. Bernoulli random variables, with $\mathbb{P}(\beta_i = 1) = \rho$ and $w_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \sigma_w^2)$, with $\sigma_w^2 = \rho/\text{SNR}$. The empirical results are generated under $p = 2048$ and $\delta = 0.5$ and we choose different ρ and SNR values: (a) $\rho = 0.1$, SNR = 0.6, (b) $\rho = 0.1$, SNR = 4, (c) $\rho = 0.2$, SNR = 4. Dash curves correspond to the observed upper bound of power achieved by LASSO.

- 3) Complementary to LASSO, the performance of BHq sequences is closed to the theoretical upper bounds at low SNRs or large sparsity levels, while it deviates from the upper bounds in other scenarios.

V. PROOF OF MAIN RESULTS

A. Asymptotic Separability

In this section, we are going to prove Proposition 1.

From (5) we have the following scaling property: $\mathcal{M}_\lambda(\mathbf{y}; \tau) = \frac{1}{\tau} \mathcal{M}_{\tau\lambda}(\mathbf{y}; 1)$. On the other hand, for any $\tau > 0$, if λ is a converging sequence with limiting measure μ_λ , it is not hard to show $\tau\lambda$ is also a converging sequence, with limiting measure $\mu_{\tau\lambda}$. Thus, to study the asymptotic limit of (5) under $(\mathbf{y}, \lambda, \tau)$, it suffices to consider $(\mathbf{y}, \tau\lambda, 1)$. As a result, without loss of generality, we will assume $\tau = 1$ in the rest of our proof.

1) *Some preliminary facts about SLOPE*: The asymptotic separability stems from the following unique properties of the SLOPE proximal minimization problem (5), which are proved in [9, Sec. 2].

Fact 1: For any $\lambda, \mathbf{y} \in \mathbb{R}^p$, with $\lambda_i \geq 0$, for all $i \in [p]$, it holds that

- (i) (*Sign consistency*) For any $i \in [p]$, $[\text{Prox}_\lambda(\mathbf{y})]_i$ has the same sign as y_i . Moreover,

$$[\text{Prox}_\lambda(\mathbf{y})]_i = \text{sign}(y_i)[\text{Prox}_\lambda(|\mathbf{y}|)]_i.$$

- (ii) (*Permutation-invariance*) For any permutation matrix Π , $\Pi \text{Prox}_\lambda(\mathbf{y}) = \text{Prox}_\lambda(\Pi \mathbf{y})$.

- (iii) (*Monotonicity and Lipschitz continuity*) For any $i, j \in [p]$, if $y_i \leq y_j$, then $0 \leq [\text{Prox}_\lambda(\mathbf{y})]_j - [\text{Prox}_\lambda(\mathbf{y})]_i \leq y_j - y_i$ and for any y_i , $[\text{Prox}_\lambda(\mathbf{y})]_i \leq |y_i|$.

An immediate yet important implication of Fact 1 is the following lemma:

Lemma 1: For any $\lambda, \mathbf{y} \in \mathbb{R}^p$, with $\lambda_i \geq 0$, there always exists an odd, non-decreasing and 1-Lipschitz function g_p such that for all $i \in [p]$, $g_p(y_i) = [\text{Prox}_\lambda(\mathbf{y})]_i$.

The proof of Lemma 1 is given in Appendix A. By Lemma 1 we know $\text{Prox}_\lambda(\mathbf{y})$ is actually the restriction of a function $g_p \in \mathcal{I}$ onto the support of $\mu_{\mathbf{y}}$. Moreover, from the permutation invariance property (Fact 1 (ii)), such g_p

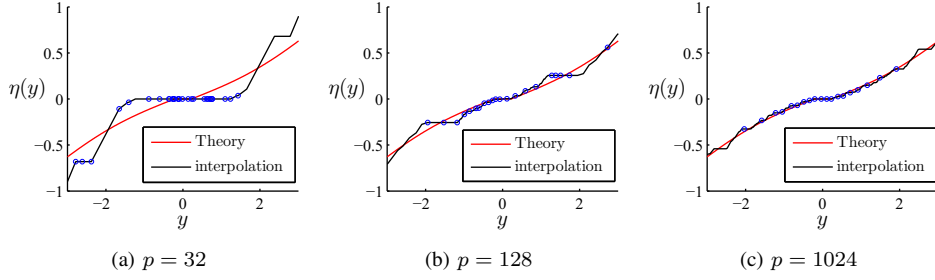


Figure 7: Comparison between $\eta(y)$ (red curve), linear interpolation (black curve) and $\{(y_i, [\text{Prox}_\lambda(\mathbf{y})]_i)\}_{i \in [p]}$ (blue dots), under three different values of p . Here $y_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, 2)$ and λ_i are i.i.d. samples from BHq distribution [9].

is only determined by the empirical measure μ_λ and μ_y . We could expect if μ_λ and μ_y both converge to some limiting distributions, g_p will also converge to certain limiting scalar function. This is exactly the essential meaning of asymptotic separability.

Before proceeding, let us take a look at a numerical justification shown in Fig. 7. Here we choose g_p to be the linear interpolation of the following set of points: $\{(|y_i|, |\hat{y}_i|), (-|y_i|, -|\hat{y}_i|)\}_{i=1}^p \cup (0, 0)$, where $\hat{\mathbf{y}} := \text{Prox}_\lambda(\mathbf{y})$. It is easy to check (as is shown in the proof of Lemma 1) such linear interpolation is a qualified candidate for g_p in Lemma 1. We compare it with the limiting scalar function $\eta(y)$ predicted by Proposition 1. It is clear that as p becomes larger, $g_p(y)$ gets increasingly close to $\eta(y)$.

2) *An equivalent form of (5)*: Based on Lemma 1, we then go on to show the equivalence between (5) and the following problem:

$$\min_{g \in \mathcal{L}} \underbrace{\frac{1}{2} \mathbb{E}_{\mu_y} [Y - g(Y)]^2 + \int_0^1 F_\lambda^{-1}(u) F_{|g(Y)|}^{-1}(u) du}_{\stackrel{\text{def}}{=} L_p(g)}. \quad (45)$$

This is formalized in the following lemma, whose proof is given in Appendix B.

Lemma 2: Denote $\mathcal{M}_\lambda^*(\mathbf{y})$ as the optimal value of (45). Then it holds that $\frac{\mathcal{M}_\lambda(\mathbf{y}; 1)}{p} = \mathcal{M}_\lambda^*(\mathbf{y})$. Besides, any optimal solution $g_p^*(y)$ of (45) satisfies: $g_p^*(\mathbf{y}) = \text{Prox}_\lambda(\mathbf{y})$.

Comparing (6) and (45), it could be now understood that the optimization in (6) is the limit of (45), as $\mu_\lambda \rightarrow \mu_\Lambda$ and $\mu_y \rightarrow \mu_Y$. Therefore, from Lemma 2, we could expect $\frac{1}{p} \mathcal{M}_\lambda(\mathbf{y}; 1) = \mathcal{M}_\lambda^*(\mathbf{y}) \rightarrow \mathcal{M}_{\mu_\Lambda}(\mu_Y, 1)$. On the other hand, $g_p^*(\cdot)$, which is the optimal solution of (45) should also converge to the optimal solution of (6): $\eta(\cdot; \mu_Y, \mu_\Lambda)$. Thus for any $\lambda, \mathbf{y}$ satisfying $\mu_\lambda \approx \mu_\Lambda$ and $\mu_y \approx \mu_Y$, we would have $\text{Prox}_\lambda(\mathbf{y}) = g_p^*(\mathbf{y}) \approx \eta(\mathbf{y}; \mu_Y, \mu_\Lambda)$, i.e., asymptotic separability holds. The final step of the proof is to make the above intuition accurate and rigorous.

3) *Taking the limit of (45)*: Recall that we have assumed $\tau = 1$. For notational simplicity, denote $\mathcal{M}_\lambda(\mathbf{y}) := \mathcal{M}_\lambda(\mathbf{y}; 1)$ and $\mathcal{M}_{\mu_\Lambda}(\mu_Y) := \mathcal{M}_{\mu_\Lambda}(\mu_Y, 1)$. Define $L(g)$ as the objective function of (6), i.e.,

$$L(g) \stackrel{\text{def}}{=} \frac{1}{2} \mathbb{E}_{\mu_Y} [Y - g(Y)]^2 + \int_0^1 F_\Lambda^{-1}(u) F_{|g(Y)|}^{-1}(u) du. \quad (46)$$

and g^* as the corresponding optimal solution. By Lemma 7 in Appendix C, we have $\sup_{g \in \mathcal{I}} |L(g) - L_p(g)| \rightarrow 0$, where $L_p(g)$ is the objective function of (45). Therefore, (8) immediately follows, since

$$\begin{aligned} |\frac{1}{p}\mathcal{M}_\lambda(\mathbf{y}; 1) - \mathcal{M}_{\mu_\Lambda}(\mu_Y, 1)| &= |\sup_{g \in \mathcal{I}} L(g) - \sup_{g \in \mathcal{I}} L_p(g)| \\ &\leq \sup_{g \in \mathcal{I}} |L(g) - L_p(g)|. \end{aligned}$$

On the other hand,

$$\begin{aligned} L_p(g^*) - L_p(g_p^*) &= L_p(g^*) - L(g^*) + L(g_p^*) - L_p(g_p^*) \\ &\quad + L(g^*) - L(g_p^*) \\ &\leq |L_p(g^*) - L(g^*)| + |L(g_p^*) - L_p(g_p^*)|, \end{aligned} \tag{47}$$

where in the last step we use the optimality of g^* . By the strong convexity of (45), we have

$$L_p(g^*) - L_p(g_p^*) \geq \frac{1}{2} \mathbb{E}_{\mu_y} |g_p^*(Y) - g^*(Y)|^2. \tag{48}$$

Combining (47) and (48) gives

$$\mathbb{E}_{\mu_y} |g_p^*(Y) - g^*(Y)|^2 \leq 2 \sup_{g \in \mathcal{I}} |L(g) - L_p(g)|.$$

By Lemma 7 again, we have $\mathbb{E}_{\mu_y} |g_p^*(Y) - g^*(Y)|^2 \rightarrow 0$, as $p \rightarrow \infty$. This is exactly (9), since $g_p^*(\mathbf{y}) = \text{Prox}_\lambda(\mathbf{y})$ by Lemma 2 and $g^*(y) = \eta(y; \mu_Y, \mu_{\tau\Lambda})$. Finally, the uniqueness of $g^*(\cdot)$ (up to a set of measure 0 with respect to μ_Y) is proved in Lemma 8. This completes our proof.

B. Asymptotic Estimation Performance

1) *Convex Gaussian Min-max Theorem*: Our proof hinges on the Convex Gaussian Min-max Theorem (CGMT). For completeness, we briefly summarize the key idea here. The CGMT studies a minimax optimization problem (PO) of the form:

$$\Phi(\mathbf{G}) = \min_{\mathbf{v} \in \mathcal{S}_v} \max_{\mathbf{u} \in \mathcal{S}_u} \mathbf{u}^\top \mathbf{G} \mathbf{v} + \psi(\mathbf{v}, \mathbf{u}), \tag{49}$$

where $\mathcal{S}_v \subset \mathbb{R}^p$, $\mathcal{S}_u \subset \mathbb{R}^n$ are two compact sets, $\psi : \mathcal{S}_v \times \mathcal{S}_u \rightarrow \mathbb{R}$ is a continuous convex-concave function w.r.t. $(\mathbf{v}, \mathbf{u})$ and $G_{ij} \stackrel{i.i.d.}{\sim} \mathcal{N}(0, 1)$. Problem (49) can be associated with the following auxiliary optimization (AO) problem:

$$\phi(\mathbf{g}, \mathbf{h}) = \min_{\mathbf{v} \in \mathcal{S}_v} \max_{\mathbf{u} \in \mathcal{S}_u} \|\mathbf{v}\|_2 \mathbf{g}^\top \mathbf{u} + \|\mathbf{u}\|_2 \mathbf{h}^\top \mathbf{v} + \psi(\mathbf{v}, \mathbf{u}), \tag{50}$$

where $\mathbf{g} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$ and $\mathbf{h} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$. Roughly speaking, CGMT shows that $\frac{1}{p} \Phi(\mathbf{G}) \approx \frac{1}{p} \phi(\mathbf{g}, \mathbf{h})$ and the optimal solutions of (49) and (50) have approximately the same empirical distributions in the large p limit. Usually, (AO) is easier to analyze, so it provides a convenient handle for analyzing (PO). For a detailed descriptions, readers can refer to [26, Theorem 3].

2) *Proof of Theorem 1:* The first step is to recast (2) into the minimax form as in (49). Letting $\mathbf{v} = \mathbf{x} - \boldsymbol{\beta}$, (2) can be equivalently written as:

$$\begin{aligned} & \min_{\mathbf{v}} \underbrace{\frac{1}{n} \left[\frac{1}{2} \|\mathbf{A}\mathbf{v} - \mathbf{w}\|^2 + J_{\lambda}(\mathbf{v} + \boldsymbol{\beta}) \right]}_{:=C(\mathbf{v})} \\ &= \min_{\mathbf{v}} \max_{\mathbf{u}} \frac{1}{n} \left[\frac{\mathbf{u}^\top}{\sqrt{n}} (\sqrt{n}\mathbf{A}) \mathbf{v} - \mathbf{u}^\top \mathbf{w} - \frac{\|\mathbf{u}\|^2}{2} + J_{\lambda}(\mathbf{v} + \boldsymbol{\beta}) \right]. \end{aligned} \quad (51)$$

Denote $\hat{\mathbf{v}} \stackrel{\text{def}}{=} \underset{\mathbf{v}}{\operatorname{argmin}} C(\mathbf{v})$ and correspondingly, $\hat{\boldsymbol{\beta}} = \hat{\mathbf{v}} + \boldsymbol{\beta}$. Now (51) has the same form as (49) and the corresponding (AO) is:

$$\begin{aligned} & \min_{\mathbf{v}} \max_{\mathbf{u}} \frac{1}{n} \left[-\frac{\|\mathbf{u}\|}{\sqrt{n}} \mathbf{h}^\top \mathbf{v} - \frac{\|\mathbf{v}\|}{\sqrt{n}} \mathbf{g}^\top \mathbf{u} - \mathbf{u}^\top \mathbf{w} - \frac{\|\mathbf{u}\|^2}{2} + J_{\lambda}(\mathbf{v} + \boldsymbol{\beta}) \right] \\ &= \min_{\mathbf{v}} \max_{\theta \geq 0} \theta \left(\left\| \frac{\|\mathbf{v}\|}{n} \mathbf{g} + \frac{\mathbf{w}}{\sqrt{n}} \right\| - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right) - \frac{1}{2} \theta^2 + \frac{J_{\lambda}(\mathbf{v} + \boldsymbol{\beta})}{n} \\ &= \min_{\mathbf{v}} \underbrace{\frac{1}{2} \left(\sqrt{\frac{\|\mathbf{v}\|^2}{n} \frac{\|\mathbf{g}\|^2}{n} + \frac{\|\mathbf{w}\|^2}{n} + 2 \frac{\|\mathbf{v}\|}{\sqrt{n}} \frac{\mathbf{g}^\top \mathbf{w}}{n} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right)^2 + \frac{J_{\lambda}(\mathbf{v} + \boldsymbol{\beta})}{n}}_{:=L(\mathbf{v})}, \end{aligned} \quad (52)$$

where $\mathbf{g} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$ and $\mathbf{h} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$. Let $D \subset \mathbb{R}^p$ be any closed set. Then by CGMT we can show for any $t \in \mathbb{R}$,

$$\mathbb{P}(\min_{\mathbf{v} \in D} C(\mathbf{v}) \leq t) \leq 2\mathbb{P}(\min_{\mathbf{v} \in D} L(\mathbf{v}) \leq t) \quad (53)$$

and if D is also convex,

$$\mathbb{P}(\min_{\mathbf{v} \in D} C(\mathbf{v}) \geq t) \leq 2\mathbb{P}(\min_{\mathbf{v} \in D} L(\mathbf{v}) \geq t). \quad (54)$$

The proof of (53) and (54) is the same as [50, Corollary 5.1] and is omitted here. We are going to apply (53) and (54) to prove (12). We will follow the proof steps in [50].

First define the following minimax problem:

$$\Psi_* \stackrel{\text{def}}{=} \min_{\sigma \geq \sigma_w} \max_{\theta \geq 0} \underbrace{\frac{\theta}{2} \left(\frac{\sigma_w^2}{\sigma} + \sigma \right) - \frac{\theta^2}{2} + \frac{1}{\delta} \left[\mathcal{F}(\sigma, \theta) - \frac{\theta\sigma}{2} \right]}_{:=\Psi(\sigma, \theta)}, \quad (55)$$

where $\mathcal{F}(\sigma, \theta) \stackrel{\text{def}}{=} \frac{\theta}{2\sigma} \mathbb{E}[Y - \eta(Y)]^2 + \int_0^1 F_{\Lambda}^{-1}(u) F_{|\eta(Y)|}^{-1}(u) du$, with $\eta(\cdot) = \eta(\cdot; \mu_Y, \mu_{\sigma\Lambda/\theta})$. To prove (12), we adopt the same perturbation argument as in [26], [29], [50]. In particular, for a pseudo-Lipschitz function $\psi(\cdot, \cdot)$, define the following set of $\mathbf{v}$:

$$D_{\nu} := \{\mathbf{v} \in \mathbb{R}^p : |\mathbb{E}_{\mu_{\mathbf{v}+\boldsymbol{\beta}, \boldsymbol{\beta}}} \psi - \mathbb{E}_{\mu^*} \psi| \geq \nu\}, \quad (56)$$

where $\nu > 0$ and μ^* denotes the joint measure of $(\eta(Y_*; \mu_{Y_*}, \mu_{\sigma_*\Lambda/\theta_*}), B)$. Here (σ_*, θ_*) is the optimal solution of (55) and $Y_* = B + \sigma_* H$, with $H \sim \mathcal{N}(0, 1)$ independent of $B \sim \mu_B$. Recall that $\hat{\boldsymbol{\beta}} = \hat{\mathbf{v}} + \boldsymbol{\beta}$, so for any $\nu > 0$ and $\varepsilon > 0$

$$\begin{aligned} & \mathbb{P}\left(\left|\frac{1}{p} \sum_{i=1}^p \psi(\hat{\beta}_i, \beta_i) - \mathbb{E}[\psi(\eta(Y_*; \mu_{Y_*}, \mu_{\sigma_*\Lambda/\theta_*}), B)]\right| \geq \nu\right) = \mathbb{P}(\hat{\mathbf{v}} \in D_{\nu}) \\ & \leq \mathbb{P}\left(\min_{\mathbf{v} \in D_{\nu}} C(\mathbf{v}) \leq \min_{\mathbf{v}} C(\mathbf{v}) + \varepsilon\right). \end{aligned} \quad (57)$$

This indicates that if we can show for any $\nu > 0$ and some $\varepsilon > 0$, $\min_{\mathbf{v} \in D_\nu} C(\mathbf{v}) \leq \min_{\mathbf{v}} C(\mathbf{v}) + \varepsilon$ occurs with vanishing probability, then (12) will immediately follow (with $\tau_* = \sigma_*/\theta_*$). In this way, proving (12) is reformulated as the perturbation analysis of $C_\lambda(\mathbf{v})$, which can be done as follows. For any $\varepsilon \geq 0$ and $K > 0$, we have

$$\begin{aligned}
& \mathbb{P}\left(\min_{\mathbf{v} \in D_\nu} C(\mathbf{v}) \leq \min_{\mathbf{v}} C(\mathbf{v}) + \varepsilon\right) \\
& \leq \mathbb{P}\left(\min_{\mathbf{v} \in D_\nu} C(\mathbf{v}) \leq \Psi_* + 2\varepsilon\right) + \mathbb{P}\left(\min_{\mathbf{v}} C(\mathbf{v}) > \Psi_* + \varepsilon\right) \\
& \leq \mathbb{P}\left(\min_{\mathbf{v} \in D_\nu \cap \mathcal{B}_{\sqrt{n}K}} C(\mathbf{v}) \leq \Psi_* + 2\varepsilon\right) + \mathbb{P}\left(\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} C(\mathbf{v}) > \Psi_* + \varepsilon\right) + 2\mathbb{P}(\hat{\mathbf{v}} \notin \mathcal{B}_{\sqrt{n}K}) \\
& \stackrel{(a)}{\leq} 2\mathbb{P}\left(\min_{\mathbf{v} \in D_\nu \cap \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) \leq \Psi_* + 2\varepsilon\right) + 2\mathbb{P}\left(\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) > \Psi_* + \varepsilon\right) + 2\mathbb{P}(\hat{\mathbf{v}} \notin \mathcal{B}_{\sqrt{n}K}), \tag{58}
\end{aligned}$$

where (a) is due to (53) and (54). Here Ψ_* is the optimal value of (55). In Appendix D, we show all the three probabilities on the RHS of (58) vanish for $K \geq \frac{\sigma_*}{\sqrt{\delta}} + \frac{\theta_{\min}}{4}$ (with $\theta_{\min}$ given in Lemma 14) :

- (i) From (88) of Lemma 10, $\mathbb{P}(|\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) - \Psi_*| \geq \varepsilon) \rightarrow 0$ for any $\varepsilon > 0$.
- (ii) From (112) of Lemma 12, $\mathbb{P}(\hat{\mathbf{v}} \notin \mathcal{B}_{\sqrt{n}K}) \rightarrow 0$.
- (iii) From Lemma 11, for any $\nu > 0$ there exists $\varepsilon_0 > 0$ such that for any $\varepsilon \leq \varepsilon_0$, $\mathbb{P}(\min_{\mathbf{v} \in D_\nu \cap \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) \leq \Psi_* + 2\varepsilon) \rightarrow 0$.

After substituting (i)-(iii) back to (58), we deduce that for any $\nu > 0$, there always exist $\varepsilon_0 > 0$ such that for any $\varepsilon \leq \varepsilon_0$, the RHS of (57) converges to 0 as $p \rightarrow \infty$. Therefore,

$$\frac{1}{p} \sum_{i=1}^p \psi(\hat{\beta}_i, \beta_i) \xrightarrow{\mathbb{P}} \mathbb{E}[\psi(\eta(Y_*; \mu_{Y_*}, \mu_{\sigma_* \Lambda / \theta_*}), B)],$$

On the other hand, by Lemma 14 in Appendix J, (σ_*, θ_*) is the unique solution of the following fixed point equation of (σ, θ) :

$$\sigma^2 = \sigma_w^2 + \frac{1}{\delta} \mathbb{E}[\eta(Y; \mu_Y, \mu_{\sigma \Lambda / \theta}) - B]^2 \tag{59}$$

$$\theta = \sigma \left[1 - \frac{1}{\delta} \mathbb{E} \eta'(Y; \mu_Y, \mu_{\sigma \Lambda / \theta}) \right]. \tag{60}$$

Therefore, letting $\tau_* = \sigma_*/\theta_*$, we can see (σ_*, τ_*) is also a solution of (13)-(14). Finally, we show such (σ_*, τ_*) is the unique solution of (13)-(14). By Lemma 14, (σ_*, θ_*) is the unique solution of (59)-(60) and it satisfies $\sigma_* \geq \sigma_w$, $\theta_* \geq \theta_{\min} > 0$. Suppose there exist two different solutions (σ_1, τ_1) and (σ_2, τ_2) to (13)-(14), then since $\sigma_1, \tau_1, \sigma_2, \tau_2 > 0$, $(\sigma_1, \sigma_1/\tau_1)$ and $(\sigma_2, \sigma_2/\tau_2)$ are two different solutions to (59)-(60), leading to a contradiction. This concludes our proof.

C. Asymptotic Variable Selection Performance

In this section, our goal is to prove Proposition 2. We first prove the convergence of $R_0^{(p)}$.

1) *Probabilistic upper bound of $R_0^{(p)}$* : To prove the convergence of $R_0^{(p)}$, the first step is to establish the following probabilistic upper bound.

Lemma 3: For any $\varepsilon > 0$, $R_0^{(p)} \leq \mathbb{P}(\eta(Y_*) = 0) + \varepsilon$, with probability approaching 1, as $p \rightarrow \infty$.

The proof of Lemma 3 is given in Appendix E, which uses a standard approximation argument (see *e.g.*, the proof of Lemma 2.2 (iii) and (v) in [53]).

2) *Probabilistic lower bound of $R_0^{(p)}$* : The second step is to prove the following matching probabilistic lower bound for $R_0^{(p)}$:

Lemma 4: Under the same setting as Proposition 2, for any $\varepsilon > 0$,

$$R_0^{(p)} \geq \mathbb{P}(\eta(Y_*) = 0) - \varepsilon, \quad (61)$$

with probability approaching 1, as $p \rightarrow \infty$.

The proof of Lemma 4, which can be found in Appendix F, is the mostly technically involved part, so we provide more detailed explanations here.

A key strategy we adopt is using the vector

$$\hat{\mathbf{s}} \stackrel{\text{def}}{=} \mathbf{A}^\top (\mathbf{y} - \mathbf{A}\hat{\boldsymbol{\beta}}) \quad (62)$$

as an indicator of zero coordinates of $\hat{\boldsymbol{\beta}}$. To give the formal statements, we need to first introduce the notion of *majorization*.

Definition 2: For two vectors $\mathbf{a}, \mathbf{b} \in \mathbb{R}^p$, we say $\mathbf{a}$ is majorized by $\mathbf{b}$ (denoted as $\mathbf{a} \prec \mathbf{b}$), if for any $j \in [p]$, $\sum_{i=j}^p |a|_{(j)} \leq \sum_{i=j}^p |b|_{(j)}$. On the other hand, we say $\mathbf{a}$ is *strictly* majorized by $\mathbf{b}$ (denoted as $\mathbf{a} \prec_S \mathbf{b}$), if for any $j \in [p]$, $\sum_{i=j}^p |a|_{(j)} < \sum_{i=j}^p |b|_{(j)}$.

Denote $|\mathbf{x}|_{(1:k)}$ as a vector formed by the largest k components of $|\mathbf{x}|$. Let us call $|\mathbf{x}|_{(1:k)}$ as the *k-dominant subvector* of $\mathbf{x}$. The following is the key lemma for establishing the probabilistic lower bound of $R_0^{(p)}$.

Lemma 5: For the optimization problem (2), suppose for some $k \in [p]$, $|\hat{\mathbf{s}}|_{(1:k)} \prec_S \boldsymbol{\lambda}_{1:k}$, where $\hat{\mathbf{s}} = \mathbf{A}^\top (\mathbf{y} - \mathbf{A}\hat{\boldsymbol{\beta}})$. Then we have $|\hat{\boldsymbol{\beta}}|_{(1:k)} = \mathbf{0}_k$.

The proof of Lemma 5 can be found in Appendix G. This characterization transform the original problem of searching zero coordinates of $\hat{\boldsymbol{\beta}}$ into a new problem of discovering whether there is a strict majorization relation between k -dominant subvectors of $\hat{\mathbf{s}}$ and $\boldsymbol{\lambda}$. The nice thing making this strategy work is that the majorization relation between two vectors is fully captured by their empirical distributions. Besides, in our setting, the empirical distributions $\mu_{\boldsymbol{\lambda}}$ and $\mu_{\hat{\mathbf{s}}}$ both have simple limits: by our assumption, $\mu_{\boldsymbol{\lambda}} \rightarrow \mu_{\Lambda}$ and in Proposition 6 of Appendix I, we show $\mu_{\hat{\mathbf{s}}} \rightarrow \mu_{\hat{\mathbf{s}}}$, with $\mu_{\hat{\mathbf{s}}}$ being the law of $\frac{Y_* - \eta(Y_*)}{\tau_*}$.

A major part of proof of Lemma 4 is to show if condition (R.1) is satisfied and $\mathbb{P}(\eta(Y_*) = 0) > 0$, then for $k = \lfloor p[\mathbb{P}(\eta(Y_*) = 0) - \varepsilon] \rfloor$, where $\varepsilon > 0$ can be arbitrarily small, we have $|\hat{\mathbf{s}}|_{(1:k)} \prec_S \boldsymbol{\lambda}_{(1:k)}$ with probability approaching 1, as $p \rightarrow \infty$. Then an application of Lemma 5 will give us the desired probabilistic lower bound for $R_0^{(p)}$ shown in Lemma 4.

Remark 10: Let us briefly explain why $\mathbf{s}$ in (62) is related with the zero coordinates of $\hat{\boldsymbol{\beta}}$. By the first order condition of (2), we can get $\hat{\mathbf{s}} \in \partial J_{\boldsymbol{\lambda}}(\hat{\boldsymbol{\beta}})$, i.e., $\mathbf{s}$ defined in (62) is a subgradient of $J_{\boldsymbol{\lambda}}(\mathbf{x})$ at $\mathbf{x} = \hat{\boldsymbol{\beta}}$. For non-smooth regularizer like $J_{\boldsymbol{\lambda}}$, the subgradient $\partial J_{\boldsymbol{\lambda}}$ at $\mathbf{x}$ can reveal some information for detecting the zero coordinates of $\mathbf{x}$. A simple example is LASSO: $J_{\boldsymbol{\lambda}}(\mathbf{x}) = \lambda \|\mathbf{x}\|_1$. In this case, we have $x_i = 0$ as long as $[\partial J_{\boldsymbol{\lambda}}(\mathbf{x})]_i \in [0, \lambda)$. This identity is used in [50] to obtain the limiting sparsity level of LASSO estimator. Here, we extend this idea to SLOPE estimator, while a key difference is that unlike the LASSO case, the zero coordinates are not determined locally: whether $x_i = 0$ or not is not completely determined by $[\partial J_{\boldsymbol{\lambda}}(\mathbf{x})]_i$. This is mainly a consequence of non-separability of $J_{\boldsymbol{\lambda}}(\mathbf{x})$.

Remark 11: We can now explain where condition (R.1) originates from. By Lemma 5, we know a sufficient condition for $R_0^{(p)}/p \geq \mathbb{P}(\eta(Y_*) = 0)$ is that $|\hat{s}|_{(1:k)} \prec_S \lambda_{(1:k)}$ holds for some k satisfying $k/p \approx \mathbb{P}(\eta(Y_*) = 0) := q_0^*$. In the asymptotic limit, this translate into the following limiting form: for any $t \in [0, q_0^*)$,

$$\int_t^{q_0^*} F_{|Y_* - \eta(Y_*)|/\tau_*}^{-1}(u) du < \int_t^{q_0^*} F_{\Lambda}^{-1}(u) du. \quad (63)$$

In fact, (63) is exactly (R.1), since for any $u \in [0, q_0^*]$ we have $F_{|Y_* - \eta(Y_*)|/\tau_*}^{-1}(u) = F_{|Y_*|/\tau_*}^{-1}(u)$.

After combining the probabilistic upper and lower bounds, we conclude that $R_0^{(p)} \xrightarrow{\mathbb{P}} \mathbb{P}(\eta(Y_*) = 0)$.

3) *Convergence of $V^{(p)}$:* Finally, we prove the convergence of $V^{(p)}$. We have the following lemma, which shows $V^{(p)} \xrightarrow{\mathbb{P}} \mathbb{P}(\eta(Y_*) \neq 0, B = 0)$ can be implied by $R_0^{(p)} \xrightarrow{\mathbb{P}} \mathbb{P}(\eta(Y_*) = 0)$.

Lemma 6: For any $\varepsilon > 0$,

$$|V^{(p)} - \mathbb{P}[\eta(Y_*) \neq 0, B = 0]| \leq |\mathbb{P}(\eta(Y_*) = 0) - R_0^{(p)}| + \varepsilon,$$

with probability approaching 1 as $p \rightarrow \infty$.

The proof of Lemma 6 can be found in Appendix H. Since the convergence of $R_0^{(p)}$ has been established, we finish our proof.

D. Optimal Estimation

In this section, we prove the fundamental estimation performance of SLOPE, as stated in Proposition 4. The proof of part (a) and (b), which justifies the uniqueness of f_σ and the existence of σ_0 , can be found in Lemma 18 and Lemma 19 in Appendix M. Here we focus on proving part (c), which is the core part of Proposition 4.

From discussions before, finding the minimum MSE is equivalent to solving (24). Indeed, we have

$$\inf_{\mu_\Lambda \in \mathcal{P}_\Lambda} \lim_{p \rightarrow \infty} \frac{\|\hat{\beta} - \beta\|_2^2}{p} = \delta(\sigma_{\text{opt}}^2 - \sigma_w^2), \quad (64)$$

where σ_{opt} is the optimal value of (24), i.e., $\sigma_{\text{opt}} \stackrel{\text{def}}{=} \inf_{\mu_\Lambda \in \mathcal{P}_\Lambda} \sigma_*$.

1) *A reformulation of $\inf_{\mu_\Lambda \in \mathcal{P}_\Lambda} \sigma_*$:* We start by noting that σ_{opt} can be equivalently expressed as:

$$\sigma_{\text{opt}} = \inf\{\sigma \mid (\sigma, \tau) \in \mathcal{D}_L, \text{ for some } \tau > 0\}, \quad (65)$$

where

$$\mathcal{D}_L \stackrel{\text{def}}{=} \{(\sigma, \tau) \in \mathbb{R}_{>0}^2 : \exists \mu_\Lambda \in \mathcal{P}_\Lambda \text{ s.t. } (\sigma, \tau) \text{ satisfies (13)-(14)}\}.$$

Geometrically, computing σ_{opt} is equivalent to searching for the leftmost point in $\mathcal{D}_L$, which is the set of all realizable (σ, τ) pair. However, characterizing $\mathcal{D}_L$ is difficult, since it is determined in a convoluted way via (13)-(14). To simplify, consider instead the following equation of (f, σ, τ)

$$\sigma^2 = \sigma_w^2 + \frac{1}{\delta} \mathbb{E}[f(B + \sigma H) - B]^2 \quad (66)$$

$$1 = \tau \left[1 - \frac{1}{\delta} \mathbb{E}f'(B + \sigma H) \right] \quad (67)$$

where $(f, \sigma, \tau) \in \mathcal{I} \times \mathbb{R}_{>0} \times \mathbb{R}_{>0}$ and $H \sim \mathcal{N}(0, 1)$ is independent of $B \sim \mu_B$. Let us emphasize that although (66)-(67) has a similar form as (13)-(14), a key difference is that unlike η in (13)-(14), f in (66)-(67) is not dependent on other parameters such as (B, H, σ, τ) .

Now define the following set of (σ, τ) :

$$\mathcal{D}_F \stackrel{\text{def}}{=} \{(\sigma, \tau) \in \mathbb{R}_{>0}^2 : \exists f \in \mathcal{I} \text{ s.t. } (f, \sigma, \tau) \text{ satisfies (66)-(67)}\}.$$

A key step of our proof is to show $\mathcal{D}_L = \mathcal{D}_F$. This can be done as follows. Clearly, we have $\mathcal{D}_L \subseteq \mathcal{D}_F$, since $\eta \in \mathcal{I}$. To prove $\mathcal{D}_F \subseteq \mathcal{D}_L$, we need to utilize Proposition 3. Suppose $(\sigma, \tau) \in \mathcal{D}_F$ and let (f, σ, τ) be the corresponding solution of (66)-(67). If $\delta > 1$, we have $\mathcal{P}_\Lambda = \mathcal{P}_2(\mathbb{R})$ and by Proposition 3 we can take $\Lambda \sim \frac{|Y| - f(|Y|)}{\tau} \in \mathcal{P}_2(\mathbb{R})$ so that $\eta(\cdot; \mu_Y, \mu_{\tau\Lambda}) = f$; if $\delta \leq 1$, then from (67) we know $f(y) \neq y$ and $\mathbb{P}(\frac{|Y| - f(|Y|)}{\tau} \neq 0) > 0$, so $\frac{|Y| - f(|Y|)}{\tau} \in \mathcal{P}_\Lambda$ and we can still take $\Lambda \sim \frac{|Y| - f(|Y|)}{\tau}$ which gives us $\eta(\cdot; \mu_Y, \mu_{\tau\Lambda}) = f$. This means $(\sigma, \tau) \in \mathcal{D}_L$. As a result, we conclude that $\mathcal{D}_F \subseteq \mathcal{D}_L$ and thus $\mathcal{D}_L = \mathcal{D}_F$. Then substituting $\mathcal{D}_L = \mathcal{D}_F$ into (65), we get the following reformulation of σ_{opt} :

$$\sigma_{\text{opt}} = \inf\{\sigma \mid (\sigma, \tau) \in \mathcal{D}_F, \text{ for some } \tau > 0\}. \quad (68)$$

2) *Lower Bound of MSE*: Note that any $(f, \sigma) \in \mathcal{I} \times \mathbb{R}_{>0}$ satisfying (66)-(67) for some $\tau > 0$, should also satisfy

$$\begin{aligned} \mathbb{E}[f(B + \sigma H) - B]^2 &= \delta(\sigma^2 - \sigma_w^2) \\ \delta^{-1} \mathbb{E}f'(B + \sigma H) &\leq 1. \end{aligned} \quad (69)$$

Therefore, if we consider the following set of σ :

$$\mathcal{A} \stackrel{\text{def}}{=} \{\sigma > 0 : \exists f \in \mathcal{I}, \text{ s.t. } (f, \sigma) \text{ satisfies (69)}\}, \quad (70)$$

then from (68) we have

$$\sigma_{\text{opt}} \geq \inf \mathcal{A}. \quad (71)$$

Compared with σ_{opt} , the lower bound $\inf \mathcal{A}$ in (71) is easier to obtain, since the variable τ is dropped. In Lemma 19 in Appendix M, we show that $\inf \mathcal{A} = \sigma_0$. Therefore, $\sigma_{\text{opt}} \geq \sigma_0$. Together with (64), we prove (30).

3) *Reaching the Lower Bound*: We now show lower bound $\sigma_{\text{opt}} \geq \sigma_0$ is tight, if $\delta^{-1} \mathbb{E}[f'_{\sigma_0}(B + \sigma_0 H)] < 1$. Recall that f_{σ_0} is the unique optimal solution of (28) when $\sigma = \sigma_0$ and (f_{σ_0}, σ_0) satisfies (69). Let $\tau_0 = [1 - \delta^{-1} \mathbb{E}f'_{\sigma_0}(B + \sigma_0 H)]^{-1}$. It is not hard to see when $\delta^{-1} \mathbb{E}[f'_{\sigma_0}(B + \sigma_0 H)] < 1$, $\tau_0 \in (0, \infty)$ and $(f_{\sigma_0}, \sigma_0, \tau_0)$ is a solution of (66)-(67). This indicates $(\sigma_0, \tau_0) \in \mathcal{D}_F$ and thus from (68) we have $\sigma_{\text{opt}} \leq \sigma_0$. Together with the lower bound $\sigma_{\text{opt}} \geq \sigma_0$, we get $\sigma_{\text{opt}} = \sigma_0$.

On the other hand, by Proposition 3 we know if μ_Λ is taken as the law of $\frac{1}{\tau_0}(|Y_0| - f_{\sigma_0}(|Y_0|))$, then $f_{\sigma_0}(y) = \eta(y; \mu_{Y_0}, \mu_{\tau_0\Lambda})$. Since $(f_{\sigma_0}, \sigma_0, \tau_0)$ is a solution of equation (66)-(67), we know $(\sigma_*, \tau_*) = (\sigma_0, \tau_0)$, where (σ_*, τ_*) is the solution of fixed-point equation (13)-(14) under this choice of μ_Λ . According to (15), we get $\lim_{p \rightarrow \infty} \frac{\|\hat{\beta} - \beta\|_2^2}{p} = \delta(\sigma_0^2 - \sigma_w^2)$. This completes our proof.

E. Optimal Variable Selection

In this section, we are going to prove Proposition 5. First, part (a) and (b) can be proved in an analogous way as in Proposition 5, which is summarized in Lemma 21. Here we focus on part (c).

1) *Upper bound of $\mathcal{P}(\alpha)$* : Directly solving the original optimization (40) is not easy. Instead, replacing by a new objective function in (40), we will first consider the following problem:

$$\begin{aligned} \overline{\mathcal{P}}(\alpha) &\stackrel{\text{def}}{=} \sup_{\mu_\Lambda \in \tilde{\mathcal{P}}_\Lambda} \mathbb{P}\left(|B + \sigma_* H| \geq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_* \mid B \neq 0\right) \\ \text{s.t. } &\mathbb{P}(|\sigma_* H| \geq y_{\text{th}}^*) \leq \alpha \end{aligned} \quad (72)$$

where $y_{\text{th}}^* = \sup_{y \geq 0} \{y \mid \eta(y; \mu_{Y_*}, \mu_{\tau_*\Lambda}) = 0\}$. It is not hard to show for any $\alpha \in [0, 1]$, we have $\overline{\mathcal{P}}(\alpha) \geq \mathcal{P}(\alpha)$. This is because the constraint $\mathbb{P}(|\sigma_* H| \geq y_{\text{th}}^*) \leq \alpha$ in (40) implies $y_{\text{th}}^* \geq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_*$ and hence the objective function of (40) is upper bounded by that of (72) for any $\mu_\Lambda \in \tilde{\mathcal{P}}_\Lambda$.

Problem (72) can be further simplified. By direct differentiation, one can check for any fixed $b \neq 0$ and $c \geq 0$, the function $\sigma \mapsto \mathbb{P}(|b + \sigma H| \geq c\sigma)$ is non-increasing on $\mathbb{R}_{\geq 0}$, where $H \sim \mathcal{N}(0, 1)$. This then implies, by conditioning on B , that $\sigma \mapsto \mathbb{P}(|B + \sigma H| \geq c\sigma \mid B \neq 0)$ is non-increasing on $\mathbb{R}_{\geq 0}$ for any distribution of B satisfying $\mathbb{P}(B \neq 0) > 0$. Therefore, solving maximization problem in (72) is equivalent to solving the minimization problem of σ_* . Meanwhile, by the definition of y_{th}^* and the fact that $\eta(y; \mu_{Y_*}, \mu_{\tau_*\Lambda}) \in \mathcal{I}$, we know: $\mathbb{P}(|\sigma_* H| \geq y_{\text{th}}^*) \leq \alpha$ if and only if $\eta(y; \mu_{Y_*}, \mu_{\tau_*\Lambda}) = 0$, for all $|y| \leq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_*$. As a result, solving (72) is equivalent to solving:

$$\begin{aligned} \sigma_{\text{opt},\alpha} &\stackrel{\text{def}}{=} \inf_{\mu_\Lambda \in \tilde{\mathcal{P}}_\Lambda} \sigma_* \\ \text{s.t. } &\eta(y; \mu_{Y_*}, \mu_{\tau_*\Lambda}) = 0, \forall |y| \leq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_* \end{aligned} \quad (73)$$

and $\overline{\mathcal{P}}(\alpha)$ in (72) can be expressed in terms of $\sigma_{\text{opt},\alpha}$ in (73) as:

$$\overline{\mathcal{P}}(\alpha) = \mathbb{P}\left(|B + \sigma_{\text{opt},\alpha} H| \geq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{\text{opt},\alpha} \mid B \neq 0\right). \quad (74)$$

So far, we arrive at the optimization problem (73), which is similar to the one that we have analyzed in the estimation setting [c.f. (24)]. Yet there are two differences: (i) a constraint on η is added to ensure type-I error is bounded by α , (ii) a constraint on μ_Λ is added to guarantee valid limits of Type-I error and power exist (see Proposition 2). It turns out that the strategy we used can still be applied. The results are parallel to Proposition 4 part (c) and are summarized in Lemma 22 in Appendix N, where it is shown that

$$\sigma_{\text{opt},\alpha} \geq \sigma_{0,\alpha} \quad (75)$$

and the lower bound can be achieved when $\mu_\Lambda = \mu_{\text{opt},\alpha}$, if $\delta^{-1}\mathbb{E}[f'_\alpha(Y_{0,\alpha})] < 1$. After combining (75) with

$$\mathcal{P}(\alpha) \leq \overline{\mathcal{P}}(\alpha) = \mathbb{P}\left(|B + \sigma_{\text{opt},\alpha} H| \geq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{\text{opt},\alpha} \mid B \neq 0\right), \quad (76)$$

and (39), we get (43).

2) *Reaching the upper bound*: Now we show for any $\alpha \in [0, 1]$, the upper bound (43) is tight, if $\delta^{-1}\mathbb{E}[f'_\alpha(Y_{0,\alpha})] < 1$ and $y_{0,\alpha} = \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{0,\alpha}$. Also it is attained by $\mu_\Lambda = \mu_{\text{opt},\alpha}$. The case of $\alpha = 0$ is easy. Indeed, in this case, both sides of (43) equal to 0. We just need to verify the case of $\alpha \in (0, 1]$. By Lemma 22 and (76), we know it suffices to show $\mathcal{P}(\alpha) = \overline{\mathcal{P}}(\alpha)$ and also $\lim_{p \rightarrow \infty} \text{Power} = \mathcal{P}(\alpha)$, when $\mu_\Lambda = \mu_{\text{opt},\alpha}$. We verify this in Lemma 23 in Appendix N, which completes our proof.

VI. CONCLUDING REMARKS

We have established the asymptotic characterization of SLOPE in the high-dimensional regime. Although SLOPE is a high-dimensional regularized regression method, asymptotically its statistical performance can be fully characterized by a few scalar random variables. The precise characterization enabled us to derive the fundamental performance limits of SLOPE for both estimation and variable selection settings. Also we showed how to design the optimal regularizing sequences that achieve these limits.

Finally, let us point out some generalizations of current results that worth exploring in the future.

- 1) One major technical assumption in the current paper is that the sensing matrix is generated from i.i.d. Gaussian. There are two possible ways to relax this assumption. The first one is to consider the Gaussian design with correlated columns, which is the setting analyzed in [6]. Under this scenario, SLOPE enjoys the nice properties of selecting all the variables associated with highly correlated columns. It would be interesting to derive a precise explanation for this phenomenon. The second direction is staying in the i.i.d. setting, while generalizing to other ensembles, *e.g.*, sub-Gaussian distribution. This is to verify the so-called *universality* phenomenon and some works have been done in the setting where the regularizer is separable [54], [55]. It would be interesting to generalize these results to non-separable regularizers such as SLOPE.
- 2) The optimal designs of λ sequences considered in this paper are based on the assumption that the true distribution of unknown signal is known. The natural question is: can we design λ sequences without (or just with partial) such prior knowledge? One related problem is designing a regularizing sequence such that the false discovery rate is always controlled under a given level. In this setting, the realistic assumption is that we do not know the sparsity of underlying signal. For this purpose, a design of λ is proposed in [9] based on some qualitative insights. It would be nice to have quantitative results utilizing the exact characterizations derived here.
- 3) From numerical simulations, we can find that in several cases, the performance of practical λ sequences such as LASSO and BHq is comparable to the optimal performance. Is it possible that the optimal performance of SLOPE can actually be approximately achieved, when we are restricted to certain sub-classes of regularizing sequences? A key step is establishing some easy-to-evaluate bounds for the performance gap between practical and optimal sequences. One benefit of using practical sequences is that we can apply some purely data-dependent methods such as cross-validation to search for the optimal tuning parameter. Note that since general λ sequence includes order $\mathcal{O}(p)$ parameters, the grid search approach that is usually used in data-dependent method is not plausible here.

APPENDIX

A. Proof of Lemma 1

First assume $0 \leq y_1 \leq \dots \leq y_p$. Denote $\hat{\mathbf{y}} := \text{Prox}_{\lambda}(\mathbf{y})$. Then consider the linear interpolation of the points $\{(y_i, \hat{y}_i)\}_{i=0}^p$, where $(y_0, \hat{y}_0) = (0, 0)$:

$$g_p^+(y) = \begin{cases} y_{i-1} + \frac{\hat{y}_i - \hat{y}_{i-1}}{y_i - y_{i-1}}(y - y_{i-1}) & y \in (y_{i-1}, y_i), \\ \hat{y}_i & y = y_i, \\ \hat{y}_p + (y - y_p) & y > y_p. \end{cases} \quad (77)$$

By Fact 1 (iii), we know $g_p^+(y)$ is non-decreasing and 1-Lipschitz continuous on $\mathbb{R}_{\geq 0}$.

For general $\mathbf{y}$, we first obtain the linear interpolation $g_p^+(y)$ of the points $\{(|y|_{(i)}, |\hat{y}|_{(i)})\}_{i=0}^p$ as above. Then $g_p(y)$ can be constructed as follows:

$$g_p(y) = \begin{cases} g_p^+(y) & y \geq 0, \\ -g_p^+(-y) & y < 0. \end{cases}$$

Clearly, such $g_p(y)$ is an odd, non-decreasing and 1-Lipschitz function. Also by Fact 1 (i) and (ii), one can easily check $g_p(y_i) = \hat{y}_i$, for all $i \in [p]$. This finishes the proof.

B. Proof of Lemma 2

For notational simplicity, denote $\mathcal{M}_{\lambda}(\mathbf{y}) := \mathcal{M}_{\lambda}(\mathbf{y}; 1)$. Let $g_p^*(\mathbf{y})$ be any minimizer of (45). Since $\mathcal{M}_{\lambda}(\mathbf{y})$ is the minimum value of (5),

$$\begin{aligned} \mathcal{M}_{\lambda}(\mathbf{y}) &\leq \frac{1}{2} \|\mathbf{y} - g_p^*(\mathbf{y})\|_2^2 + \sum_{i=1}^p \lambda_i |g_p^*(\mathbf{y})|_{(i)} \\ &= p \mathcal{M}_{\lambda}^*(\mathbf{y}). \end{aligned} \quad (78)$$

Next, we show $p \mathcal{M}_{\lambda}^*(\mathbf{y}) \leq \mathcal{M}_{\lambda}(\mathbf{y})$. Given Lemma 1, this is immediate. Indeed,

$$\begin{aligned} \mathcal{M}_{\lambda}^*(\mathbf{y}) &\leq L_p(g_p) \\ &= \frac{1}{2} \|\mathbf{y} - \text{Prox}_{\lambda}(\mathbf{y})\|_2^2 + \sum_{i=1}^p \lambda_i |\text{Prox}_{\lambda}(\mathbf{y})|_{(i)} \\ &= \frac{\mathcal{M}_{\lambda}(\mathbf{y})}{p}, \end{aligned} \quad (79)$$

where g_p is the function we construct in Lemma 1, which satisfies $g_p \in \mathcal{I}$ and $g_p(\mathbf{y}) = \text{Prox}_{\lambda}(\mathbf{y})$. Combining (78) and (79), we get $\mathcal{M}_{\lambda}^*(\mathbf{y}) = \frac{\mathcal{M}_{\lambda}(\mathbf{y})}{p}$.

Substituting $\mathcal{M}_{\lambda}^*(\mathbf{y}) = \frac{\mathcal{M}_{\lambda}(\mathbf{y})}{p}$ into (78), we have for any minimizer g_p^* of (45), $g_p^*(\mathbf{y})$ is also a minimizer of (5). Since $\text{Prox}_{\lambda}(\mathbf{y})$ is the unique minimizer of (5), we know any minimizer of (45) should satisfy: $g_p^*(\mathbf{y}) = \text{Prox}_{\lambda}(\mathbf{y})$.

C. Auxiliary Results for Proving Proposition 1

Lemma 7: Suppose $\{\mathbf{y}^{(p)}\}_{p \in \mathbb{Z}^+}$ and $\{\boldsymbol{\lambda}^{(p)}\}_{p \in \mathbb{Z}^+}$ are converging sequences with limiting measure μ_Y and μ_Λ . Then

$$\sup_{g \in \mathcal{I}} |L(g) - L_p(g)| \rightarrow 0, \quad (80)$$

where $L(g)$ and $L_p(g)$ are defined in (46) and (45).

Proof: The first step is to establish the following uniform convergence of a class of Pseudo-Lipschitz functions. Let Ψ be the set of all functions $\psi : \mathbb{R} \rightarrow \mathbb{R}$ satisfying: $\psi(0) = 0$ and $|\psi(x) - \psi(y)| \leq (1 + |x| + |y|)|x - y|$ for any $x, y \in \mathbb{R}$. Then

$$\sup_{\psi \in \Psi} |\mathbb{E}_{\mu_{\mathbf{y}}} \psi(Y) - \mathbb{E}_{\mu_Y} \psi(Y)| \rightarrow 0. \quad (81)$$

To prove (81), first for any $\psi \in \Psi$ consider the truncation:

$$\hat{\psi}_A(y) = \begin{cases} \psi(-A) & y < -A, \\ \psi(y) & |y| \leq A, \\ \psi(A) & y > A, \end{cases}$$

where $A > 0$ is a constant. It is easy to check $\hat{\psi}_A(y)$ is $(1 + 2A)$ -Lipschitz continuous, so

$$|\mathbb{E}_{\mu_{\mathbf{y}}} \hat{\psi}_A(Y) - \mathbb{E}_{\mu_Y} \hat{\psi}_A(Y)| \stackrel{(a)}{\leq} (1 + 2A)W_1(\mu_{\mathbf{y}}, \mu_Y) \stackrel{(b)}{\leq} (1 + 2A)W_2(\mu_{\mathbf{y}}, \mu_Y),$$

where (a) follows from Kantonovich duality theorem ([56, Theorem 1.3]) and (b) follows from Holder's inequality.

Therefore,

$$\begin{aligned} |\mathbb{E}_{\mu_{\mathbf{y}}} \psi(Y) - \mathbb{E}_{\mu_Y} \psi(Y)| &\leq |\mathbb{E}_{\mu_{\mathbf{y}}} \hat{\psi}_A(Y) - \mathbb{E}_{\mu_Y} \hat{\psi}_A(Y)| \\ &\quad + |\mathbb{E}_{\mu_{\mathbf{y}}} \hat{\psi}_A(Y) - \mathbb{E}_{\mu_{\mathbf{y}}} \psi(Y)| + |\mathbb{E}_{\mu_Y} \hat{\psi}_A(Y) - \mathbb{E}_{\mu_Y} \psi(Y)| \\ &\leq (1 + 2A)W_2(\mu_{\mathbf{y}}, \mu_Y) \\ &\quad + 2\mathbb{E}_{\mu_{\mathbf{y}}} [\mathbb{I}_{|Y| \geq A}(Y^2 + |Y|)] + 2\mathbb{E}_{\mu_Y} [\mathbb{I}_{|Y| \geq A}(Y^2 + |Y|)]. \end{aligned} \quad (82)$$

For any $\varepsilon > 0$, each term on the RHS of (82) can be bounded as follows. Since $(\mathbb{E}_{\mu_Y} |Y|)^2 \leq \mathbb{E}_{\mu_Y} Y^2 < \infty$, by DCT there always exists $A > 0$ such that $\mathbb{E}_{\mu_Y} [\mathbb{I}_{|Y| \geq A}(Y^2 + |Y|)] \leq \frac{\varepsilon}{4}$. On the other hand, for any given $A > 0$ and $\varepsilon > 0$, there always exists $p_0 \in \mathbb{N}$ such that for any $p \geq p_0$, (i) $W_2(\mu_{\mathbf{y}}, \mu_Y) \leq \frac{\varepsilon}{4(1+2A)}$, since $W_2(\mu_{\mathbf{y}}, \mu_Y) \rightarrow 0$ and (ii) $\mathbb{E}_{\mu_{\mathbf{y}}} [\mathbb{I}_{|Y| \geq A}(Y^2 + |Y|)] \leq \mathbb{E}_{\mu_Y} [\mathbb{I}_{|Y| \geq A}(Y^2 + |Y|)] + \frac{\varepsilon}{4}$ by Theorem 7.12 (iv) in [56]. Note that the RHS of (82) does not depend on ψ , so for any $\varepsilon > 0$, there exists $p_0 \in \mathbb{N}$ such that for any $p \geq p_0$, $\sup_{\psi \in \Psi} |\mathbb{E}_{\mu_{\mathbf{y}}} \psi(Y) - \mathbb{E}_{\mu_Y} \psi(Y)| \leq \varepsilon$. Therefore, (81) is proved.

We are now ready to show (80). Recalling the definitions of $L(g)$ and $L_p(g)$ in (46) and (45), we have

$$\begin{aligned} \sup_{g \in \mathcal{I}} |L(g) - L_p(g)| &\leq \frac{1}{2} \sup_{g \in \mathcal{I}} |\mathbb{E}_{\mu_{\mathbf{y}}} [Y - g(Y)]^2 - \mathbb{E}_{\mu_Y} [Y - g(Y)]^2| \\ &\quad + \sup_{g \in \mathcal{I}} \int_0^1 |F_{\boldsymbol{\lambda}}^{-1}(u) - F_{\Lambda}^{-1}(u)| F_{|g(Y)|}^{-1}(u) du \\ &\quad + \sup_{g \in \mathcal{I}} \int_0^1 F_{\boldsymbol{\lambda}}^{-1}(u) |F_{|g(Y)|}^{-1}(u) - F_{|g(\mathbf{y})|}^{-1}(u)| du. \end{aligned} \quad (83)$$

Therefore, it remains to control each term on the RHS of (83). The first term can be handled by using (81), since $y \mapsto [y - g(y)]^2$ belongs to Ψ for $g \in \mathcal{I}$; the second term can be controlled as : Term II $\leq W_2(\mu_\lambda, \mu_\Lambda) \sqrt{\mathbb{E}_{\mu_Y} Y^2}$ using (86) and Cauchy-Swartz inequality; similarly for the third term, we have

$$\begin{aligned} \text{Term III} &\leq \sup_{g \in \mathcal{I}} \sqrt{\mathbb{E}_{\mu_\lambda} \Lambda^2} W_2(\mu_{|g(\mathbf{y})|}, \mu_{|g(Y)|}) \\ &\leq \sqrt{\mathbb{E}_{\mu_\lambda} \Lambda^2} W_2(\mu_{\mathbf{y}}, \mu_Y), \end{aligned}$$

where the last inequality follows from the definition of Wasserstein distance and the Lipschitz continuity of g :

$$\begin{aligned} W_2(\mu_{|g(\mathbf{y})|}, \mu_{|g(Y)|})^2 &= \inf_{\pi \in \Pi(\mu_{|g(\mathbf{y})|}, \mu_{|g(Y)|})} \int (g - h)^2 d\pi(g, h) \\ &= \inf_{\pi \in \Pi(\mu_{\mathbf{y}}, \mu_Y)} \int [|g(x)| - |g(y)|]^2 d\pi(x, y) \\ &\leq \inf_{\pi \in \Pi(\mu_{\mathbf{y}}, \mu_Y)} \int (x - y)^2 d\pi(x, y) \\ &= W_2(\mu_{\mathbf{y}}, \mu_Y)^2. \end{aligned}$$

Substituting the above bounds back to (83) and using the assumption that $W_2(\mu_\lambda, \mu_\Lambda), W_2(\mu_{\mathbf{y}}, \mu_Y) \rightarrow 0$, we obtain the desired results. $\blacksquare$

Lemma 8: The optimization problem (6) has an optimal solution and it is unique (up to a set of measure 0 with respect to μ_Y).

Proof: Without loss of generality, we assume $\tau = 1$. The objective function $L(g)$ of (6) is defined on the following L^2 space:

$$\mathcal{H}_{\mu_Y} \stackrel{\text{def}}{=} \{g(y) \mid g(y) \text{ is measurable and } \|g\|_{\mu_Y} < \infty\} \quad (84)$$

where $\|g\|_{\mu_Y} \stackrel{\text{def}}{=} [\mathbb{E}_{\mu_Y} g^2(Y)]^{1/2}$. It is known that in L^2 space (and more generally in all normed linear spaces), the convention is to work with *equivalence class* of functions [57, p.135-136]. The equivalence class of a function $f \in \mathcal{H}_{\mu_Y}$, denoted as $[f]$, is the collection of all functions $g \in \mathcal{H}_{\mu_Y}$ satisfying $\|g - f\|_{\mu_Y} = 0$. As a notational convention, we will write $[f]$ as f , and the set $\{[f] : f \in \mathcal{I}\}$ as $\mathcal{I}$. Also $\|g - f\|_{\mu_Y} = 0$ will be denoted as $g = f$.

We first show $L(g)$ is 1-strongly convex on $\mathcal{H}_{\mu_Y}$, i.e., for any $g_1, g_2 \in \mathcal{H}_{\mu_Y}$,

$$L(\theta g_1 + (1 - \theta)g_2) \leq \theta L(g_1) + (1 - \theta)L(g_2) - \frac{\theta(1 - \theta)}{2} \|g_2(Y) - g_1(Y)\|^2. \quad (85)$$

First, for any $\theta \in [0, 1]$,

$$\begin{aligned} \int_0^1 F_\Lambda^{-1}(u) F_{|\theta g_1(Y) + (1 - \theta)g_2(Y)|}^{-1}(u) du &\leq \int_0^1 F_\Lambda^{-1}(u) F_{|\theta g_1(Y)| + (1 - \theta)|g_2(Y)|}^{-1}(u) du \\ &= \theta \int_0^1 F_\Lambda^{-1}(u) F_{|g_1(Y)|}^{-1}(u) du + (1 - \theta) \int_0^1 F_\Lambda^{-1}(u) F_{|g_2(Y)|}^{-1}(u) du, \end{aligned}$$

which implies that $L_1(g) := \int_0^1 F_\Lambda^{-1}(u) F_{|g(Y)|}^{-1}(u) du$ is convex. Also, it is not hard to check $L_2(g) := \frac{1}{2} \mathbb{E}_\mu[Y - g(Y)]^2$ is 1-strongly convex by definition (85). Then the strong convexity of $L(g)$ follows, since $L(g) = L_1(g) +$

$L_2(g)$. On the other hand, we can show $L(g)$ is continuous on $\mathcal{H}_{\mu_Y}$. Indeed,

$$\begin{aligned} |L(g_2) - L(g_1)| &\leq \frac{1}{2} \|g_2 - g_1\|_{\mu_Y} \cdot \|2y - g_1 - g_2\|_{\mu_Y} \\ &\quad + \sqrt{\mathbb{E}\Lambda^2} \left| \|g_2\|_{\mu_Y} - \|g_1\|_{\mu_Y} \right| \\ &\leq \|g_2 - g_1\|_{\mu_Y} (2\|y\|_{\mu_Y} + 2\|g_1\|_{\mu_Y} + \|g_2 - g_1\|_{\mu_Y} + \sqrt{\mathbb{E}\Lambda^2}). \end{aligned}$$

Since $\|y\|_{\mu_Y}, \|g_1\|_{\mu_Y}, \|g_2\|_{\mu_Y} < \infty$, we conclude that $L(g)$ is continuous.

Next we are going to show the set $\mathcal{I}$ is convex, bounded and closed in $\mathcal{H}_{\mu_Y}$. The convexity can be directly checked by definition. Choose any $g_1, g_2 \in \mathcal{I}$. Then there exists $S \subseteq \mathbb{R}$ with $\mu_Y(S) = 1$ such that for any $y_1, y_2 \in S$, $y_1 \leq y_2$ and any $y \in S$, we have $0 \leq g_i(y_2) - g_i(y_1) \leq y_2 - y_1$ and $g_i(y) = -g_i(-y)$, where $i = 1, 2$. Then for any $\theta \in [0, 1]$, function $\theta g_1 + (1 - \theta)g_2$ also satisfy (i) and (ii) on S , so $\theta g_1 + (1 - \theta)g_2 \in \mathcal{I}$. The boundedness directly follows from the fact that for any $g \in \mathcal{I}$, $g(y) \leq |y|$ on some $S \subseteq \mathbb{R}$ with $\mu_Y(S) = 1$. To show closedness, suppose $g_k(y) \in \mathcal{I}, k = 1, 2, \dots$ is a sequence of functions that converge to some $g(y) \in \mathcal{H}_{\mu_Y}$. Then by Riesz-Fischer Theorem, there exists a sub-sequence of $\{g_k(y)\}_{k \in \mathbb{Z}^+}$ that converges point-wise to $g(y)$ on some $S \subseteq \mathbb{R}$ with $\mu_Y(S) = 1$. By this μ_Y -almost everywhere convergence of $g_k(y)$ to $g(y)$, we know there exists some $S' \subseteq \mathbb{R}$ with $\mu_Y(S') = 1$, such that for any $y_1, y_2 \in S', y_1 \leq y_2$ and any $y \in S'$, it holds that $0 \leq g(y_2) - g(y_1) \leq y_2 - y_1$ and $g(y) = -g(-y)$. Therefore, $g(y) \in \mathcal{I}$ and thus $\mathcal{I}$ is closed.

The final step is to apply Theorem 17 in [57, Chap. 8] to conclude that (6) has an optimal solution $g^* \in \mathcal{I}$. Also the uniqueness of g^* can be easily checked by the strong convexity of $L(g)$. Suppose there exists two different optimal solutions, g_1^*, g_2^* with $L(g_1^*) = L(g_2^*)$ and $g_1^* \neq g_2^*$. Then by (85), for $g = \frac{g_1^* + g_2^*}{2} \in \mathcal{I}$ we have $L(g) < L(g_1^*) = L(g_2^*)$, which leads to a contradiction. ■

The following result provides the explicit formula for calculating Wasserstein-2 distance between probability measure on $\mathbb{R}$. Readers can find a proof in Theorem 2.18 in [56].

Lemma 9: Suppose $\mu_1, \mu_2 \in \mathcal{P}_2(\mathbb{R})$ and the corresponding quantile functions are F_1^{-1} and F_2^{-1} . Then

$$W_2(\mu_1, \mu_2)^2 = \int_0^1 (F_1^{-1}(t) - F_2^{-1}(t))^2 dt. \quad (86)$$

D. Auxiliary Results for Proving Theorem 1

In this section, we prove three auxiliary lemmas used in the proof of Theorem 1.

The first two results are on the asymptotic properties of auxiliary problem (52). To state these asymptotic results, similar as Theorem 1, we will consider a sequence of auxiliary problems described by the instances $\{\mathbf{g}^{(p)}, \mathbf{h}^{(p)}, \boldsymbol{\beta}^{(p)}, \mathbf{w}^{(p)}, \boldsymbol{\lambda}^{(p)}\}_{p \in \mathbb{Z}^+}$. They satisfy the following: (i) $\mathbf{g}^{(p)} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$, $\mathbf{h}^{(p)} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$, $p \in \mathbb{Z}^+$ are all independent, (ii) $\{\boldsymbol{\beta}^{(p)}\}_{p \in \mathbb{Z}^+}, \{\mathbf{w}^{(p)}\}_{p \in \mathbb{Z}^+}, \{\boldsymbol{\lambda}^{(p)}\}_{p \in \mathbb{Z}^+}$ are the same converging sequences as in Theorem 1. Here the requirement that $\{\mathbf{g}^{(p)}\}_{p \in \mathbb{Z}^+}$ and $\{\mathbf{h}^{(p)}\}_{p \in \mathbb{Z}^+}$ are independent is not completely necessary, since we are only aiming for results regarding convergence in probability. The independence assumption simply allows us to directly apply some results obtained in Appendix K.

The first lemma is about the minimum value and the minimizer of $L(\mathbf{v})$ over a bounded Euclidean ball. Recall that

$$L(\mathbf{v}) \stackrel{\text{def}}{=} \frac{1}{2} \left(\sqrt{\frac{\|\mathbf{v}\|^2}{n} \frac{\|\mathbf{g}\|^2}{n} + \frac{\|\mathbf{w}\|^2}{n} + 2 \frac{\|\mathbf{v}\|}{\sqrt{n}} \frac{\mathbf{g}^\top \mathbf{w}}{n} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right)_+^2 + \frac{J_\lambda(\mathbf{v} + \boldsymbol{\beta})}{n}. \quad (87)$$

Lemma 10: Let Ψ_* and (σ_*, θ_*) be the optimal value and solution of the minimax problem in (55) and $\theta_{\min} > 0$ be the lower bound of θ_* obtained in Lemma 14. For any $\varepsilon > 0$ and $K \geq \frac{\sigma_*}{\sqrt{\delta}} + \frac{\theta_{\min}}{4}$, we have

$$\mathbb{P} \left(\left| \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) - \Psi_* \right| \leq \varepsilon \right) \rightarrow 1 \quad (88)$$

and

$$\mathbb{P} \left(\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{32n\varepsilon/\gamma}}(\hat{\mathbf{v}}) \cap \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) \geq \Psi_* + \varepsilon \right) \rightarrow 1, \quad (89)$$

where $\gamma = \frac{\theta_{\min} \sigma_w^2}{4(K^2 + \sigma_w^2)^{3/2}}$ and

$$\hat{\mathbf{v}} := \eta(\boldsymbol{\beta} + \sigma_* \mathbf{h}) - \boldsymbol{\beta}. \quad (90)$$

Here in (90), $\eta(\cdot) := \eta(\boldsymbol{\beta} + \sigma_* \mathbf{h}; \mu_{Y_*}, \mu_{\sigma_* \Lambda / \theta_*})$ and $Y_* = B + \sigma_* H$, with $H \sim \mathcal{N}(0, 1)$ independent of $B \sim \mu_B$.

Proof: We follow the proof of Proposition B.2 in [50]. First introduce the event $\mathcal{A} = \bigcap_{i=1}^5 \mathcal{A}_i$, where

$$\begin{aligned} \mathcal{A}_1 &:= \left\{ \frac{\|\mathbf{g}\|^2}{n}, \frac{\|\mathbf{h}\|^2}{p}, \frac{\|\mathbf{w}\|^2}{\sigma_w^2 n} \in [1 - \varsigma, 1 + \varsigma], \left| \frac{\mathbf{g}^\top \mathbf{w}}{n} \right| \leq \varsigma \right\}, \\ \mathcal{A}_2 &:= \{ \|\hat{\mathbf{v}}\| / \sqrt{p} \leq \sigma_* \}, \\ \mathcal{A}_3 &:= \left\{ \sqrt{\frac{\|\hat{\mathbf{v}}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \hat{\mathbf{v}}}{n} \geq \frac{\theta_{\min}}{2} \right\}, \\ \mathcal{A}_4 &:= \{ |\tilde{L}(\hat{\mathbf{v}}) - \Psi_*| \leq \varepsilon \}, \\ \mathcal{A}_5 &:= \left\{ \sup_{\sigma \in [\sigma_w, \sqrt{\sigma_w^2 + K^2}]} \left| \left(\mathcal{F}_p(\sigma, \theta_*) - \frac{\sigma \theta_* \|\mathbf{h}\|^2}{2p} \right) - \left(\mathcal{F}(\sigma, \theta_*) - \frac{\sigma \theta_*}{2} \right) \right| \leq \delta \varepsilon \right\}, \end{aligned} \quad (91)$$

with $\varepsilon > 0$ and $\varsigma \in (0, \frac{1}{2})$. In (91), $\mathcal{F}_p$ and $\mathcal{F}$ are the same as in (160) and (163) and $\tilde{L}(\mathbf{v})$ is defined as:

$$\tilde{L}(\mathbf{v}) \stackrel{\text{def}}{=} \frac{1}{2} \left(\sqrt{\frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right)_+^2 + \frac{J_\lambda(\mathbf{v} + \boldsymbol{\beta})}{n}. \quad (92)$$

Based on the event $\mathcal{A}$, our subsequent analysis will become fully deterministic: we will condition on fixed $\mathbf{h}$ and $\mathbf{g}$ in $\mathcal{A}$. Before doing so, let us first show each of the events $\mathcal{A}_1 \sim \mathcal{A}_5$ occurs with probability approaching 1 as $p \rightarrow \infty$, so $\mathbb{P}(\mathcal{A}) \rightarrow 1$. This will ensure all the results obtained by conditioning on $\mathcal{A}$ hold with probability approaching 1.

$\mathcal{A}_1$: By the law of large number and the fact that $\{\mathbf{w}\}_{p \in \mathbb{N}}$ is a converging sequence with limiting variance $\sigma_w^2 > 0$, it is not hard to show $\mathbb{P}(\mathcal{A}_1) \rightarrow 1$.

$\mathcal{A}_2$: From (179), we have

$$\frac{\|\hat{\mathbf{v}}\|^2}{p} \xrightarrow{a.s.} \mathbb{E}[B - \eta(B + \sigma_* H)]^2. \quad (93)$$

Then together with (151), we get $\frac{\|\hat{\mathbf{v}}\|^2}{p} \xrightarrow{a.s.} (\sigma_*)^2 - \sigma_w^2$. Therefore, $\mathbb{P}(\mathcal{A}_2) \rightarrow 1$, since $\sigma_w^2 > 0$ by assumption.

$\mathcal{A}_3$: From (179) and (181), we can get

$$\sqrt{\frac{\|\hat{\mathbf{v}}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \hat{\mathbf{v}}}{n} \xrightarrow{a.s.} \sqrt{\frac{1}{\delta} \mathbb{E}[B - \eta(B + \sigma_* H)]^2 + \sigma_w^2} - \frac{\sigma_* \mathbb{E} \eta'(B + \sigma_* H)}{\delta} = \theta_*, \quad (94)$$

where the last step follows from (151). From Lemma 14, there exists $\theta_{\min} > 0$ such that $\theta_* \geq \theta_{\min}$. Therefore, $\mathbb{P}(\mathcal{A}_3) \rightarrow 1$.

$\mathcal{A}_4$: From the definition of $\mathcal{F}(\sigma, \theta)$ in (163),

$$\begin{aligned}\mathcal{F}(\sigma_*, \theta_*) &= \frac{\theta_*}{2\sigma_*} \mathbb{E}[(Y_*) - \eta(Y_*)]^2 + \int_0^1 F_\Lambda^{-1}(u) F_{|\eta(Y_*)|}^{-1}(u) du. \\ &= \frac{\theta_*}{2\sigma_*} \left[\mathbb{E}(\eta(Y_*) - B)^2 - 2(\sigma_*)^2 \mathbb{E}(\eta'(Y_*)) \right] + \frac{\theta_* \sigma_*}{2} \\ &\quad + \int_0^1 F_\Lambda^{-1}(u) F_{|\eta(Y_*)|}^{-1}(u) du \\ &= \frac{\theta_* \delta}{2\sigma_*} (- (\sigma_*)^2 - \sigma_w^2 + 2\sigma_* \sigma_w) + \int_0^1 F_\Lambda^{-1}(u) F_{|\eta(Y_*)|}^{-1}(u) du + \frac{\theta_* \sigma_*}{2},\end{aligned}\tag{95}$$

where in the last step we use (151). Then substituting (95) into (55), we get

$$\Psi_* = \Psi(\sigma_*, \theta_*) = \frac{(\theta_*)^2}{2} + \int_0^1 F_\Lambda^{-1}(u) F_{|\eta(Y_*)|}^{-1}(u) du.\tag{96}$$

On the other hand,

$$\begin{aligned}\tilde{L}(\mathring{v}) &= \frac{1}{2} \left(\sqrt{\frac{\|\mathring{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathring{v}}{n} \right)_+^2 + \frac{J_\lambda(\mathring{v} + \beta)}{n} \\ &\xrightarrow{a.s.} \frac{(\theta_*)^2}{2} + \int_0^1 F_\Lambda^{-1}(u) F_{|\eta(Y_*)|}^{-1}(u) du.\end{aligned}\tag{97}$$

where we use (94). From (96) and (97), we have $\tilde{L}(\mathring{v}) \xrightarrow{a.s.} \Psi_*$. Therefore, $\mathbb{P}(\mathcal{A}_4) \rightarrow 1$ for any $\varepsilon > 0$.

$\mathcal{A}_5$: From (161) and strong law of large number for triangular array [58, Theorem 2.1], we have $\mathcal{F}_p(\sigma, \theta_*) - \frac{\sigma \theta_* \|\mathbf{h}\|^2}{2p} \xrightarrow{a.s.} \mathcal{F}(\sigma, \theta_*) - \frac{\sigma \theta_*}{2}$ for any $\sigma \in [\sigma_w, \sqrt{\sigma_w^2 + K^2}]$. From (167) in Lemma 16, we know $\mathcal{F}(\cdot, \theta_*)$ is Lipschitz continuous on $[\sigma_w, \sqrt{\sigma_w^2 + K^2}]$. This indicates $\sigma \mapsto \mathcal{F}(\sigma, \theta_*) - \frac{\sigma \theta_*}{2}$ is also Lipschitz continuous on $[\sigma_w, \sqrt{\sigma_w^2 + K^2}]$. On the other hand, in the proof of Lemma 16 [line below (174)], we show $\mathcal{F}_p(\cdot, \theta_*)$ is continuously differentiable on $[\sigma_w, \infty)$ with derivative satisfying $\left| \frac{\partial \mathcal{F}_p(\sigma, \theta_*)}{\partial \sigma} \right| \leq \frac{3\theta_* \|\beta\|^2 + \sigma^2(2 + \theta_*) \|\mathbf{h}\|^2}{\sigma^2 p}$. Then by [58, Theorem 2.1] again and the fact that $\{\beta\}_{p \in \mathbb{Z}^+}$ is a converging sequence, we have almost surely $\limsup_{p \rightarrow \infty} \left| \frac{\partial \mathcal{F}_p(\sigma, \theta_*)}{\partial \sigma} - \frac{\theta_* \|\mathbf{h}\|^2}{2p} \right| \leq C$ for any $\sigma \in [\sigma_w, \sqrt{\sigma_w^2 + K^2}]$, where $C > 0$ is some constant. This indicates that almost surely, $\sigma \mapsto \frac{\partial \mathcal{F}_p(\sigma, \theta_*)}{\partial \sigma} - \frac{\theta_* \|\mathbf{h}\|^2}{2p}$ is C -Lipschitz continuous on $[\sigma_w, \sqrt{\sigma_w^2 + K^2}]$ for any large enough p . Then by the same epsilon net argument as in the proof of Lemma 16, we can show

$$\sup_{\sigma \in [\sigma_w, \sqrt{\sigma_w^2 + K^2}]} \left| \left(\mathcal{F}_p(\sigma, \theta_*) - \frac{\sigma \theta_* \|\mathbf{h}\|^2}{2p} \right) - \left(\mathcal{F}(\sigma, \theta_*) - \frac{\sigma \theta_*}{2} \right) \right| \xrightarrow{a.s.} 0.$$

Therefore, $\mathbb{P}(\mathcal{A}_5) \rightarrow 1$ for any $\varepsilon > 0$.

Now we are ready to start the deterministic analysis conditioned on the event $\mathcal{A}$. It is more convenient to work with $\tilde{L}(v)$ than $L(v)$, since it is locally strongly convex (the precise meaning will be given below). In the sequel,

we will start by studying the limiting properties of $\tilde{L}(\mathbf{v})$ and then associate them to $L(\mathbf{v})$, by showing $L(\mathbf{v})$ can be well-approximated by $\tilde{L}(\mathbf{v})$ as $p \rightarrow \infty$. For $\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}$, $\tilde{L}(\mathbf{v})$ can be equivalently written as:

$$\begin{aligned}\tilde{L}(\mathbf{v}) &= \max_{\theta \geq 0} \theta \left(\sqrt{\frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right) - \frac{\theta^2}{2} + \frac{J_\lambda(\mathbf{v} + \boldsymbol{\beta})}{n} \\ &= \max_{\theta \geq 0} \theta \left(\min_{\sigma \in [\sigma_w, \sqrt{\sigma_w^2 + K^2}]} \left\{ \frac{\sigma}{2} + \frac{\|\mathbf{v}\|^2/n + \sigma_w^2}{2\sigma} \right\} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right) - \frac{\theta^2}{2} + \frac{J_\lambda(\mathbf{v} + \boldsymbol{\beta})}{n} \\ &= \max_{\theta \geq 0} \min_{\sigma \in [\sigma_w, \sqrt{\sigma_w^2 + K^2}]} \frac{\theta}{2} \left(\frac{\sigma_w^2}{\sigma} + \sigma \right) - \frac{\theta^2}{2} + \frac{1}{n} \left[\frac{\theta \|\mathbf{v}\|^2}{2\sigma} - \theta \mathbf{h}^\top \mathbf{v} + J_\lambda(\mathbf{v} + \boldsymbol{\beta}) \right].\end{aligned}\quad (98)$$

Therefore,

$$\begin{aligned}\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) &\stackrel{(a)}{\geq} \min_{\sigma \in [\sigma_w, \sqrt{\sigma_w^2 + K^2}]} \frac{\theta_*}{2} \left(\frac{\sigma_w^2}{\sigma} + \sigma \right) - \frac{(\theta_*)^2}{2} + \frac{1}{\delta} \left[\mathcal{F}_p(\sigma, \theta_*) - \frac{\sigma \theta_* \|\mathbf{h}\|^2}{2p} \right], \\ &\stackrel{(b)}{\geq} \min_{\sigma \in [\sigma_w, \sqrt{\sigma_w^2 + K^2}]} \Psi(\sigma, \theta_*) - \varepsilon \\ &\geq \Psi_* - \varepsilon \\ &\stackrel{(c)}{\geq} \tilde{L}(\hat{\mathbf{v}}) - 2\varepsilon,\end{aligned}\quad (99)$$

where (a) follows from (98) and (160), (b) is due to $\mathcal{A}_5$ and (c) is due to $\mathcal{A}_4$. Besides, since $\frac{\|\hat{\mathbf{v}}\|_2}{\sqrt{n}} \leq \frac{\sigma_*}{\sqrt{\delta}}$ under $\mathcal{A}_2$ and $\frac{\sigma_*}{\sqrt{\delta}} \leq K - \frac{\theta_{\min}}{4}$ by assumption, we have $\frac{\|\hat{\mathbf{v}}\|_2}{\sqrt{n}} \leq K$ and thus $\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) \leq \tilde{L}(\hat{\mathbf{v}}) \leq \Psi_* + \varepsilon$. Therefore, combining it with (99) yields

$$\left| \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) - \Psi_* \right| \leq 2\varepsilon. \quad (100)$$

On the other hand, $\mathbf{v} \mapsto \sqrt{\frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathbf{v}}{n}$ is $\frac{1}{\sqrt{n}}(\sqrt{\frac{3}{2\delta}} + 1)$ -Lipschitz continuous under $\mathcal{A}_1$ and $\sqrt{\frac{\|\hat{\mathbf{v}}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \hat{\mathbf{v}}}{n} \geq \frac{\theta_{\min}}{2}$ under $\mathcal{A}_3$. Therefore, for $r = (\sqrt{\frac{3}{2\delta}} + 1)^{-1} \frac{\theta_{\min}}{4}$, $\sqrt{\frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \geq \frac{\theta_{\min}}{4}$ for all $\mathbf{v} \in \mathcal{B}_{\sqrt{n}r}(\hat{\mathbf{v}})$. Then using Lemma F.14 in [50] and $\frac{\|\hat{\mathbf{v}}\|_2}{\sqrt{n}} + r \leq \frac{\sigma_*}{\sqrt{n}} + r < K$ (hence $\mathcal{B}_{\sqrt{n}r}(\hat{\mathbf{v}}) \subset \mathcal{B}_{\sqrt{n}K}$) under $\mathcal{A}_2$, we can show $\tilde{L}(\mathbf{v})$ is $\frac{\gamma}{n}$ -strongly convex on $\mathcal{B}_{\sqrt{n}r}(\hat{\mathbf{v}})$, where $\gamma = \frac{\theta_{\min} \sigma_w^2}{4(K^2 + \sigma_w^2)^{3/2}}$. In other words, $\tilde{L}(\mathbf{v})$ is locally strongly convex in $\mathcal{B}_{\sqrt{n}r}(\hat{\mathbf{v}})$. Also since $\mathcal{B}_{\sqrt{n}r}(\hat{\mathbf{v}}) \subset \mathcal{B}_{\sqrt{n}K}$, together with (99) we have $\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}r}(\hat{\mathbf{v}})} \tilde{L}(\mathbf{v}) \geq \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) \geq \tilde{L}(\hat{\mathbf{v}}) - 2\varepsilon$. By Lemma B.1 in [50] we know if $0 < 2\varepsilon \leq \frac{(\sqrt{n}r)^2 \gamma}{8n}$, then $\|\tilde{\mathbf{v}} - \hat{\mathbf{v}}\|^2 \leq \frac{4n\varepsilon}{\gamma} \leq \frac{nr^2}{4}$, where $\tilde{\mathbf{v}} = \underset{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}}{\operatorname{argmin}} \tilde{L}(\mathbf{v})$. Moreover, for any $\mathbf{v} \in \mathcal{B}_{4\sqrt{n\varepsilon/\gamma}}^o(\hat{\mathbf{v}})$, we have $\tilde{L}(\mathbf{v}) \geq \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) + 2\varepsilon$. This implies

$$\min_{\mathbf{v} \in \mathcal{B}_{4\sqrt{n\varepsilon/\gamma}}^o(\hat{\mathbf{v}}) \cap \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) \geq \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) + 2\varepsilon. \quad (101)$$

Finally, we show $L(\mathbf{v})$ is well-approximated by $\tilde{L}(\mathbf{v})$ under event $\mathcal{A}_1$. Note that $L(\mathbf{v}) - \tilde{L}(\mathbf{v}) = g(\Delta)$, where $g(t) = \frac{1}{2}[(\sqrt{x+t} - y)_+^2 - (\sqrt{x} - y)_+^2]$, with $x = \frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2$, $y = \frac{\mathbf{h}^\top \mathbf{v}}{n}$ and $\Delta = \frac{\|\mathbf{v}\|^2}{n} \left(\frac{\|\mathbf{g}\|^2}{n} - 1 \right) + \left(\frac{\|\mathbf{w}\|^2}{n} - \sigma_w^2 \right) + 2 \frac{\|\mathbf{v}\|}{\sqrt{n}} \frac{\mathbf{g}^\top \mathbf{w}}{n}$. Also it is not hard to show under event $\mathcal{A}_1$, $|g(t)| \leq \frac{|t|}{2} \left(1 + \frac{|y|}{\sigma_w} \right)$, $|y| \leq \sqrt{\frac{3}{2\delta}} K$ and $|\Delta| \leq (K^2 + \sigma_w^2 + 2K)\varsigma$ for any $\mathbf{v} \in \mathcal{S}_v(K)$. Therefore,

$$\sup_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} |L(\mathbf{v}) - \tilde{L}(\mathbf{v})| \leq \underbrace{\frac{K^2 + \sigma_w^2 + 2K}{2} \left(1 + \sqrt{\frac{3}{2\delta}} \frac{K}{\sigma_w} \right)}_{:= C_K} \varsigma \quad (102)$$

and thus

$$\left| \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) - \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) \right| \leq C_K \varsigma. \quad (103)$$

Now we are ready to turn back to $L(\mathbf{v})$ to show (88) and (89). Substituting (102) and (103) into (100) and (101) gives

$$\left| \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) - \Psi_* \right| \leq C_K \varsigma + \varepsilon \quad (104)$$

and

$$\begin{aligned} \min_{\mathbf{v} \in \mathcal{B}_{\frac{\varepsilon}{4\sqrt{n\varepsilon/\gamma}}(\hat{\mathbf{v}}) \cap \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) &\geq \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) + 2(\varepsilon - C_K \varsigma). \\ &\stackrel{(a)}{\geq} \Psi_* + \varepsilon - 3C_K \varsigma, \end{aligned} \quad (105)$$

where in (a) we use (104). For any $\varepsilon > 0$, choose $\varsigma \leq \min\{\frac{1}{2}, \frac{\varepsilon}{6C_K}\}$ in (104) and (105). Then (88) and (89) immediately follows, since $\mathbb{P}(\mathcal{A}) \rightarrow 1$. $\blacksquare$

The second lemma is on the asymptotic empirical distribution of the optimal solution of auxiliary problem.

Lemma 11: Let $\psi(\cdot, \cdot)$ be a pseudo-Lipschitz function with constant L and $D_\nu := \{\mathbf{v} \in \mathbb{R}^p : |\mathbb{E}_{\mu_{\mathbf{v}+\beta, \beta}} \psi - \mathbb{E}_{\mu^*} \psi| \geq \nu\}$ as defined in (56). Then for any $K \geq \frac{\sigma_*}{\sqrt{\delta}} + \frac{\theta_{\min}}{4}$, $\nu > 0$ and $\varepsilon \leq \frac{\gamma \nu^2}{192L^2(1+2\delta K^2+32(\mathbb{E}B^2+\sigma_w^2))\delta}$,

$$\mathbb{P}\left(\min_{\mathbf{v} \in D_\nu \cap \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) \leq \Psi_* + 2\varepsilon\right) \rightarrow 0, \quad (106)$$

where Ψ_* is defined in (55).

Proof: For any $\nu > 0$, we will consider the following event:

$$\mathcal{E} = \left\{ |\mathbb{E}_{\mu_{\hat{\mathbf{v}}+\beta, \beta}} \psi - \mathbb{E}_{\mu^*} \psi| \leq \frac{\nu}{2} \right\} \cap \left\{ \frac{1}{p} \|\hat{\mathbf{v}}\|^2, \frac{1}{p} \|\beta\|^2 \leq 4(\mathbb{E}B^2 + \sigma_w^2) \right\},$$

where $\hat{\mathbf{v}} := \eta(\beta + \sigma_* \mathbf{h}) - \beta$ is the same as in (90) and μ^* is the joint measure of $(\eta(B + \sigma_* H), B)$, with $\eta(\cdot) := \eta(\cdot; \mu_{Y_*}, \mu_{\sigma_* \Lambda / \theta_*})$ and $H \sim \mathcal{N}(0, 1)$ independent of $B \sim \mu_B$.

We first show $\mathbb{P}(\mathcal{E}) \rightarrow 1$, as $p \rightarrow \infty$. From (164) in the proof of Lemma 15, we have $W_2(\mu_{\mathbf{h}, \beta}, H \otimes B) \xrightarrow{a.s.} 0$, with $H \sim \mathcal{N}(0, 1)$ and $B \sim \mu_B$. Meanwhile, $(h, b) \mapsto (\eta(b + \sigma_* h), b)$ is a $\sqrt{3 + 2(\sigma_*)^2}$ -Lipschitz continuous mapping. Hence similar as (165), $W_2(\mu_{\hat{\mathbf{v}}+\beta, \beta}, \mu^*) \xrightarrow{a.s.} 0$ and by Theorem 7.12 (iv) in [56], $\mathbb{E}_{\mu_{\hat{\mathbf{v}}+\beta, \beta}} \psi \xrightarrow{a.s.} \mathbb{E}_{\mu^*} \psi$. Similarly, we can show

$$\frac{1}{p} \|\hat{\mathbf{v}}\|^2 \xrightarrow{a.s.} \mathbb{E}[\eta(B + \sigma_* H) - B]^2 < 4(\mathbb{E}B^2 + \sigma_w^2).$$

Also since $\{\beta\}_{p \in \mathbb{N}}$ is a converging sequence, $\frac{1}{p} \|\beta\|^2 \rightarrow \mathbb{E}B^2$. As a result, $\mathbb{P}(\mathcal{E}) \rightarrow 1$ for any $\nu > 0$.

Next we show conditioned on $\mathcal{E}$, it holds that

$$D_\nu \cap \mathcal{B}_{\sqrt{n}K} \subseteq \mathcal{B}_{\sqrt{n\varepsilon_0}}^o(\hat{\mathbf{v}}) \cap \mathcal{B}_{\sqrt{n}K}, \quad (107)$$

where $K \geq \frac{\sigma_*}{\sqrt{\delta}} + \frac{\theta_{\min}}{4}$, $D_\nu := \{v \in \mathbb{R}^p : |\mathbb{E}_{\mu_{v+\beta, \beta}} \psi - \mathbb{E}_{\mu^*} \psi| \geq \nu\}$ and $\varepsilon_0^2 = \frac{\nu^2}{3L^2(1+2\delta K^2+32(\mathbb{E}B^2+\sigma_w^2))\delta}$. Since ψ is L pseudo-Lipschitz, we can get

$$\begin{aligned} |\mathbb{E}_{\mu_{v+\beta, \beta}} \psi - \mathbb{E}_{\mu_{\hat{v}+\beta, \beta}} \psi| &= \left| \frac{1}{p} \sum_{i=1}^p \psi(v_i + \beta_i, \beta_i) - \frac{1}{p} \sum_{i=1}^p \psi(\hat{v}_i + \beta_i, \beta_i) \right| \\ &\leq \frac{L}{p} \sum_{i=1}^p (1 + \sqrt{(v_i + \beta_i)^2 + \beta_i^2} + \sqrt{(\hat{v}_i + \beta_i)^2 + \beta_i^2}) |v_i - \hat{v}_i| \\ &\stackrel{(a)}{\leq} \frac{L}{p} [3(p + 2\|v\|^2 + 2\|\hat{v}\|^2 + 6\|\beta\|^2)]^{1/2} \|v - \hat{v}\|, \end{aligned} \quad (108)$$

where in (a) we use Cauchy-Swartz inequality and $1 + \sqrt{x} + \sqrt{y} \leq \sqrt{3(1+x+y)}$. Meanwhile, conditioned on event $\mathcal{E}$, if $v \in D_\nu$, then

$$|\mathbb{E}_{\mu_{v+\beta, \beta}} \psi - \mathbb{E}_{\mu_{\hat{v}+\beta, \beta}} \psi| \geq |\mathbb{E}_{\mu_{v+\beta, \beta}} \psi - \mathbb{E}_{\mu^*} \psi| - |\mathbb{E}_{\mu_{\hat{v}+\beta, \beta}} \psi - \mathbb{E}_{\mu^*} \psi| \geq \frac{\nu}{2}. \quad (109)$$

Combining (108) and (109), we know conditioned on event $\mathcal{E}$, $\frac{1}{n} \|v - \hat{v}\|^2 \geq \frac{\nu^2}{3L^2(1+2\delta K^2+32(\mathbb{E}B^2+\sigma_w^2))\delta} = \varepsilon_0^2$ for any $v \in D_\nu \cap \mathcal{B}_{\sqrt{n}K}$.

From (107), we have

$$\mathbb{P}\left(\min_{v \in D_\nu \cap \mathcal{B}_{\sqrt{n}K}} L(v) \leq \Psi_* + 2\varepsilon\right) \leq \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}\varepsilon_0}^o(\hat{v}) \cap \mathcal{B}_{\sqrt{n}K}} L(v) \leq \Psi_* + 2\varepsilon\right) + \mathbb{P}(\mathcal{E}^C). \quad (110)$$

On the other hand, from (89) in Lemma 10 we have

$$\mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}\varepsilon_0}^o(\hat{v}) \cap \mathcal{B}_{\sqrt{n}K}} L(v) \leq \Psi_* + 2\varepsilon\right) \rightarrow 0, \quad (111)$$

if $\sqrt{n}\varepsilon_0 \geq 8\sqrt{\frac{n\varepsilon}{\gamma}}$. Therefore, for any $\nu > 0$ we can choose $\varepsilon \leq \frac{\gamma\nu^2}{192L^2(1+2\delta K^2+32(\mathbb{E}B^2+\sigma_w^2))\delta}$ and from (110) and (111) we can get (106). $\blacksquare$

The last lemma in this section shows that the optimal solution of the original problem is bounded with probability converging to 1.

Lemma 12: For $K \geq \frac{\sigma_*}{\sqrt{\delta}} + \frac{\theta_{\min}}{4}$, we have as $p \rightarrow \infty$,

$$\mathbb{P}(\hat{v} \notin \mathcal{B}_{\sqrt{n}K}) \rightarrow 0. \quad (112)$$

Proof: To show $\hat{v} = \underset{v}{\operatorname{argmin}} C(v)$ is bounded with probability approaching 1, we use the following property: for any $a \leq K$,

$$\min_{v \in \mathcal{B}_{\sqrt{n}a}^o \cap \mathcal{B}_{\sqrt{n}K}} C(v) > \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon \Rightarrow \min_{v \in \mathcal{B}_{\sqrt{n}a}^o} C(v) > \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon. \quad (113)$$

One can prove (113) by contradiction. If $\min_{v \in \mathcal{B}_{\sqrt{n}a}^o} C(v) \leq \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon$, it must hold that $\min_{v \in \mathcal{B}_{\sqrt{n}K}^C} C(v) \leq \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon$. Then by convexity of $C(v)$, we can always find $v_0 \in \mathcal{B}_{\sqrt{n}a}^o \cap \mathcal{B}_{\sqrt{n}K}$ such that $C(v_0) \leq \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon$, which leads to a contradiction with $\min_{v \in \mathcal{B}_{\sqrt{n}a}^o \cap \mathcal{B}_{\sqrt{n}K}} C(v) > \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon$. As a result, for any $a \leq K$,

$$\begin{aligned} \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}K}^C} C(v) \leq \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon\right) &\leq \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}a}^o} C(v) \leq \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon\right) \\ &\stackrel{(a)}{\leq} \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}a}^o \cap \mathcal{B}_{\sqrt{n}K}} C(v) \leq \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon\right), \end{aligned} \quad (114)$$

where (a) follows from (113). Now we choose $a \in (\frac{\sigma_*}{\sqrt{\delta}}, K)$. Then in the probability space of auxiliary problem (52), under event $\mathcal{A}_2$ [c.f. (91)] we can get $\mathcal{B}_{\sqrt{na}}^o \subset \mathcal{B}_{\sqrt{n}\Delta_a}^o(\hat{v})$, with $\Delta_a = a - \frac{\sigma_*}{\sqrt{\delta}}$. Therefore, for any $\varepsilon > 0$ we have

$$\begin{aligned} & \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{na}}^o \cap \mathcal{B}_{\sqrt{n}K}} C(v) \leq \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon\right) \\ & \leq \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{na}}^o \cap \mathcal{B}_{\sqrt{n}K}} C(v) \leq \Psi_* + 2\varepsilon\right) + \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) \geq \Psi_* + \varepsilon\right) \\ & \stackrel{(a)}{\leq} \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{na}}^o \cap \mathcal{B}_{\sqrt{n}K}} L(v) \leq \Psi_* + 2\varepsilon\right) + \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}K}} L(v) \geq \Psi_* + \varepsilon\right) \\ & \leq \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}\Delta_a}^o(\hat{v}) \cap \mathcal{B}_{\sqrt{n}K}} L(v) \leq \Psi_* + 2\varepsilon\right) + \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}K}} L(v) \geq \Psi_* + \varepsilon\right) + \mathbb{P}(\mathcal{A}_2^C), \end{aligned} \quad (115)$$

where Ψ_* is defined in (55), in (a) we use (53) and (54). Combining (114) and (115), we have for any $\varepsilon > 0$,

$$\begin{aligned} \mathbb{P}(\hat{v} \notin \mathcal{B}_{\sqrt{n}K}) & \leq \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}K}^C} C(v) \leq \min_{v \in \mathcal{B}_{\sqrt{n}K}} C(v) + \varepsilon\right) \\ & \leq \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}\Delta_a}^o(\hat{v}) \cap \mathcal{B}_{\sqrt{n}K}} L(v) \leq \Psi_* + 2\varepsilon\right) + \mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}K}} L(v) > \Psi_* + \varepsilon\right) + \mathbb{P}(\mathcal{A}_2^C). \end{aligned} \quad (116)$$

It remains to show all the three terms on the RHS of (116) converge to 0 for some $\varepsilon > 0$. From (89), we know for $\varepsilon > 0$ if we choose an a such that $\sqrt{n}\Delta_a \geq 8\sqrt{\frac{n\varepsilon}{\gamma}}$, where $\gamma = \frac{\theta_{\min}\sigma_w^2}{4(K^2 + \sigma_w^2)^{3/2}}$, then

$$\mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}\Delta_a}^o(\hat{v}) \cap \mathcal{B}_{\sqrt{n}K}} L(v) \leq \Psi_* + 2\varepsilon\right) \rightarrow 0. \quad (117)$$

Clearly, such a always exists if $\varepsilon \in (0, \frac{\gamma^2\theta_{\min}^2}{1024})$ since $K \geq \frac{\sigma_*}{\sqrt{\delta}} + \frac{\theta_{\min}}{4}$. On the other hand, in the proof of Lemma 10 we show $\mathbb{P}(\mathcal{A}_2^C) \rightarrow 0$ and from (88) we have for any $\varepsilon > 0$, $\mathbb{P}\left(\min_{v \in \mathcal{B}_{\sqrt{n}K}} L(v) > \Psi_* + \varepsilon\right) \rightarrow 0$. Substituting these results back to (116), we reach (112). $\blacksquare$

E. Proof of Lemma 3

To obtain the probabilistic upper bound for $R_0^{(p)}$, we can approximate indicator function $\mathbb{I}_{x=0}$ by a series of envelope functions:

$$\psi_h(x) = \begin{cases} 1 - h^{-1}x & 0 \leq x < h, \\ 1 + h^{-1}x & -h \leq x < 0, \\ 0 & |x| \geq h, \end{cases} \quad (118)$$

where $h > 0$. We can see $\psi_h(x)$ is an upper bound of $\mathbb{I}_{x=0}$ and satisfies $0 \leq \psi_h(x) - \mathbb{I}_{x=0} \leq \mathbb{I}_{0 < |x| < h}$, so

$$\begin{aligned} R_0^{(p)} - \mathbb{P}[\eta(Y_*) = 0] & = \mathbb{E}_{\mu_{\hat{\beta}}} \mathbb{I}_{x=0} - \mathbb{E}_{\mu_{\hat{B}}} \mathbb{I}_{x=0} \\ & \leq \mathbb{E}_{\mu_{\hat{\beta}}} \psi_h(x) - \mathbb{E}_{\mu_{\hat{B}}} \psi_h(x) + \mathbb{E}_{\mu_{\hat{B}}} \psi_h(x) - \mathbb{E}_{\mu_{\hat{B}}} \mathbb{I}_{x=0} \\ & \leq |\mathbb{E}_{\mu_{\hat{\beta}}} \psi_h(x) - \mathbb{E}_{\mu_{\hat{B}}} \psi_h(x)| + \mathbb{P}[|\eta(Y_*)| \in (0, h)], \end{aligned} \quad (119)$$

where $\mu_{\hat{B}}$ denotes the distribution of $\eta(Y_*)$. Moreover, $\psi_h(x)$ is h^{-1} -Lipschitz (and hence pseudo-Lipschitz by definition), so (12) can now be applied, which gives us $|\mathbb{E}_{\mu_{\hat{\beta}}} \psi_h(x) - \mathbb{E}_{\mu_{\hat{B}}} \psi_h(x)| \xrightarrow{\mathbb{P}} 0$ for any fixed $h > 0$.

Meanwhile, by continuity of probability, we have $\lim_{h \rightarrow 0} \mathbb{P}(|\eta(Y_*)| \in (0, h)) = 0$. As a result, on both sides of (119) taking $p \rightarrow \infty$ and then $h \rightarrow 0$, we get for any $\varepsilon > 0$,

$$R_0^{(p)} \leq \mathbb{P}(\eta(Y_*) = 0) + \varepsilon \quad (120)$$

with probability approaching 1 as $p \rightarrow \infty$.

F. Proof of Lemma 4

If $\mathbb{P}(\eta(Y_*) = 0) = 0$, then since $R_0^{(p)} \geq 0$, (61) trivially holds. Thus, it only remains to address the case when $\mathbb{P}(\eta(Y_*) = 0) > 0$. Towards this end, we will utilize Lemma 5. To apply Lemma 5, we need to verify for sufficiently large $k \in [p]$, $|\hat{s}|_{(1:k)} \prec_S \lambda_{1:k}$ with probability approaching 1. For any $p, k, \ell \in \mathbb{Z}^+$, with $0 \leq k < \ell \leq p$,

$$\begin{aligned} \sum_{i=k+1}^{\ell} |\hat{s}|_{(i)} - \sum_{i=k+1}^{\ell} \lambda_i &= \int_{k/p}^{\ell/p} F_{|\hat{s}|}^{-1}(u) du - \int_{k/p}^{\ell/p} F_{|\lambda|}^{-1}(u) du \\ &\leq \int_{k/p}^{\ell/p} F_{|\hat{s}|}^{-1}(u) du - \int_{k/p}^{\ell/p} F_{\Lambda}^{-1}(u) du + \left| \int_{k/p}^{\ell/p} F_{|\hat{s}|}^{-1}(u) du - \int_{k/p}^{\ell/p} F_{|\hat{S}|}^{-1}(u) du \right| \\ &\quad + \left| \int_{k/p}^{\ell/p} F_{\Lambda}^{-1}(u) du - \int_{k/p}^{\ell/p} F_{|\lambda|}^{-1}(u) du \right| \\ &\leq \int_{k/p}^{\ell/p} F_{|\hat{S}|}^{-1}(u) du - \int_{k/p}^{\ell/p} F_{\Lambda}^{-1}(u) du + W_2(\mu_{\hat{s}}, \mu_{\hat{S}}) + W_2(\mu_{\lambda}, \mu_{\Lambda}), \end{aligned} \quad (121)$$

where in the last step, we use (86). Here, $\mu_{\hat{S}}$ is the law of $\tau_*^{-1}[Y_* - \eta(Y_*)]$. By condition (R.1), we know for any $\varepsilon \in (0, q_0^*]$ (recall that $q_0^* = \mathbb{P}(\eta(Y_*) = 0)$), there exists $\varsigma > 0$ such that

$$\begin{aligned} &\max_{t \in [0, q_0^* - \varepsilon]} \int_t^{q_0^*} F_{|\hat{s}|}^{-1}(u) du - \int_t^{q_0^*} F_{\Lambda}^{-1}(u) du \\ &\stackrel{(a)}{=} \tau_*^{-1} \max_{t \in [0, q_0^* - \varepsilon]} \int_t^{q_0^*} F_{|Y_*|}^{-1}(u) du - \int_t^{q_0^*} F_{\tau_* \Lambda}^{-1}(u) du \\ &\stackrel{(b)}{\leq} -\tau_*^{-1} \varsigma, \end{aligned} \quad (122)$$

where to reach (a), we use $\hat{S} \stackrel{\text{law}}{=} \tau_*^{-1}[Y_* - \eta(Y_*)]$ and the fact that $\eta(y) = 0$ for $|y| \leq F_{|Y_*|}^{-1}(q_0^*)$ and (b) is due to (R.1) and the fact that $t \mapsto \int_t^{q_0^*} F_{|Y_*|}^{-1}(u) du$ and $t \mapsto \int_t^{q_0^*} F_{\tau_* \Lambda}^{-1}(u) du$ are both continuous. For any fixed $\varepsilon \in (0, q_0^*]$, $\lfloor p(q_0^* - \varepsilon) \rfloor \leq \lfloor pq_0^* \rfloor - 1$ for large enough p . Then substituting (122) into (121) we can get for large enough p and any $0 \leq k \leq \lfloor p(q_0^* - \varepsilon) \rfloor$,

$$\begin{aligned} &\sum_{i=k+1}^{\lfloor pq_0^* \rfloor} |\hat{s}|_{(i)} - \sum_{i=k+1}^{\lfloor pq_0^* \rfloor} \lambda_i \\ &\stackrel{(a)}{\leq} \int_{k/p}^{q_0^*} F_{|\hat{S}|}^{-1}(u) du - \int_{k/p}^{q_0^*} F_{\Lambda}^{-1}(u) du - \left[\int_{\lfloor pq_0^* \rfloor / p}^{q_0^*} F_{|\hat{s}|}^{-1}(u) du - \int_{\lfloor pq_0^* \rfloor / p}^{q_0^*} F_{\Lambda}^{-1}(u) du \right] \\ &\quad + W_2(\mu_{\hat{s}}, \mu_{\hat{S}}) + W_2(\mu_{\lambda}, \mu_{\Lambda}) \\ &\stackrel{(b)}{\leq} -\tau_*^{-1} \varsigma + \tau_*^{-1} \left(\int_{\lfloor pq_0^* \rfloor / p}^{q_0^*} F_{|Y_*|}^{-1}(u) du + \int_{\lfloor pq_0^* \rfloor / p}^{q_0^*} F_{\tau_* \Lambda}^{-1}(u) du \right) + W_2(\mu_{\hat{s}}, \mu_{\hat{S}}) + W_2(\mu_{\lambda}, \mu_{\Lambda}), \end{aligned} \quad (123)$$

where (a) follows from (121) and (b) follows from (122). We now show the last four terms in (123) vanish as $p \rightarrow \infty$: $\int_{\lfloor pq_0^* \rfloor / p}^{q_0^*} F_{|Y_*|}^{-1}(u) du$ and $\int_{\lfloor pq_0^* \rfloor / p}^{q_0^*} F_{\tau_* \Lambda}^{-1}(u) du$ converges to 0 by DCT; $W_2(\mu_{\hat{s}}, \mu_S) \xrightarrow{\mathbb{P}} 0$ is proved in

Proposition 6; $W_2(\mu_\lambda, \mu_\Lambda) \rightarrow 0$ since $\{\lambda\}_{p \in \mathbb{Z}^+}$ is a converging sequence. As a result, for any fixed $\varepsilon \in (0, q_0^*]$, there exists $\varepsilon > 0$ such that

$$\max_{0 \leq k \leq \lfloor p(q_0^* - \varepsilon) \rfloor} \sum_{i=k+1}^{\lfloor pq_0^* \rfloor} |\hat{s}|_{(i)} - \sum_{i=k+1}^{\lfloor pq_0^* \rfloor} \lambda_i \leq -\varsigma, \quad (124)$$

with probability approaching 1, as $p \rightarrow \infty$. Now we show conditioned on (124), there exists $k_0 \in [\lfloor p(q_0^* - \varepsilon) \rfloor, \lfloor pq_0^* \rfloor - 1]$ such that

$$\max_{0 \leq k \leq k_0} \sum_{i=k+1}^{k_0+1} |\hat{s}|_{(i)} - \sum_{i=k+1}^{k_0+1} \lambda_i < 0. \quad (125)$$

In other words, $|\hat{s}|_{(1:k_0+1)} \prec_S \lambda_{1:k_0+1}$. Such k_0 can be retrieved as follows:

Step 0. Let k_s denote the candidate for k_0 and initialize $k_s = \lfloor p(q_0^* - \varepsilon) \rfloor$. From (124) we know at the initial step,

$$\max_{0 \leq k \leq k_s} \sum_{i=k+1}^{\lfloor pq_0^* \rfloor} |\hat{s}|_{(i)} - \sum_{i=k+1}^{\lfloor pq_0^* \rfloor} \lambda_i \leq -\varsigma. \quad (126)$$

Step 1. If $k_s + 1 = \lfloor pq_0^* \rfloor$, then we output $k_0 = k_s$; otherwise we go to step 2.

Step 2. If $\sum_{i=k_s+2}^{\lfloor pq_0^* \rfloor} |\hat{s}|_{(i)} - \sum_{i=k_s+2}^{\lfloor pq_0^* \rfloor} \lambda_i > -\frac{\varsigma}{2}$, then together with (126) we get $\max_{0 \leq k \leq k_s} \sum_{i=k+1}^{k_s+1} |\hat{s}|_{(i)} - \sum_{i=k+1}^{k_s+1} \lambda_i \leq -\frac{\varsigma}{2}$. Hence, we output $k_0 = k_s$; otherwise, we update k_s and ς as: $k_s \leftarrow k_s + 1$, $\varsigma \leftarrow \frac{\varsigma}{2}$ and return back to step 1. Clearly, (126) still holds under the updated k_s and ς .

Then from Lemma 5, (125) implies that $|\hat{\beta}|_{(1:k_0+1)} = 0$ and thus $R_0^{(p)} \geq q_0^* - \varepsilon$ since $k_0 \geq p(q_0^* - \varepsilon) - 1$ by construction. Summing up, if $q_0^* > 0$ then for any $\varepsilon \in (0, q_0^*]$, $R_0^{(p)} \geq q_0^* - \varepsilon = \mathbb{P}(\eta(Y_*) = 0) - \varepsilon$ with probability approaching 1 as $p \rightarrow \infty$. Therefore (61) is verified.

G. Proof of Lemma 5

The key is to establish the following result: for any p -dimensional vectors $\mathbf{a}$ and λ with $0 \leq \lambda_1 \leq \dots \leq \lambda_p$, if $|\mathbf{a} - \text{Prox}_\lambda(\mathbf{a})|_{(1:k)} \prec_S \lambda_{1:k}$ for some $k \in [p]$, then $|\text{Prox}_\lambda(\mathbf{a})|_{(1:k)} = \mathbf{0}$. To prove this, it suffices to show $|\text{Prox}_\lambda(\mathbf{a})|_{(k)} = 0$. Assume $|\text{Prox}_\lambda(\mathbf{a})|_{(k)} \neq 0$. Define the index set $I_k \stackrel{\text{def}}{=} \{i \mid |\text{Prox}_\lambda(\mathbf{a})|_{(i)} = |\text{Prox}_\lambda(\mathbf{a})|_{(k)}\}$ and denote $\ell := \min I_k$ and $m := \max I_k$. According to the formula for ∂J_λ [47, Fact V.3], we have: if $|\text{Prox}_\lambda(\mathbf{a})|_{(k)} \neq 0$, then for any $\mathbf{g} \in \partial J_\lambda(\text{Prox}_\lambda(\mathbf{a}))$ it holds that $|\mathbf{g}|_{(\ell:m)} \prec \lambda_{\ell:m}$ and $\sum_{i=\ell}^m |\mathbf{g}|_{(i)} = \sum_{i=\ell}^m \lambda_i$. On the other hand, by the first order optimality condition, $\mathbf{a} - \text{Prox}_\lambda(\mathbf{a}) \in \partial J_\lambda(\text{Prox}_\lambda(\mathbf{a}))$. Hence, $|\mathbf{a} - \text{Prox}_\lambda(\mathbf{a})|_{(\ell:m)} \prec \lambda_{\ell:m}$, i.e., for any $\ell \leq q \leq m$,

$$\sum_{i=q}^m |\mathbf{a} - \text{Prox}_\lambda(\mathbf{a})|_{(i)} \leq \sum_{i=q}^m \lambda_i \quad (127)$$

and also

$$\sum_{i=\ell}^m |\mathbf{a} - \text{Prox}_\lambda(\mathbf{a})|_{(i)} = \sum_{i=\ell}^m \lambda_i. \quad (128)$$

Therefore,

$$\begin{aligned}
\sum_{i=\ell}^m |\mathbf{a} - \text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})|_{(i)} &= \sum_{i=\ell}^k |\mathbf{a} - \text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})|_{(i)} + \sum_{i=k+1}^m |\mathbf{a} - \text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})|_{(i)} \\
&< \sum_{i=\ell}^k \lambda_i + \sum_{i=k+1}^m \lambda_i \\
&= \sum_{i=\ell}^m \lambda_i,
\end{aligned} \tag{129}$$

where the inequality follows from the condition that $|\mathbf{a} - \text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})|_{(1:k)} \prec_S \boldsymbol{\lambda}_{1:k}$ and (127). Clearly, (129) contradicts (128). Therefore, $|\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})|_{(k)} = 0$ and thus $|\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})|_{(1:k)} = \mathbf{0}_k$.

Now we are ready to prove Lemma 5. To apply the above result, we just need to express $\hat{\boldsymbol{\beta}}$ as

$$\hat{\boldsymbol{\beta}} = \text{Prox}_{\boldsymbol{\lambda}}(\hat{\boldsymbol{\beta}} + \hat{\mathbf{s}}), \tag{130}$$

using the first order optimality condition of (2). Therefore, we can let $\mathbf{a} = \hat{\boldsymbol{\beta}} + \hat{\mathbf{s}}$ and thus $\hat{\mathbf{s}} = \mathbf{a} - \text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})$. The desired result then immediately follows.

H. Proof of Lemma 6

Similar as proof of Lemma 3, the idea is again approximating the indicator function $\mathbb{I}_{x \neq 0, y=0}$ by some Lipschitz continuous functions. Here we use:

$$\phi_h(x, y) = [1 - \psi_h(x)]\psi_h(y),$$

where $\psi_h(x)$ is defined in (118). We have

$$|\phi_h(x, y) - \mathbb{I}_{x \neq 0, y=0}| \leq \mathbb{I}_{0 < |x| < h} + \mathbb{I}_{0 < |y| < h}. \tag{131}$$

Therefore, for any $h > 0$,

$$\begin{aligned}
&|V^{(p)} - \mathbb{P}[\eta(Y_*) \neq 0, B = 0]| \\
&= |\mathbb{E}_{\mu_{\hat{\boldsymbol{\beta}}, \boldsymbol{\beta}}} \mathbb{I}_{x \neq 0, y=0} - \mathbb{E}_{\mu_{\hat{B}, B}} \mathbb{I}_{x \neq 0, y=0}| \\
&\leq \mathbb{E}_{\mu_{\hat{\boldsymbol{\beta}}, \boldsymbol{\beta}}} |\mathbb{I}_{x \neq 0, y=0} - \phi_h(x, y)| + |\mathbb{E}_{\mu_{\hat{\boldsymbol{\beta}}, \boldsymbol{\beta}}} \phi_h(x, y) - \mathbb{E}_{\mu_{\hat{B}, B}} \phi_h(x, y)| \\
&\quad + \mathbb{E}_{\mu_{\hat{B}, B}} |\mathbb{I}_{x \neq 0, y=0} - \phi_h(x, y)| \\
&\leq \frac{1}{p} \sum_{i=1}^p (\mathbb{I}_{|\hat{\beta}_i| < h} - \mathbb{I}_{\hat{\beta}_i=0}) + \frac{1}{p} \sum_{i=1}^p (\mathbb{I}_{|\beta_i| < h} - \mathbb{I}_{\beta_i=0}) + \mathbb{P}(0 < |\eta(Y_*)| < h) \\
&\quad + \mathbb{P}(0 < |B| < h) + |\mathbb{E}_{\mu_{\hat{\boldsymbol{\beta}}, \boldsymbol{\beta}}} \phi_h(x, y) - \mathbb{E}_{\mu_{\hat{B}, B}} \phi_h(x, y)|,
\end{aligned} \tag{132}$$

where $\mu_{\hat{B}, B}$ denotes the joint distribution of $(\eta(Y_*), B)$ and in the last step we use (131) and $\mathbb{I}_{0 < |x| < h} = \mathbb{I}_{|x| < h} - \mathbb{I}_{x=0}$. Let us compute the limit of each term on the RHS of (132). The last term converges in probability to zero due to (12). By continuity of probability, $\mathbb{P}(0 < |\eta(Y_*)| < h)$ and $\mathbb{P}(0 < |B| < h)$ converge to 0 as $h \rightarrow 0$. Following similar steps leading to (120), we can get $\frac{1}{p} \sum_{i=1}^p \mathbb{I}_{|\hat{\beta}_i| < h} \xrightarrow{\mathbb{P}} \mathbb{P}(|\eta(Y_*)| < h)$ and $\frac{1}{p} \sum_{i=1}^p \mathbb{I}_{|\beta_i| < h} \xrightarrow{\mathbb{P}} \mathbb{P}(|B| < h)$ if

h satisfies $\mathbb{P}(|\eta(Y_*)| = h) = \mathbb{P}(B = h) = 0$. Therefore, for such $h > 0$, (132) yields the following: for any $\varepsilon > 0$ there exists p_0 such that when $p > p_0$,

$$\begin{aligned} |V^{(p)} - \mathbb{P}[\eta(Y_*) \neq 0, B = 0]| &\leq [\mathbb{P}(|B| = 0) + \mathbb{P}(|\eta(Y_*)| = 0)] - (r_0^{(p)} + R_0^{(p)}) \\ &\quad + 2\mathbb{P}(0 < |B| < h) + 2\mathbb{P}(0 < |\eta(Y_*)| < h) + \frac{\varepsilon}{2}. \\ &\leq 2\mathbb{P}(0 < |B| < h) + 2\mathbb{P}(0 < |\eta(Y_*)| < h) + |\mathbb{P}(|\eta(Y_*)| = 0) - R_0^{(p)}| + \varepsilon, \end{aligned} \quad (133)$$

where in the last step we use Assumption (A.2). Then taking $h \rightarrow 0$ along a sequence $\{h_i\}_{i \in \mathbb{Z}^+}$ with $\mathbb{P}(|\eta(Y_*)| = h_i) = \mathbb{P}(B = h_i) = 0$ in (133), we get for any $\varepsilon > 0$,

$$|V^{(p)} - \mathbb{P}[\eta(Y_*) \neq 0, B = 0]| \leq |\mathbb{P}(\eta(Y_*) = 0) - R_0^{(p)}| + \varepsilon, \quad (134)$$

with probability approaching 1 as $p \rightarrow \infty$.

I. Asymptotic Properties of $\hat{s}$

In this section, we study the limiting properties of the following vector: $\hat{s} = \mathbf{A}^\top(\mathbf{w} - \mathbf{A}\hat{\mathbf{v}})$, where $\hat{\mathbf{v}} = \text{argmin}_{\mathbf{v}} C(\mathbf{v})$. Recall that $C(\mathbf{v})$ is the objective function of primary problem defined in (51):

$$C(\mathbf{v}) = \frac{1}{n} \left[\frac{1}{2} \|\mathbf{A}\mathbf{v} - \mathbf{w}\|^2 + J_\lambda(\mathbf{v} + \boldsymbol{\beta}) \right]$$

and $\hat{\mathbf{v}}$ is the optimal solution. The main goal is to prove Proposition 6, which characterizes the limiting empirical distribution of $(\hat{s}, \boldsymbol{\beta})$. We will follow the proof strategy in [50, Appendix E].

Proposition 6: Under the same setting as Theorem 1, define $\mu_{\hat{s}, B}$ as the joint measure of $(\tau_*^{-1}[Y - \eta(Y; \mu_{Y_*}, \mu_{\tau_*\Lambda})], B)$. It holds that $W_2(\mu_{\hat{s}, \boldsymbol{\beta}}, \mu_{\hat{s}, B}) \xrightarrow{\mathbb{P}} 0$.

Proof: The first step is to obtain an alternative representation of $\hat{s}$. Consider the event $E = \{\hat{\mathbf{v}} \in \mathcal{B}_{\sqrt{n_K}}\}$, for some $K \geq \frac{\sigma_*}{\sqrt{\delta}} + \frac{\theta_{\min}}{4}$, where $\theta_{\min}$ is given in Lemma 14. It is shown in Lemma 12 that $\mathbb{P}(E) \rightarrow 1$ as $p \rightarrow \infty$. Under event E , we have

$$\begin{aligned} \min_{\mathbf{v} \in \mathbb{R}^p} C(\mathbf{v}) &= \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n_K}}} \frac{1}{n} \left[\frac{1}{2} \|\mathbf{A}\mathbf{v} - \mathbf{w}\|^2 + J_\lambda(\mathbf{v} + \boldsymbol{\beta}) \right] \\ &= \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n_K}}} \max_{\mathbf{s} \in C_\lambda} \frac{1}{n} \left[\frac{1}{2} \|\mathbf{A}\mathbf{v} - \mathbf{w}\|^2 + \mathbf{s}^\top(\mathbf{v} + \boldsymbol{\beta}) \right] \\ &= \max_{\mathbf{s} \in C_\lambda} \underbrace{\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n_K}}} \frac{1}{n} \left[\frac{1}{2} \|\mathbf{A}\mathbf{v} - \mathbf{w}\|^2 + \mathbf{s}^\top(\mathbf{v} + \boldsymbol{\beta}) \right]}_{\stackrel{\text{def}}{=} S(\mathbf{s})}, \end{aligned} \quad (135)$$

where C_λ is defined in (187) and the last step follows from Sion's minimax theorem [59]. By the first order optimality condition of , we know $\hat{s} \in \partial J_\lambda(\hat{\mathbf{v}} + \boldsymbol{\beta})$. On the other hand, it is not hard to show for any $\mathbf{x} \in \mathbb{R}^p$ and

$s \in \partial J_\lambda(x)$, it holds that $s \in C_\lambda$ and $J_\lambda(x) = s^\top x$. Therefore, $\hat{s} \in C_\lambda$ and $\hat{s}^\top(\hat{v} + \beta) = J_\lambda(\hat{v} + \beta)$. Then

$$\begin{aligned} \mathcal{S}(\hat{s}) &= \min_{v \in \mathcal{B}_{\sqrt{n}K}} \frac{1}{n} \left[\frac{1}{2} \|Av - w\|^2 + \hat{s}^\top(v + \beta) \right] \\ &\stackrel{(a)}{=} \frac{1}{n} \left[\frac{1}{2} \|A\hat{v} - w\|^2 + \hat{s}^\top(\hat{v} + \beta) \right] \\ &= \frac{1}{n} \left[\frac{1}{2} \|A\hat{v} - w\|^2 + J_\lambda(\hat{v} + \beta) \right] \\ &\stackrel{(b)}{=} \max_{s \in C_\lambda} \mathcal{S}(s), \end{aligned}$$

where (a) follows from first order optimality condition and the fact that $\hat{v} \in \mathcal{B}_{\sqrt{n}K}$ under event E and (b) follows from (135) and $\hat{v} \in \min_{v \in \mathbb{R}^p} C(v)$. This implies that under event E , which happens with probability approaching 1, $\hat{s} \in \arg\max_{s \in C_\lambda} \mathcal{S}(s)$. Therefore, in order to study the limiting behavior of $\hat{s}$, we can instead study $\arg\max_{s \in C_\lambda} \mathcal{S}(s)$.

The analysis of $\arg\max_{s \in C_\lambda} \mathcal{S}(s)$ can be carried out based on CGMT framework. First, similar as Proposition E.1 in [50], we can get for any closed set D and $t \in \mathbb{R}$,

$$\mathbb{P}(\max_{s \in D} \mathcal{S}(s) \geq t) \leq 2\mathbb{P}(\max_{s \in D} S(s) \geq t) \quad (136)$$

and if D is also convex,

$$\mathbb{P}(\max_{s \in D} \mathcal{S}(s) \leq t) \leq 2\mathbb{P}(\max_{s \in D} S(s) \leq t). \quad (137)$$

Here

$$S(s) = \min_{v \in \mathcal{B}_{\sqrt{n}K}} \frac{1}{2} \left(\sqrt{\frac{\|v\|^2}{n} \frac{\|g\|^2}{n} + \frac{\|w\|^2}{n} + 2 \frac{\|v\|}{\sqrt{n}} \frac{g^\top w}{n} - \frac{h^\top v}{n}} \right)^2_+ + \frac{s^\top(v + \beta)}{n}. \quad (138)$$

Then for any $\varepsilon, \nu > 0$ and $D_\nu \stackrel{\text{def}}{=} \left\{ s : W_2(\mu_{s, \beta}, \mu_{\hat{s}, B}) \geq \nu \right\} \cap C_\lambda$, we have

$$\begin{aligned} \mathbb{P}(\max_{s \in D_\nu} \mathcal{S}(s) \geq \max_{s \in C_\lambda} \mathcal{S}(s) - \varepsilon) &\leq \mathbb{P}(\max_{s \in D_\nu} \mathcal{S}(s) \geq \Psi^* - 2\varepsilon) + \mathbb{P}(\max_{s \in C_\lambda} \mathcal{S}(s) < \Psi^* - \varepsilon) \\ &\stackrel{(a)}{\leq} 2\mathbb{P}(\max_{s \in D_\nu} S(s) \geq \Psi^* - 2\varepsilon) + 2\mathbb{P}(\max_{s \in C_\lambda} S(s) < \Psi^* - \varepsilon) \end{aligned} \quad (139)$$

where Ψ^* is defined in (55) and step (a) follows from (136) and (137). Then combining (139) with Lemma 13, we get for any $\nu > 0$, there exists $\varepsilon > 0$ such that $\mathbb{P}(\max_{s \in D_\nu} \mathcal{S}(s) \geq \max_{s \in C_\lambda} \mathcal{S}(s) - \varepsilon) \rightarrow 0$. Therefore, for any $\nu > 0$ and $\varepsilon > 0$,

$$\begin{aligned} \mathbb{P}(W_2(\mu_{\hat{s}, \beta}, \mu_{\hat{s}, B}) \geq \nu) &= \mathbb{P}(\hat{s} \in D_\nu) \\ &\leq \mathbb{P}(\hat{s} \in D_\nu \cap E) + \mathbb{P}(E^C) \\ &\leq \mathbb{P}\left[\hat{s} \in D_\nu \text{ and } \mathcal{S}(\hat{s}) = \max_{s \in C_\lambda} \mathcal{S}(s)\right] + \mathbb{P}(E^C) \\ &\leq \mathbb{P}\left[\max_{s \in D_\nu} \mathcal{S}(s) = \max_{s \in C_\lambda} \mathcal{S}(s)\right] + \mathbb{P}(E^C) \\ &\leq \mathbb{P}(\max_{s \in D_\nu} \mathcal{S}(s) \geq \max_{s \in C_\lambda} \mathcal{S}(s) - \varepsilon) + \mathbb{P}(E^C). \end{aligned} \quad (140)$$

By the discussion above, the RHS of (140) converges to 0 for some $\varepsilon > 0$. Since, the LHS of (140) does not depend on ε , this concludes the proof. $\blacksquare$

Lemma 13: Under the same setting as Proposition 6, as $p \rightarrow \infty$,

$$\max_{s \in C_\lambda} S(s) \xrightarrow{\mathbb{P}} \Psi^*, \quad (141)$$

where $S(\mathbf{s})$ is defined in (138) and Ψ^* is defined in (55). Denote $D_\nu := \left\{ \mathbf{s} : W_2(\mu_{\mathbf{s}, \beta}, \mu_{\hat{S}, B}) \geq \nu \right\} \cap C_\lambda$. For any $\nu > 0$, there exists $\varepsilon > 0$ such that

$$\mathbb{P}\left(\max_{\mathbf{s} \in D_\nu} S(\mathbf{s}) \geq \Psi^* - \varepsilon\right) \rightarrow 0. \quad (142)$$

Proof: For the similar reason as introducing $\tilde{L}(\mathbf{v})$ when analyzing $L(\mathbf{v})$ [c.f. (87) and (92) in the proof of Lemma 10], we consider the following approximation of $S(\mathbf{s})$:

$$\tilde{S}(\mathbf{s}) = \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \frac{1}{2} \left(\sqrt{\frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right)_+^2 + \frac{\mathbf{s}^\top (\mathbf{v} + \beta)}{n}. \quad (143)$$

Note that $S(\mathbf{s}) = \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} L(\mathbf{v}) - \Delta(\mathbf{v})$ and $\tilde{S}(\mathbf{s}) = \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) - \Delta(\mathbf{v})$, where $\Delta(\mathbf{v}) = \frac{J_\lambda(\mathbf{v} + \beta) - \mathbf{s}^\top (\mathbf{v} + \beta)}{n}$. Therefore, by (102)

$$\sup_{\mathbf{s} \in \mathbb{R}^p} |S(\mathbf{s}) - \tilde{S}(\mathbf{s})| \leq \sup_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} |L(\mathbf{v}) - \tilde{L}(\mathbf{v})| \xrightarrow{\mathbb{P}} 0. \quad (144)$$

On the other hand, similar as (135) we have

$$\begin{aligned} \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) &= \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \frac{1}{2} \left(\sqrt{\frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right)_+^2 + \frac{J_\lambda(\mathbf{v} + \beta)}{n} \\ &= \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \max_{\mathbf{s} \in C_\lambda} \frac{1}{2} \left(\sqrt{\frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right)_+^2 + \frac{\mathbf{s}^\top (\mathbf{v} + \beta)}{n} \\ &= \max_{\mathbf{s} \in C_\lambda} \min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \frac{1}{2} \left(\sqrt{\frac{\|\mathbf{v}\|^2}{n} + \sigma_w^2} - \frac{\mathbf{h}^\top \mathbf{v}}{n} \right)_+^2 + \frac{\mathbf{s}^\top (\mathbf{v} + \beta)}{n} \\ &= \max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}), \end{aligned} \quad (145)$$

where C_λ is defined in (187). Using successively (144) and (145), we have for any $\varepsilon > 0$,

$$\begin{aligned} \mathbb{P}\left(\max_{\mathbf{s} \in C_\lambda} S(\mathbf{s}) < \Psi^* - \varepsilon\right) &\leq \mathbb{P}\left(\max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) < \Psi^* - \frac{\varepsilon}{2}\right) + \delta_p \\ &= \mathbb{P}\left(\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) < \Psi^* - \frac{\varepsilon}{2}\right) + \delta_p, \end{aligned} \quad (146)$$

where $\delta_p \rightarrow 0$ as $p \rightarrow \infty$. Similarly on the other direction, we can also get for any $\varepsilon > 0$, there is some $\delta'_p \rightarrow 0$ such that

$$\mathbb{P}\left(\max_{\mathbf{s} \in C_\lambda} S(\mathbf{s}) \geq \Psi^* + \varepsilon\right) \leq \mathbb{P}\left(\min_{\mathbf{v} \in \mathcal{B}_{\sqrt{n}K}} \tilde{L}(\mathbf{v}) \geq \Psi^* + \frac{\varepsilon}{2}\right) + \delta'_p \quad (147)$$

Then combining (146) and (147) with (100), we get $\max_{\mathbf{s} \in C_\lambda} S(\mathbf{s}) \xrightarrow{\mathbb{P}} \Psi^*$ and from (144) we get $\max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) \xrightarrow{\mathbb{P}} \Psi^*$.

Next, we show (142). First, the following bound holds

$$\begin{aligned} \mathbb{P}\left(\max_{\mathbf{s} \in D_\nu} S(\mathbf{s}) \geq \Psi^* - \varepsilon\right) &\leq \mathbb{P}\left(\max_{\mathbf{s} \in D_\nu} \tilde{S}(\mathbf{s}) \geq \Psi^* - 2\varepsilon\right) + \mathbb{P}\left(\sup_{\mathbf{s} \in D_\nu} |S(\mathbf{s}) - \tilde{S}(\mathbf{s})| \geq \varepsilon\right) \\ &\leq \mathbb{P}\left(\max_{\mathbf{s} \in D_\nu} \tilde{S}(\mathbf{s}) \geq \max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) - 3\varepsilon\right) + \mathbb{P}\left(\max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) \geq \Psi^* + \varepsilon\right) \\ &\quad + \mathbb{P}\left(\sup_{\mathbf{s} \in D_\nu} |S(\mathbf{s}) - \tilde{S}(\mathbf{s})| \geq \varepsilon\right). \end{aligned} \quad (148)$$

Recall that we have already shown that the last two terms on the RHS of (148) vanish as $p \rightarrow \infty$. Therefore, it remains to show the first term also converges to 0. The main step is to establish there exist $\varsigma_{\max}, \gamma > 0$ such that for any $\varsigma \in (0, \varsigma_{\max})$,

$$\mathbb{P}\left(\max_{\mathbf{s} \in C_\lambda \cap \mathcal{B}_{\sqrt{p}\varsigma}^\circ(\hat{\mathbf{s}})} \tilde{S}(\mathbf{s}) \geq \max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) - \gamma\varsigma\right) \rightarrow 0, \quad (149)$$

where

$$\begin{aligned}\hat{s} &\stackrel{\text{def}}{=} \frac{1}{\tau_*} [(\beta + \sigma_* \mathbf{h}) - (\beta + \hat{v})] \\ &= \frac{1}{\tau_*} [(\beta + \sigma_* \mathbf{h}) - \eta(\beta + \sigma_* \mathbf{h}; \mu_{Y_*}, \mu_{\tau_* \Lambda})]\end{aligned}$$

and $\hat{v}$ is defined in (90). The convergence in (149) can be proved in exactly the same way as Theorem E.7 in [50], which deals with the case of LASSO. For simplicity, we do not re-present the proof details here. Here C_λ (the unit ball of the dual norm of J_λ) plays the role of the set $\{\mathbf{x} \mid \|\mathbf{x}\|_\infty \leq \lambda\}$ in [50], which is the unit ball of the dual norm of $\lambda \|\cdot\|_1$. Now consider the event $E = \{W_2(\mu_{\hat{s}, \beta}, \mu_{\hat{s}, B}) \leq \frac{\nu}{2}\}$. Conditioned on E , it holds that for $\mathbf{s} \in D_\nu$,

$$\begin{aligned}\frac{1}{p} \|\mathbf{s} - \hat{s}\|^2 &\geq W_2(\mu_{\mathbf{s}, \beta}, \mu_{\hat{s}, \beta})^2 \\ &\geq [W_2(\mu_{\mathbf{s}, \beta}, \mu_{\hat{s}, B}) - W_2(\mu_{\hat{s}, \beta}, \mu_{\hat{s}, B})]^2 \\ &\geq \frac{\nu^2}{4},\end{aligned}$$

which indicates that $D_\nu \subseteq C_\lambda \cap B_{\frac{\sqrt{p}\nu}{2}}^o(\hat{s})$. Therefore,

$$\begin{aligned}\mathbb{P}\left(\max_{\mathbf{s} \in D_\nu} \tilde{S}(\mathbf{s}) \geq \max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) - \varepsilon\right) &\leq \mathbb{P}\left(\max_{\mathbf{s} \in D_\nu} \tilde{S}(\mathbf{s}) \geq \max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) - \varepsilon \bigcap E\right) + \mathbb{P}(E^C) \\ &\leq \mathbb{P}\left(\max_{\mathbf{s} \in C_\lambda \cap B_{\frac{\sqrt{p}\nu}{2}}^o(\hat{s})} \tilde{S}(\mathbf{s}) \geq \max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) - \varepsilon\right) + \mathbb{P}(E^C).\end{aligned}\tag{150}$$

According to (149), the first term in (150) vanishes if $\varepsilon \leq \frac{\gamma\nu}{2}$. For the second term, it is not hard to show $(H, B) \mapsto (\tau_*^{-1}[(B + \sigma_* H) - \eta(B + \sigma_* H; \mu_{Y_*}, \mu_{\tau_* \Lambda})], B)$ is a Lipschitz continuous mapping and from (164), $W_2(\mu_{\mathbf{h}, \beta}, H \otimes B) \xrightarrow{\mathbb{P}} 0$, so similar as (165) we can get $W_2(\mu_{\hat{s}, \beta}, \mu_{\hat{s}, B}) \xrightarrow{\mathbb{P}} 0$ and thus $\mathbb{P}(E^C) \rightarrow 0$. Therefore, from (150), for any $\nu > 0$, there exists $\varepsilon_0 > 0$ such that any $\varepsilon \leq \varepsilon_0$ satisfies $\mathbb{P}(\max_{\mathbf{s} \in D_\nu} \tilde{S}(\mathbf{s}) \geq \max_{\mathbf{s} \in C_\lambda} \tilde{S}(\mathbf{s}) - \varepsilon)$. Substituting this back to (148), we finish the proof. $\blacksquare$

J. Properties of Limiting Scalar Problem

It turns out that the limiting behavior of (2) is fully captured by (55). In this section, we study the key properties of (55).

Lemma 14: The minimax problem (55) has a unique optimal solution (σ_*, θ_*) , which is also the unique solution to the equation:

$$\begin{aligned}\sigma^2 &= \sigma_w^2 + \frac{1}{\delta} \mathbb{E}[\eta(Y; \mu_Y, \mu_{\sigma \Lambda / \theta}) - B]^2, \\ \theta &= \sigma \left[1 - \frac{1}{\delta} \mathbb{E}\eta'(Y; \mu_Y, \mu_{\sigma \Lambda / \theta})\right],\end{aligned}\tag{151}$$

where $Y = B + \sigma H$, with $H \sim \mathcal{N}(0, 1)$ independent of $B \sim \mu_B$. Besides, there exists $\theta_{\min} > 0$ such that $\theta^* \geq \theta_{\min}$.

Proof: The proof includes two steps: (I) show the saddle point of $\Psi(\sigma, \theta)$ exists and is unique and it is also the unique optimal solution of the minimax problem (55), (II) show (σ_*, θ_*) is the saddle point of $\Psi(\sigma, \theta)$ if and only if it is the solution to (151).

We first show the set of saddle points of $\Psi(\sigma, \theta)$ is nonempty and compact, using Proposition 5.5.7 of [60]. To apply this result, it suffices to check: (i) $\Psi(\cdot, \theta)$ and $-\Psi(\sigma, \cdot)$ are convex and closed for any fixed $\theta \geq 0$ and

$\sigma \geq \sigma_w$, (ii) there exists some $\bar{\theta} \geq 0$, $\bar{\sigma} \geq \sigma_w$ and $\gamma_1, \gamma_2 \in \mathbb{R}$ such that the level sets $\{\sigma \geq \sigma_w \mid \Psi(\sigma, \bar{\theta}) \leq \gamma_1\}$ and $\{\theta \geq 0 \mid -\Psi(\bar{\sigma}, \theta) \leq \gamma_2\}$ are both non-empty and compact. From Lemma 16, $\mathcal{F}(\sigma, \theta)$ is convex-concave and continuously differentiable with respect to σ and θ . Therefore, condition (i) is satisfied. Also partial derivatives of $\mathcal{F}(\sigma, \theta)$ can be computed as

$$\frac{\partial \Psi}{\partial \sigma} = \frac{\theta}{2\sigma^2} \left(\sigma^2 - \sigma_w^2 - \frac{1}{\delta} \lim_{p \rightarrow \infty} \frac{1}{p} \mathbb{E} \|\beta - \text{Prox}_{\sigma\lambda/\theta}(\beta + \sigma\mathbf{h})\|^2 \right), \quad (152)$$

$$\frac{\partial \Psi}{\partial \theta} = \frac{1}{2} \left(\frac{\sigma_w^2}{\sigma} + \sigma \right) - \theta + \frac{1}{\delta} \left[\lim_{p \rightarrow \infty} \mathbb{E} \frac{\|\beta + \sigma\mathbf{h} - \text{Prox}_{\sigma\lambda/\theta}(\beta + \sigma\mathbf{h})\|^2}{2\sigma p} - \frac{\sigma}{2} \right], \quad (153)$$

using (55), (178) and (177). Next we show $\{\sigma \geq \sigma_w \mid \Psi(\sigma, \bar{\theta}) \leq \gamma_1\}$ is non-empty and compact for some $\bar{\theta} \geq 0$ and $\gamma_1 \in \mathbb{R}$. First, we have

$$\begin{aligned} \mathbb{E} \|\beta - \text{Prox}_{\sigma\lambda/\theta}(\beta + \sigma\mathbf{h})\|^2 &\stackrel{(a)}{=} \sigma^2 \mathbb{E} \left\| \frac{\beta}{\sigma} - \text{Prox}_{\lambda/\theta} \left(\frac{\beta}{\sigma} + \mathbf{h} \right) \right\|^2 \\ &\leq 2\sigma^2 \mathbb{E} \left(\frac{\|\beta\|^2}{\sigma^2} \right) + 2\|\text{Prox}_{\lambda/\theta}(\mathbf{h})\|^2 + 2\|\text{Prox}_{\lambda/\theta}(\frac{\beta}{\sigma} + \mathbf{h}) - \text{Prox}_{\lambda/\theta}(\mathbf{h})\|^2 \\ &\stackrel{(b)}{\leq} 4\sigma^2 \mathbb{E} \|\text{Prox}_{\lambda/\theta}(\mathbf{h})\|^2 + 6\|\beta\|^2 \\ &\stackrel{(c)}{\leq} 4\sigma^2 \sum_{i=1}^p \mathbb{E} \max \{ |h_i| - \frac{\bar{\lambda}}{\theta}, 0 \}^2 + 6\|\beta\|^2, \end{aligned} \quad (154)$$

where (a) follows from the identity $\text{Prox}_{\tau\lambda}(x) = \tau \text{Prox}_{\lambda}(x/\tau)$, (b) follows from the non-expansiveness of proximal operator and (c) is a consequence of (188), where $\bar{\lambda} = \frac{1}{p} \sum_{i=1}^p \lambda_i$. Plugging the bound (154) into (152) gives

$$\frac{\partial \Psi}{\partial \sigma} \geq \frac{\theta}{2\sigma^2} \left[\sigma^2 \left(1 - \frac{8}{\delta} \int_{\frac{\mathbb{E}\Lambda}{\theta}}^{\infty} (z - \frac{\mathbb{E}\Lambda}{\theta}) d\Phi(z) \right) - \sigma_w^2 - \frac{6}{\delta} \mathbb{E} B^2 \right], \quad (155)$$

where $\Phi(z)$ is CDF of standard Gaussian. When $\mathbb{E}\Lambda > 0$, from (155) we know there exists $\theta_1 > 0$ and $\sigma_1 \geq \sigma_w$ such that $\frac{\partial \Psi(\sigma, \theta_1)}{\partial \sigma} \geq \frac{\theta_1}{8}$ for all $\sigma \geq \sigma_1$; when $\mathbb{E}\Lambda = 0$, by our assumption we must have $\delta > 1$, so from (157) we have $\frac{\partial \Psi(\sigma, \theta)}{\partial \sigma} = \frac{\theta}{2\sigma^2} [(1 - \frac{1}{\delta})\sigma^2 - \sigma_w^2]$ for any $\theta > 0$, $\sigma \geq \sigma_w$ implying $\frac{\partial \Psi(\sigma, \theta)}{\partial \sigma} \geq \frac{\theta}{4}(1 - \frac{1}{\delta})$ for all $\sigma \geq \sqrt{\frac{2\delta}{\delta-1}}\sigma_w$. Therefore, there exists $\bar{\theta} > 0$, $c > 0$ and $K \geq \sigma_w$ such that $\frac{\partial \Psi(\sigma, \bar{\theta})}{\partial \sigma} \geq c$ for any $\sigma \geq K$. This means that $\Psi(\sigma, \bar{\theta}) > \Psi(K, \bar{\theta})$ for all $\sigma > K$, so the set $\{\sigma \geq \sigma_w \mid \Psi(\sigma, \bar{\theta}) \leq \Psi(K, \bar{\theta})\} \subset [\sigma_w, K]$ and it is non-empty (include at least one point $\sigma = K$) and closed since $\Psi(\cdot, \bar{\theta})$ is a closed function. As a result, we can take $\gamma_1 = \Psi(K, \bar{\theta})$ and the level set $\{\sigma \geq \sigma_w \mid \Psi(\sigma, \bar{\theta}) \leq \gamma_1\}$ is non-empty and compact. On the other hand, we can show $\{\theta \geq 0 \mid -\Psi(\sigma_w, \theta) \leq 0\}$ is non-empty and compact. First since $\Psi(\sigma_w, \cdot)$ is 1-strongly concave and continuously differentiable, we have for any $\theta \geq 0$,

$$\begin{aligned} \Psi(\sigma_w, \theta) &\leq \Psi(\sigma_w, 0) + \frac{\partial \Psi(\sigma_w, 0)}{\partial \theta} \theta - \frac{1}{2} \theta^2 \\ &\leq \left(\sigma_w + \frac{\mathbb{E} B^2}{2\sigma_w \delta} \right) |\theta| - \frac{1}{2} \theta^2, \end{aligned}$$

where in the last step we use $\Psi(\sigma_w, 0) = 0$ and $\frac{\partial \Psi(\sigma_w, 0)}{\partial \theta} = \sigma_w + \frac{\mathbb{E} B^2}{2\sigma_w \delta}$, which can be deduced from (55) and (153). Then the level set $\{\theta \geq 0 \mid -\Psi(\sigma_w, \theta) \leq 0\} \subset \left[0, 2\left(\sigma_w + \frac{\mathbb{E} B^2}{2\sigma_w \delta} \right) \right]$ and it is non-empty (include at least one point $\theta = 0$) and closed since $\Psi(\sigma_w, \cdot)$ is a closed function. Letting $\gamma_2 = 0$, we verify condition (ii).

Up to now, we have proved the existence and boundedness of saddle points of $\Psi(\sigma, \theta)$. Next we prove the uniqueness. To do this, it suffices to show the optimal solution of $\min_{\sigma \geq \sigma_w} \max_{\theta \geq 0} \Psi(\sigma, \theta)$ is bounded and unique, then the uniqueness of saddle points follows due to the fact that each saddle point of $\Psi(\sigma, \theta)$ is also an optimal

solution of $\min_{\sigma \geq \sigma_w} \max_{\theta \geq 0} \Psi(\sigma, \theta)$ [60, Proposition 3.4.1]. First, we show that any σ_* should be bounded. Indeed, from the verification of condition (ii) above, we know there exists $c > 0$ and $K \geq \sigma_w$ such that for any $\sigma \geq K$,

$$\max_{\theta \geq 0} \Psi(\sigma, \theta) \geq \Psi(\sigma, \bar{\theta}) \geq \Psi(K, \bar{\theta}) + c(\sigma - K),$$

so we must have

$$\sigma_* \leq \underbrace{K + 1 + \frac{\max_{\sigma \in [\sigma_w, K]} \max_{\theta \geq 0} \Psi(\sigma, \theta) - \Psi(K, \bar{\theta})}{c}}_{:=C_1},$$

otherwise, $\min_{\sigma \geq \sigma_w} \max_{\theta \geq 0} \Psi(\sigma, \theta) > \min_{\sigma \in [\sigma_w, K]} \max_{\theta \geq 0} \Psi(\sigma, \theta)$ leading to a contradiction. On the other hand, we can also show $\theta_*(\sigma) := \arg\max_{\theta \geq 0} \Psi(\sigma, \theta)$ is uniformly bounded for $\sigma \in [\sigma_w, C_1]$. To see this, note from (153)

$$\begin{aligned} \frac{\partial \Psi}{\partial \theta} &\leq -\theta + 2 \left[\max_{\sigma \in [\sigma_w, C_1]} \sigma \left(1 + \frac{1}{\delta}\right) + \frac{\mathbb{E}B^2 + \sigma_w^2 \delta}{\sigma \delta} \right], \\ \Rightarrow \theta_*(\sigma) &\leq 2 \underbrace{\left[\max_{\sigma \in [\sigma_w, C_1]} \sigma \left(1 + \frac{1}{\delta}\right) + \frac{\mathbb{E}B^2 + \sigma_w^2 \delta}{\sigma \delta} \right]}_{:=C_2}, \quad \forall \sigma \in [\sigma_w, C_1]. \end{aligned}$$

As a result, $\max_{\theta \geq 0} \Psi(\sigma, \theta) = \max_{\theta \in [0, C_2]} \Psi(\sigma, \theta)$ for $\sigma \in [\sigma_w, C_1]$. Therefore, by Berge Maximum Theorem [61, Theorem 17.31], $\theta_*(\sigma)$ is an upper hemicontinuous correspondence on $[\sigma_w, C_1]$. By strong concavity of $\Psi(\sigma, \cdot)$, $\sigma \mapsto \theta_*(\sigma)$ is a function (*i.e.*, single-valued correspondence). As a result, one can easily check by definition that $\theta_*(\sigma)$ is continuous on $[\sigma_w, C_1]$. Besides, we can get $\theta_*(\sigma) > 0$ for any $\sigma \in [\sigma_w, C_1]$. Indeed from (153) when $\mathbb{P}(\Lambda = 0) < 1$, $\frac{\partial \Psi(\sigma, 0)}{\partial \theta} = \frac{1}{2} \left(\frac{\sigma_w^2}{\sigma} + \sigma \right) + \frac{\mathbb{E}B^2}{2\sigma\delta} > 0$ and when $\mathbb{P}(\Lambda = 0) = 1$, $\delta \geq 1$, $\frac{\partial \Psi(\sigma, 0)}{\partial \theta} \geq \frac{\sigma_w^2}{2\sigma} + \frac{\sigma}{2} \left(1 - \frac{1}{\delta}\right) > 0$. Therefore, there exists $\theta_{\min} > 0$ such that

$$\theta_*(\sigma) \geq \theta_{\min}, \quad (156)$$

for any $\sigma \in [\sigma_w, C_1]$. Since $\theta^*(\sigma) \in [\theta_{\min}, C_2]$ for any $\sigma \in [\sigma_w, C_1]$, we get $\max_{\theta \geq 0} \Psi(\sigma, \theta) = \max_{\theta \in [\theta_{\min}, C_2]} \Psi(\sigma, \theta)$ for any $\sigma \in [\sigma_w, C_1]$. On the other hand, $\Psi(\cdot, \theta)$ is $\frac{\sigma_w^2 \theta_{\min}}{C_1^3}$ -strongly convex on $[\sigma_w, C_1]$ for any fixed $\theta \in [\theta_{\min}, C_2]$, so we can check by definition that the function $\max_{\theta \in [\theta_{\min}, C_2]} \Psi(\cdot, \theta)$ is also $\frac{\sigma_w^2 \theta_{\min}}{C_1^3}$ -strongly convex on $[\sigma_w, C_1]$. We conclude that $\sigma \mapsto \max_{\theta \geq 0} \Psi(\sigma, \theta)$ is $\frac{\sigma_w^2 \theta_{\min}}{C_1^3}$ -strongly convex on $[\sigma_w, C_1]$, since $\max_{\theta \geq 0} \Psi(\sigma, \theta) = \max_{\theta \in [\theta_{\min}, C_2]} \Psi(\sigma, \theta)$. Recall that any optimal solution $\sigma_* = \arg\min_{\sigma \geq \sigma_w} \max_{\theta \geq 0} \Psi(\sigma, \theta)$ should lie in $[\sigma_w, C_1]$, so the uniqueness holds.

Finally, we show (σ_*, θ_*) is a saddle point of $\Psi(\sigma, \theta)$ if and only if it is a solution of (151). From (152) and (153), $\frac{\partial \Psi(\sigma_w, \theta)}{\partial \sigma} \leq 0$ for any $\theta \geq 0$ and $\frac{\partial \Psi(\sigma, 0)}{\partial \theta} \geq 0$ for any $\sigma \geq \sigma_w$. Since $\Psi(\sigma, \theta)$ is convex-concave and continuously differentiable, by first order optimality condition we know (σ_*, θ_*) is a saddle point if and only if $\frac{\partial \Psi(\sigma_*, \theta_*)}{\partial \sigma} = \frac{\partial \Psi(\sigma_*, \theta_*)}{\partial \theta} = 0$. On the other hand, from (55), (167) and (168) we can get

$$\frac{\partial \Psi}{\partial \sigma} = \frac{\theta}{2\sigma^2} \left(\sigma^2 - \sigma_w^2 - \frac{1}{\delta} \mathbb{E}(\eta(B + \sigma H; \mu_Y, \mu_{\sigma\Lambda/\theta}) - B)^2 \right), \quad (157)$$

$$\frac{\partial \Psi}{\partial \theta} = \frac{1}{2} \left(\frac{\sigma_w^2}{\sigma} + \sigma \right) - \theta + \frac{1}{\delta} \left[\frac{\mathbb{E}(\eta(B + \sigma H; \mu_Y, \mu_{\sigma\Lambda/\theta}) - B)^2}{2\sigma} - \sigma \mathbb{E}\eta'(B + \sigma H; \mu_Y, \mu_{\sigma\Lambda/\theta}) \right]. \quad (158)$$

Note that (157) and (158) are actually the scalar representation of (152) and (153). Setting the RHS of (157) and (158) to be zero, we can get (151). $\blacksquare$

K. Moreau Envelope of J_λ

Recall that for $\tau > 0$, the Moreau envelope of $J_\lambda(\mathbf{x})$ is:

$$\mathcal{M}_\lambda(\mathbf{y}; \tau) = \min_{\mathbf{x}} \frac{1}{2\tau} \|\mathbf{x} - \mathbf{y}\|^2 + J_\lambda(\mathbf{x}) \quad (159)$$

and the corresponding optimal solution is the proximal operator $\text{Prox}_{\tau\lambda}(\mathbf{y})$. In this section, we study the limiting behavior of the following function:

$$\begin{aligned} \mathcal{F}_p(\sigma, \theta) &\stackrel{\text{def}}{=} \frac{1}{p} \mathcal{M}_\lambda(\beta + \sigma \mathbf{h}; \frac{\sigma}{\theta}) \\ &= \frac{1}{p} \min_{\mathbf{x}} \frac{\theta}{2\sigma} \|\mathbf{x} - (\beta + \sigma \mathbf{h})\|^2 + J_\lambda(\mathbf{x}), \end{aligned} \quad (160)$$

where $(\sigma, \theta) \in \mathbb{R}_{>0} \times \mathbb{R}_{\geq 0}$, $\mathbf{h} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$ and β and λ are the same as in (1) and (2).

Lemma 15: Consider a sequence of instances $\{\mathbf{h}^{(p)}, \beta^{(p)}, \lambda^{(p)}\}_{p \in \mathbb{Z}^+}$, where $\mathbf{h}^{(p)} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$ are all independent and $\{\beta^{(p)}\}_{p \in \mathbb{Z}^+}$, $\{\lambda^{(p)}\}_{p \in \mathbb{Z}^+}$ are both converging sequences with limiting measure μ_B and μ_Λ . As $p \rightarrow \infty$, for every $(\sigma, \theta) \in \mathbb{R}_{>0} \times \mathbb{R}_{\geq 0}$,

$$\mathcal{F}_p(\sigma, \theta) \xrightarrow{a.s.} \mathcal{F}(\sigma, \theta), \quad (161)$$

and

$$\mathbb{E}_{\mathbf{h}} \mathcal{F}_p(\sigma, \theta) \rightarrow \mathcal{F}(\sigma, \theta), \quad (162)$$

where

$$\mathcal{F}(\sigma, \theta) \stackrel{\text{def}}{=} \min_{g \in \mathcal{I}} \frac{\theta}{2\sigma} \mathbb{E}[Y - g(Y)]^2 + \int_0^1 F_\Lambda^{-1}(u) F_{|g(Y)|}^{-1}(u) du. \quad (163)$$

Here, $Y = B + \sigma H$, with $H \sim \mathcal{N}(0, 1)$, $B \sim \mu_B$ independent and $\eta(\cdot)$ is the limiting scalar function in Proposition 1.

Proof: For notational simplicity, we will omit the superscript “(p)” in $\mathbf{h}^{(p)}$, $\beta^{(p)}$ and $\lambda^{(p)}$. We first show the empirical measure $\mu_{\mathbf{h}, \beta}$ converges almost surely to $H \otimes B$ under Wasserstein-2 distance. Let $g(x, y)$ be a bounded and continuous test function. By strong law of large number for triangular array [58, Theorem 2.1], we can get $\frac{1}{p} \sum_{i=1}^p [g(h_i, \beta_i) - \mathbb{E}_H g(H, \beta_i)] \xrightarrow{a.s.} 0$. For any y , $\mathbb{E}_H g(H, y)$ is bounded and it is not hard to show $y \mapsto \mathbb{E}_H g(H, y)$ is continuous: we know $g(h, y)$ is uniformly continuous over any compact set in $\mathbb{R}^2$, so for $C = \Phi^{-1}(1 - \frac{\varepsilon}{8\|g\|_\infty})$ and any $y_0 \in \mathbb{R}$, $\varepsilon > 0$, there exists $\delta > 0$ such that $|g(h, y) - g(h, y_0)| \leq \frac{\varepsilon}{2}$ whenever $|y - y_0| \leq \delta$ and $|h| \leq C$. Hence,

$$|\mathbb{E}_H g(H, y) - \mathbb{E}_H g(H, y_0)| \leq \mathbb{E}_H [\mathbb{I}_{|H| \leq C} |g(H, y) - g(H, y_0)|] + 2\|g\|_\infty \mathbb{E}_H (\mathbb{I}_{|H| > C}) \leq \varepsilon.$$

This shows the continuity of $y \mapsto \mathbb{E}_H [g(H, y)]$. Therefore, $\frac{1}{p} \sum_{i=1}^p \mathbb{E}_H [g(H, \beta_i)] \rightarrow \mathbb{E}_{H, B} [g(H, B)]$ and we get for any bounded and continuous $g(x, y)$, $\frac{1}{p} \sum_{i=1}^p g(h_i, \beta_i) \xrightarrow{a.s.} \mathbb{E}_{H, B} [g(H, B)]$ indicating the almost sure weak convergence of $\mu_{\mathbf{h}, \beta}$ to $H \otimes B$. On the other hand, by strong law of large number again, $\frac{1}{p} \sum_{i=1}^p h_i^2 + \beta_i^2 \xrightarrow{a.s.} 1 + \mathbb{E} B^2$. Therefore, by Theorem 7.12 (iii) in [56],

$$W_2(\mu_{\mathbf{h}, \beta}, H \otimes B) \xrightarrow{a.s.} 0. \quad (164)$$

Then we can show $W_2(\mu_{\mathbf{y}}, \mu_Y) \xrightarrow{a.s.} 0$, where $\mathbf{y} := \beta + \sigma \mathbf{h}$. Indeed,

$$\begin{aligned}
W_2(\mu_{\mathbf{y}}, \mu_Y)^2 &= \inf_{\pi \in \Pi(\mu_{\mathbf{y}}, \mu_Y)} \int (x - y)^2 d\pi(x, y) \\
&= \inf_{\pi \in \Pi(\mu_{\mathbf{h}, \beta}, H \otimes B)} \int [(x_2 + \sigma x_1) - (y_2 + \sigma y_1)]^2 d\pi(\mathbf{x}, \mathbf{y}) \\
&\leq 2(\sigma^2 + 1) \inf_{\pi \in \Pi(\mu_{\mathbf{h}, \beta}, H \otimes B)} \int \|\mathbf{x} - \mathbf{y}\|^2 d\pi(\mathbf{x}, \mathbf{y}) \\
&= 2(\sigma^2 + 1) W_2(\mu_{\mathbf{h}, \beta}, H \otimes B)^2
\end{aligned} \tag{165}$$

and since $W_2(\mu_{\mathbf{h}, \beta}, H \otimes B) \xrightarrow{a.s.} 0$, we get $W_2(\mu_{\mathbf{y}}, \mu_Y) \xrightarrow{a.s.} 0$. Similarly, for $\theta > 0$ we can show $W_2(\mu_{\sigma \lambda / \theta}, \mu_{\sigma \Lambda / \theta}) \rightarrow 0$ from our assumption that $W_2(\mu_{\lambda}, \mu_{\Lambda}) \rightarrow 0$.

Now we can prove (161). For $\theta = 0$, it directly follows from (160) and (163) that $\mathcal{F}_p(\sigma, \theta) = \mathcal{F}(\sigma, \theta) = 0$. For $\theta > 0$, we have $W_2(\mu_{\mathbf{y}}, \mu_Y) \xrightarrow{a.s.} 0$ and $W_2(\mu_{\sigma \lambda / \theta}, \mu_{\sigma \Lambda / \theta}) \rightarrow 0$. Then from Proposition 1, we can get (161) by letting $\tau = \sigma / \theta$.

To prove (162), first observe that:

$$\mathcal{F}_p(\sigma, \theta) = \frac{\min_{\mathbf{x}} \frac{\theta}{2\sigma} \|\sigma \mathbf{h} + \beta - \mathbf{x}\|^2 + J_{\lambda}(\mathbf{x})}{p} \leq \theta \frac{\sigma^2 \|\mathbf{h}\|^2 + \|\beta\|^2}{\sigma p}.$$

By strong law of large number (SLLN), we have

$$\theta \frac{\sigma^2 \|\mathbf{h}\|^2 + \|\beta\|^2}{\sigma p} \xrightarrow{a.s.} \frac{\theta(\sigma^2 + \mathbb{E}B^2)}{\sigma} = \lim_{p \rightarrow \infty} \mathbb{E}_{\mathbf{h}} \theta \frac{\sigma^2 \|\mathbf{h}\|^2 + \|\beta\|^2}{\sigma p}.$$

In other words,

$$\lim_{p \rightarrow \infty} \mathbb{E}_{\mathbf{h}} \theta \frac{\sigma^2 \|\mathbf{h}\|^2 + \|\beta\|^2}{\sigma p} = \mathbb{E} \lim_{p \rightarrow \infty} \theta \frac{\sigma^2 \|\mathbf{h}\|^2 + \|\beta\|^2}{\sigma p}.$$

Then by (161) and generalized dominated convergence theorem (GDCT) (Theorem 19 in Sec. 4.4 of [57]), we have

$$\lim_{p \rightarrow \infty} \mathbb{E} \mathcal{F}_p(\sigma, \theta) = \mathbb{E} \lim_{p \rightarrow \infty} \mathcal{F}_p(\sigma, \theta), \tag{166}$$

which is exactly (162). ■

Lemma 16: $\mathcal{F}(\sigma, \theta)$ is convex-concave and continuously differentiable with respect to both σ and θ on $\mathbb{R}_{>0} \times \mathbb{R}_{\geq 0}$, with partial derivatives:

$$\frac{\partial \mathcal{F}(\sigma, \theta)}{\partial \sigma} = \frac{-\theta \mathbb{E}[B - \eta(Y)]^2}{2\sigma^2} + \frac{\theta}{2}, \tag{167}$$

$$\frac{\partial \mathcal{F}(\sigma, \theta)}{\partial \theta} = \frac{\mathbb{E}[\eta(Y) - B]^2}{2\sigma} - \sigma \mathbb{E} \eta'(Y) + \frac{\sigma}{2}, \tag{168}$$

where $B \sim \mu_B$, $\Lambda \sim \mu_{\Lambda}$, $Y = B + \sigma H$, with $H \sim \mathcal{N}(0, 1)$ independent of B . Here, when $\theta > 0$, $\eta(\cdot) = \eta(\cdot; F_Y, F_{\sigma \Lambda / \theta})$; when $\theta = 0$, we let $\eta(y) = 0$, if $\mathbb{P}(\Lambda = 0) < 1$ and $\eta(y) = y$, if $\mathbb{P}(\Lambda = 0) = 1$.

Proof: When $\mathbb{P}(\Lambda = 0) = 1$, from (163) we have $\mathcal{F}(\sigma, \theta) = 0$ for any $(\sigma, \theta) \in \mathbb{R}_{>0} \times \mathbb{R}_{\geq 0}$, so $\frac{\partial \mathcal{F}(\sigma, \theta)}{\partial \sigma} = \frac{\partial \mathcal{F}(\sigma, \theta)}{\partial \theta} = 0$. In this case, all the results trivially hold. Therefore, it suffices to consider $\mathbb{P}(\Lambda = 0) < 1$.

We first prove the convexity of $\mathcal{F}(\sigma, \theta)$. Note $\mathcal{F}_p(\sigma, \theta)$ [defined in (160)] can be rewritten as:

$$\mathcal{F}_p(\sigma, \theta) = \frac{1}{p} \min_{\mathbf{v}} \frac{\sigma \theta}{2} \|\mathbf{v} / \sigma - \mathbf{h}\|^2 + J_{\lambda}(\mathbf{v} + \beta). \tag{169}$$

Since $h(\mathbf{v}) = \frac{\theta}{2}\|\mathbf{v} - \mathbf{h}\|^2$ is a convex function (due to $\theta \geq 0$), $(\sigma, \mathbf{v}) \mapsto \frac{\sigma\theta}{2}\|\mathbf{v}/\sigma - \mathbf{h}\|^2$ is convex since it is the perspective function of $h(\mathbf{v})$. Therefore, the objective function in (169) is jointly convex w.r.t. $(\sigma, \mathbf{v})$ and after partial minimization over $\mathbf{v}$, $\mathcal{F}_p(\cdot, \theta)$ is still convex. On the other hand, $\mathcal{F}_p(\sigma, \theta)$ as a function of θ is the infimum of a family of linear functions of θ , so $\mathcal{F}_p(\sigma, \cdot)$ is concave. Denote $\overline{\mathcal{F}}_p(\sigma, \theta) := \mathbb{E}_{\mathbf{h}} \mathcal{F}_p(\sigma, \theta)$. Clearly, $\overline{\mathcal{F}}_p(\sigma, \theta)$ and $F(\sigma, \theta)$ are still convex-concave, since taking expectation and limit preserves convexity.

Then we show for fixed $\sigma \in \mathbb{R}_{>0}$, $\mathcal{F}_p(\sigma, \cdot)$ is continuously differentiable on $\mathbb{R}_{\geq 0}$. We follow the same argument as Theorem 2.26 in [62]. Denote $\mathbf{y} := \boldsymbol{\beta} + \sigma \mathbf{h}$ and

$$g(\varepsilon) := \mathcal{F}_p(\sigma, \theta + \varepsilon) - \mathcal{F}_p(\sigma, \theta) - \frac{\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma p} \varepsilon, \quad (170)$$

where $\varepsilon \geq -\theta$. On one hand, we have

$$\begin{aligned} g(\varepsilon) &\leq \frac{(\theta + \varepsilon)\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma p} + \frac{J_{\boldsymbol{\lambda}}(\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y}))}{p} \\ &\quad - \left[\frac{\theta\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma p} + \frac{J_{\boldsymbol{\lambda}}(\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y}))}{p} \right] - \frac{\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma p} \varepsilon = 0 \end{aligned} \quad (171)$$

On the other hand,

$$\begin{aligned} g(\varepsilon) &\geq \frac{(\theta + \varepsilon)\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/(\theta + \varepsilon)}(\mathbf{y})\|^2}{2\sigma p} + \frac{J_{\boldsymbol{\lambda}}(\text{Prox}_{\sigma\boldsymbol{\lambda}/(\theta + \varepsilon)}(\mathbf{y}))}{p} \\ &\quad - \left[\frac{\theta\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/(\theta + \varepsilon)}(\mathbf{y})\|^2}{2\sigma p} + \frac{J_{\boldsymbol{\lambda}}(\text{Prox}_{\sigma\boldsymbol{\lambda}/(\theta + \varepsilon)}(\mathbf{y}))}{p} \right] - \frac{\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma p} \varepsilon \\ &= \frac{\varepsilon}{\sigma} \left(\frac{\|\text{Prox}_{\sigma\boldsymbol{\lambda}/(\theta + \varepsilon)}(\mathbf{y})\|^2 - \|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2p} \right). \end{aligned} \quad (172)$$

From Lemma 17 there exists $C > 0$ such that $\frac{1}{p} \left| \frac{\partial \|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{\partial \theta} \right| \leq C$ for any $\theta \geq 0$. Combined with (171) and (172), this indicates $-\frac{C\varepsilon^2}{2\sigma} \leq g(\varepsilon) \leq 0$. Substituting the boundedness of $g(\varepsilon)$ into (170) yields for any $\theta \geq 0$,

$$\frac{\partial \mathcal{F}_p(\sigma, \theta)}{\partial \theta} = \frac{\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma p}. \quad (173)$$

Also $\frac{\partial \mathcal{F}_p(\sigma, \theta)}{\partial \theta}$ is continuous, which follows from (173), (183) and (184).

Following the same procedure as above, we can also show for fixed $\theta \in \mathbb{R}_{\geq 0}$, $\mathcal{F}_p(\cdot, \theta)$ is continuously differentiable on $\mathbb{R}_{>0}$ with

$$\frac{\partial \mathcal{F}_p(\sigma, \theta)}{\partial \sigma} = -\frac{\theta\|\boldsymbol{\beta} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma^2 p} + \frac{\theta\|\mathbf{h}\|^2}{2p}. \quad (174)$$

From (173) and (174), we can get $\left| \frac{\partial \mathcal{F}_p(\sigma, \theta)}{\partial \theta} \right| \leq \frac{4(\sigma^2\|\mathbf{h}\|^2 + \|\boldsymbol{\beta}\|^2)}{\sigma p}$ and $\left| \frac{\partial \mathcal{F}_p(\sigma, \theta)}{\partial \sigma} \right| \leq \frac{3\theta\|\boldsymbol{\beta}\|^2 + \sigma^2(2 + \theta)\|\mathbf{h}\|^2}{\sigma^2 p}$. Since both bounds have finite expectation, by dominated convergence theorem (DCT) we have:

$$\frac{\partial \overline{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} = \frac{\mathbb{E}\|\mathbf{y} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma p} \quad (175)$$

and

$$\frac{\partial \overline{\mathcal{F}}_p(\sigma, \theta)}{\partial \sigma} = \frac{-\theta\mathbb{E}\|\boldsymbol{\beta} - \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma^2 p} + \frac{\theta}{2}. \quad (176)$$

We now show $\mathcal{F}(\sigma, \cdot)$ and $\mathcal{F}(\cdot, \theta)$ are continuously differentiable on $\mathbb{R}_{\geq 0}$ and $\mathbb{R}_{>0}$, respectively. In particular, we only present the proof for $\mathcal{F}(\sigma, \cdot)$ and the case of $\mathcal{F}(\cdot, \theta)$ can be derived following same approach. The key is to establish the uniform convergence of $\frac{\partial \overline{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta}$ to $\frac{\partial \overline{\mathcal{F}}_{\infty}(\sigma, \theta)}{\partial \theta}$ on $\theta \geq 0$, where $\frac{\partial \overline{\mathcal{F}}_{\infty}(\sigma, \theta)}{\partial \theta} := \lim_{p \rightarrow \infty} \frac{\partial \overline{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta}$.

First consider the case $\theta > 0$. Let $[a, b] \subset \mathbb{R}_{>0}$ and $\theta_1, \theta_2 \in [a, b]$. We have

$$\begin{aligned}
\left| \frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta_1)}{\partial \theta} - \frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta_2)}{\partial \theta} \right| &\leq \mathbb{E} \left| \frac{\partial \mathcal{F}_p(\sigma, \theta_1)}{\partial \theta} - \frac{\partial \mathcal{F}_p(\sigma, \theta_2)}{\partial \theta} \right| \\
&\stackrel{(a)}{\leq} \frac{1}{\sigma} \left(\frac{1}{2p} \mathbb{E} \left| \frac{\partial \|\text{Prox}_{\sigma \lambda / \theta}(\mathbf{y})\|^2}{\partial \theta} \right|_{\theta=\theta'} + \frac{1}{p} \mathbb{E} \left| \frac{\partial \mathbf{y}^\top \text{Prox}_{\sigma \lambda / \theta}(\mathbf{y})}{\partial \theta} \right|_{\theta=\theta'} \right) |\theta_1 - \theta_2| \\
&\stackrel{(b)}{\leq} \frac{2}{\sigma \sqrt{p}} \frac{\mathbb{E} \|\mathbf{y}\|}{a^2} |\theta_1 - \theta_2| \\
&\stackrel{(c)}{\leq} C |\theta_1 - \theta_2|,
\end{aligned}$$

where (a) follows from (173) and intermediate value theorem, with $\theta' \in [\theta_1, \theta_2]$, (b) follows from (183) and (184), and in (c), $C > 0$ is some constant that only depends on σ , a and b . Therefore, $\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta}$ is C -Lipschitz continuous on $\theta \in [a, b]$ for any $p \in \mathbb{Z}^+$. Meanwhile, we can show $\{\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta}\}_{p \in \mathbb{Z}^+}$ converges pointwise to $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ on $\theta \geq 0$ following the proof of (166). Then for any $\varepsilon > 0$ we can construct a $\frac{\varepsilon}{3C}$ -epsilon net $\mathcal{E}$ such that for any $\theta \in [a, b]$, there exists $z \in \mathcal{E}$ satisfying $|\theta - z| \leq \frac{\varepsilon}{3C}$. Clearly, the cardinality $|\mathcal{E}| < \infty$. Therefore, for any $\varepsilon > 0$ there exists p_0 such that for any $m, p \geq p_0$, $\max_{z \in \mathcal{E}} \left| \frac{\partial \bar{\mathcal{F}}_m(\sigma, z)}{\partial \theta} - \frac{\partial \bar{\mathcal{F}}_p(\sigma, z)}{\partial \theta} \right| \leq \frac{\varepsilon}{3}$. Hence for any $\varepsilon > 0$, $\theta \in [a, b]$ and $m, p \geq p_0$,

$$\begin{aligned}
\left| \frac{\partial \bar{\mathcal{F}}_m(\sigma, \theta)}{\partial \theta} - \frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} \right| &\leq \left| \frac{\partial \bar{\mathcal{F}}_m(\sigma, \theta)}{\partial \theta} - \frac{\partial \bar{\mathcal{F}}_m(\sigma, z)}{\partial \theta} \right| + \left| \frac{\partial \bar{\mathcal{F}}_m(\sigma, z)}{\partial \theta} - \frac{\partial \bar{\mathcal{F}}_p(\sigma, z)}{\partial \theta} \right| + \left| \frac{\partial \bar{\mathcal{F}}_p(\sigma, z)}{\partial \theta} - \frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} \right| \\
&\leq C \times \frac{\varepsilon}{3C} + \frac{\varepsilon}{3} + C \times \frac{\varepsilon}{3C} = \varepsilon,
\end{aligned}$$

which implies the uniform convergence of $\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta}$ when $\theta \in [a, b]$. Meanwhile, by DCT and continuity of $\frac{\partial \mathcal{F}_p(\sigma, \theta)}{\partial \theta}$ for any fixed $\mathbf{h}$, $\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta}$ is also continuous w.r.t. θ . Then we can apply Theorem 7.12 in [63] to obtain $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ is continuous on $[a, b]$. Since $[a, b]$ above can be arbitrary subset of $\mathbb{R}_{>0}$, we actually obtain that $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ is continuous when $\theta > 0$.

Next we handle the case when $\theta = 0$. From (175), (183) and (184), we can get for any $\theta \geq 0$

$$\frac{\partial^2 \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta^2} = \frac{2}{\theta^2} \mathbb{E} \left\{ \sum_{i=1}^p \mathbb{I}_{[\text{Prox}_{\sigma \lambda / \theta}(\mathbf{y})]_i \neq 0} \left[[\text{Prox}_{\sigma \lambda / \theta}(|\mathbf{y}|)]_i - |y_i| \right] \right\} \leq 0,$$

where in the first equality we use DCT and the last inequality follows from Fact 1. Hence, $\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta}$ is non-increasing on $\theta \geq 0$, with $\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} = \mathbb{E} \frac{\|\mathbf{y}\|^2}{2\sigma p}$ and $\lim_{\theta \rightarrow \infty} \frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} = 0$. Since $\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta}$ converges to $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ pointwise, $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ is also non-increasing with $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, 0)}{\partial \theta} = \frac{\mathbb{E} B^2 + \sigma^2}{2\sigma}$ and $\lim_{\theta \rightarrow \infty} \frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta} = 0$. On the other hand, from (175)

$$\begin{aligned}
\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} &\geq \mathbb{E} \left(\frac{\|\mathbf{y}\|^2 - 2\|\mathbf{y}\| \|\text{Prox}_{\sigma \lambda / \theta}(\mathbf{y})\|}{2\sigma p} \right) \\
&\stackrel{(a)}{\geq} \mathbb{E} \left(\frac{\|\mathbf{y}\|^2 - 2\|\mathbf{y}\| \|\text{Prox}_{\sigma \bar{\lambda} / \theta}(\mathbf{y})\|}{2\sigma p} \right) \\
&\geq \frac{\mathbb{E} \|\mathbf{y}\|^2}{2\sigma p} - \frac{1}{\sigma} \left(\frac{\mathbb{E} \|\mathbf{y}\|^2}{p} \right)^{1/2} \left[\frac{1}{p} \sum_{i=1}^p \mathbb{E} (|y_i| - \frac{\sigma}{\theta} \bar{\lambda})_+^2 \right]^{1/2},
\end{aligned}$$

where step (a) follows from (188), with $\bar{\lambda} = \frac{\sum_{i=1}^p \lambda_i}{p} \mathbf{1}_p$. Hence taking $p \rightarrow \infty$ on both sides above, we have

$$\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta} \geq \frac{1}{2\sigma} \left[\mathbb{E} B^2 + \sigma^2 - 2\sqrt{(\mathbb{E} B^2 + \sigma^2) \mathbb{E} (|Y| - \frac{\sigma}{\theta} \mathbb{E} \Lambda)_+^2} \right].$$

Then taking $\theta \rightarrow 0^+$, we get

$$\begin{aligned} \lim_{\theta \rightarrow 0^+} \frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta} &\geq \lim_{\theta \rightarrow 0^+} \frac{1}{2\sigma} \left[\mathbb{E}B^2 + \sigma^2 - \frac{1}{\sigma} \sqrt{(\mathbb{E}B^2 + \sigma^2) \mathbb{E}(|Y| - \frac{\sigma}{\theta} \mathbb{E}\Lambda)_+^2} \right] \\ &\stackrel{(a)}{=} \frac{\mathbb{E}B^2 + \sigma^2}{2\sigma} = \frac{\partial \bar{\mathcal{F}}_\infty(\sigma, 0)}{\partial \theta}, \end{aligned}$$

where in step (a) we use DCT and $\mathbb{E}\Lambda > 0$. Since $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ is non-increasing on $\theta \geq 0$, we can get $\lim_{\theta \rightarrow 0^+} \frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta} = \frac{\partial \bar{\mathcal{F}}_\infty(\sigma, 0)}{\partial \theta}$. Recall that we have already proved $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ is continuous when $\theta > 0$, so $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ is continuous on $\mathbb{R}_{\geq 0}$.

Up to now, we have shown $\left\{ \frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} \right\}_{p \in \mathbb{Z}^+}$, $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ are all bounded, non-increasing and continuous functions on $\mathbb{R}_{\geq 0}$ and also $\left\{ \frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} \right\}_{p \in \mathbb{Z}^+}$ converges to $\frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ pointwise on $\theta \geq 0$. Therefore, following the proof of Glivenko-Cantelli theorem (for a reference, see Theorem 19.1 in [53]) we can show $\frac{\partial \bar{\mathcal{F}}_p(\sigma, \theta)}{\partial \theta} \rightarrow \frac{\partial \bar{\mathcal{F}}_\infty(\sigma, \theta)}{\partial \theta}$ converges uniformly on $\theta \geq 0$. In (162) we prove the pointwise convergence $\bar{\mathcal{F}}_p(\sigma, \theta) \rightarrow \mathcal{F}(\sigma, \theta)$. Then using Theorem 7.17 and 7.12 in [63] together with (175), we get $\mathcal{F}(\sigma, \cdot)$ is continuously differentiable on $\mathbb{R}_{\geq 0}$, with

$$\frac{\partial \mathcal{F}(\sigma, \theta)}{\partial \theta} = \lim_{p \rightarrow \infty} \frac{\mathbb{E} \|\mathbf{y} - \text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma p}, \quad (177)$$

where $\mathbf{y} = \boldsymbol{\beta} + \sigma \mathbf{h}$. Repeating the same procedure, we can also prove that $\mathcal{F}(\cdot, \theta)$ is continuously differentiable on $\mathbb{R}_{> 0}$, with

$$\frac{\partial \mathcal{F}(\sigma, \theta)}{\partial \sigma} = \lim_{p \rightarrow \infty} \frac{-\theta \mathbb{E} \|\boldsymbol{\beta} - \text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{2\sigma^2 p} + \frac{\theta}{2}. \quad (178)$$

Finally, we compute the limit in (177) and (157). From Proposition 1 and (165), we have $\frac{\|\text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y}) - \eta(\mathbf{y})\|_2^2}{p} \xrightarrow{a.s.} 0$ as $p \rightarrow \infty$, where $\eta := \eta(\cdot; \mu_{B+\sigma H}, \mu_{\sigma \Lambda/\tau})$. In addition,

$$\frac{1}{p} \left| \|\text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y}) - \boldsymbol{\beta}\|_2^2 - \|\eta(\mathbf{y}) - \boldsymbol{\beta}\|_2^2 \right| \leq \frac{2}{p} \|\text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y}) - \eta(\mathbf{y})\| (\|\mathbf{y}\| + \|\boldsymbol{\beta}\|),$$

so we can get

$$\frac{1}{p} \|\text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y}) - \boldsymbol{\beta}\|^2 \xrightarrow{a.s.} \frac{1}{p} \|\eta(\mathbf{y}) - \boldsymbol{\beta}\|^2 \xrightarrow{a.s.} \mathbb{E}[\eta(Y) - B]^2, \quad (179)$$

where the last step follows from Theorem 7.12 (iv) in [56] after combining (164) with the fact that $[\eta(b + \sigma h) - b]^2 \leq C(1 + h^2 + b^2)$ for any $h, b \in \mathbb{R}$ and some $C > 0$. Then similar as (166), we can obtain

$$\frac{1}{p} \mathbb{E} \|\text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y}) - \boldsymbol{\beta}\|^2 \rightarrow \mathbb{E}[\eta(Y) - B]^2. \quad (180)$$

On the other hand, we can also get

$$\frac{1}{p} \mathbf{h}^\top \text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y}) \xrightarrow{a.s.} \frac{1}{p} \mathbf{h}^\top \eta(\mathbf{y}) \xrightarrow{a.s.} \mathbb{E}[\eta(Y)H] \quad (181)$$

and

$$\frac{1}{p} \mathbb{E}[\mathbf{h}^\top \text{Prox}_{\sigma \boldsymbol{\lambda}/\theta}(\mathbf{y})] \rightarrow \mathbb{E}[\eta(Y)H] = \sigma \mathbb{E}\eta'(Y). \quad (182)$$

where in the last step we use Stein's lemma. Substituting (180) and (182) into (177) and (178), we reach at (167) and (168). $\blacksquare$

Lemma 17: For any $\theta \geq 0$, $\mathbf{y} \in \mathbb{R}^p$ and $\boldsymbol{\lambda} \in \mathbb{R}_{\geq 0}^p$, we have

$$\frac{\partial \|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{\partial \theta} = \begin{cases} 0 & \boldsymbol{\lambda} = \mathbf{0}, \\ \frac{2 \cdot \mathbf{1}^\top \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(|\mathbf{y}|)}{\theta^2} & \boldsymbol{\lambda} \neq \mathbf{0}, \end{cases} \quad (183)$$

and

$$\frac{\partial \mathbf{y}^\top \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})}{\partial \theta} = \begin{cases} 0 & \boldsymbol{\lambda} = \mathbf{0}, \\ \frac{1}{\theta^2} \sum_{i=1}^p |y_i| \mathbb{I}_{[\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})]_i \neq 0} & \boldsymbol{\lambda} \neq \mathbf{0}. \end{cases} \quad (184)$$

Here when $\boldsymbol{\lambda} \neq \mathbf{0}$ and $\theta = 0$, we let $\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y}) := \mathbf{0}$.

Proof: When $\boldsymbol{\lambda} = \mathbf{0}$, we have $\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y}) = \mathbf{y}$, so $\|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2 = \mathbf{y}^\top \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y}) = \|\mathbf{y}\|^2$ and $\frac{\partial \|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{\partial \theta} = \frac{\partial \mathbf{y}^\top \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})}{\partial \theta} = 0$.

Next we consider $\boldsymbol{\lambda} \neq \mathbf{0}$. From Lemma 2.3 and 2.4 of [9], for any $\mathbf{a} \in \mathbb{R}^p$ satisfying $0 \leq a_1 \leq a_2 \leq \dots \leq a_p$,

$$\frac{\partial [\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})]_i}{\partial \lambda_j} = -\frac{\mathbb{I}_{i \in I_j}}{\max\{|I_j|, 1\}},$$

where $I_j \stackrel{\text{def}}{=} \{k \in [p] \mid |[\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})]_k| = |[\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})]_j| \text{ and } [\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})]_k \neq 0\}$. Therefore,

$$\frac{\partial \|\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})\|^2}{\partial \lambda_j} = \sum_{i=1}^p 2[\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})]_i \frac{\partial [\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})]_i}{\partial \lambda_j} = -2[\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})]_j \quad (185)$$

and

$$\frac{\partial \mathbf{a}^\top \text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})}{\partial \lambda_j} = \sum_{i=1}^p a_i \frac{\partial [\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})]_i}{\partial \lambda_j} = -\frac{1}{\max\{|I_j|, 1\}} \sum_{i \in I_j} a_i. \quad (186)$$

On the other hand, by Fact 1, $\|\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})\|^2$ and $\mathbf{a}^\top \text{Prox}_{\boldsymbol{\lambda}}(\mathbf{a})$ only depend on $\mu_{\boldsymbol{\lambda}}$ and $\mu_{|\mathbf{a}|}$. Therefore, for any $\mathbf{y} \in \mathbb{R}^p$ and $\theta > 0$, it holds that $\frac{\partial \|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2}{\partial \theta} = \frac{2}{\theta^2} \mathbf{1}^\top \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(|\mathbf{y}|)$ and $\frac{\partial \mathbf{y}^\top \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})}{\partial \theta} = \frac{1}{\theta^2} \sum_{i=1}^p |y_i| \mathbb{I}_{[\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})]_i \neq 0}$ by chain rule. For $\theta = 0$, we need to study the behavior of $\|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|$ when θ is closed to 0. It can be shown (for a proof, see (A.11) in [7])

$$\text{Prox}_{\boldsymbol{\lambda}}(\mathbf{y}) = \mathbf{y} - \text{Proj}_{C_{\boldsymbol{\lambda}}}(\mathbf{y}), \quad (187)$$

where $C_{\boldsymbol{\lambda}} = \{\boldsymbol{\nu} \in \mathbb{R}^p \mid \boldsymbol{\nu} \prec \boldsymbol{\lambda}\}$ (“ $\prec$ ” denotes majorization, see Definition 2) is the unit ball of the dual norm of $J_{\boldsymbol{\lambda}}$ [9, Proposition 1.1] and $\text{Proj}_{C_{\boldsymbol{\lambda}}}$ is the orthogonal projection onto $C_{\boldsymbol{\lambda}}$. Take $\bar{\boldsymbol{\lambda}} = \frac{1}{p} \sum_{i=1}^p \lambda_i$ and it is not hard to show $\bar{\boldsymbol{\lambda}} := \bar{\lambda} \mathbf{1}_p$ satisfies $\bar{\boldsymbol{\lambda}} \prec \boldsymbol{\lambda}$. Clearly, $\frac{\sigma}{\theta} \bar{\boldsymbol{\lambda}} \prec \frac{\sigma}{\theta} \boldsymbol{\lambda}$ for $\sigma, \theta > 0$ and $C_{\sigma\bar{\boldsymbol{\lambda}}/\theta} \subset C_{\sigma\boldsymbol{\lambda}/\theta}$, so from (187) we have

$$\|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2 \leq \|\text{Prox}_{\sigma\bar{\boldsymbol{\lambda}}/\theta}(\mathbf{y})\|^2 = \sum_{i=1}^p \max\{|y_i| - \frac{\sigma\bar{\lambda}}{\theta}, 0\}^2. \quad (188)$$

Since $\bar{\lambda} > 0$ (due to $\boldsymbol{\lambda} \neq \mathbf{0}$), (188) indicates that when $0 < \theta \leq \frac{\sigma\bar{\lambda}}{\max\{|y_i|, 1\}}$, $\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y}) = \mathbf{0}$. On the other hand, we let $\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y}) := \mathbf{0}$, when $\boldsymbol{\lambda} \neq \mathbf{0}$ and $\theta = 0$. As a result, $\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y}) = \mathbf{0}$ for $\theta \in [0, \frac{\sigma\bar{\lambda}}{\max\{|y_i|, 1\}}]$ and combining the partial derivatives of $\|\text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})\|^2$ and $\mathbf{y}^\top \text{Prox}_{\sigma\boldsymbol{\lambda}/\theta}(\mathbf{y})$ on $\theta > 0$ obtained above, we can get (183) and (184). $\blacksquare$

L. Proof of Proposition 3

It directly follows from Proposition 1 that $\mathcal{M}_{\mu_Y} \subseteq \mathcal{I}$, so we just need to show $\mathcal{I} \subseteq \mathcal{M}_{\mu_Y}$.

For any $f \in \mathcal{I}$, consider the function $r(y) = y - f(y)$. It can be easily verified that $r(y) \in \mathcal{I}$. We claim that if we choose $\Lambda \sim r(|Y|)$ with $Y \sim \mu_Y$, then f is the optimal solution of (6) (when $\tau = 1$). Indeed, when $\tau = 1$ and $\Lambda \sim r(|Y|)$, (6) can be equivalently written as

$$\begin{aligned} \text{Problem (6)} &\stackrel{(a)}{=} \min_{g \in \mathcal{I}} \frac{1}{2} \mathbb{E}_{\mu_Y} [|Y| - g(|Y|)]^2 + \int_0^1 F_{r(|Y|)}^{-1}(u) F_{g(|Y|)}^{-1}(u) du \\ &\stackrel{(b)}{=} \min_{g \in \mathcal{I}} \frac{1}{2} \mathbb{E}_{\mu_Y} [|Y| - g(|Y|)]^2 + \mathbb{E}_{\mu_Y} [|Y| - f(|Y|)] g(|Y|) \\ &= \min_{g \in \mathcal{I}} \frac{1}{2} \mathbb{E}_{\mu_Y} [f(|Y|) - g(|Y|)]^2 + \frac{1}{2} \mathbb{E}_{\mu_Y} [Y^2 - f^2(|Y|)], \end{aligned} \quad (189)$$

where (a) follows from $g \in \mathcal{I}$ and in (b) we substitute $r(y) = y - f(y)$ and use the fact that $r \in \mathcal{I}$. From (189), we can see f is the optimal solution of (6). On the other hand, since $f \in \mathcal{I}$ and $\mathbb{E}Y^2 < \infty$, we can verify that $\mathbb{E}\Lambda^2 < \infty$ and hence $\Lambda \in \mathcal{P}_2(\mathbb{R})$. In conclusion, for any $f \in \mathcal{I}$, we can always choose $\Lambda \sim |Y| - f(|Y|)$ satisfying $\Lambda \in \mathcal{P}_2(\mathbb{R})$, such that f is the optimal solution of (6) (when $\tau = 1$). By Proposition 1, this means $f(y) = \eta(y; \mu_Y, \mu_\Lambda)$ and hence $f \in \mathcal{M}_{\mu_Y}$. Therefore, $\mathcal{I} \subseteq \mathcal{M}_{\mu_Y}$.

M. Auxiliary Results for Proving Proposition 4

Lemma 18: For any $\sigma > 0$, we have: (I) optimization problem (28) is convex and always has a unique optimal solution $f_\sigma \in \mathcal{I}$, (II) $\mathcal{L}(\sigma)$ defined in (28) is continuous at σ .

Proof: (I) Optimization problem (28) can be equivalently written as:

$$\begin{aligned} &\min_{f \in \mathcal{I}} \mathbb{E}_{\mu_Y} [f(Y) - \mathbb{E}(B | Y)]^2 + \mathbb{E}_{\mu_Y} [\text{Var}(B | Y)] \\ &\text{s.t. } \mathbb{E}_{\mu_Y} f'(Y) \leq \delta. \end{aligned}$$

Then by the same arguments in the proof of Lemma 8, it is not hard to check for any $\sigma > 0$, it is (strongly) convex and has a unique solution $f_\sigma \in \mathcal{I}$.

(II) Next, we prove the continuity of $\mathcal{L}(\sigma)$ at any $\sigma > 0$. Define the following set

$$\mathcal{I}_\sigma \stackrel{\text{def}}{=} \{f \mid f \in \mathcal{I} \text{ and } \mathbb{E}f'(B + \sigma H) \leq \delta\}, \quad (190)$$

where $H \sim \mathcal{N}(0, 1)$ and $B \sim \mu_B$ are independent. Note that for any $\sigma > 0$, we have $\mathcal{I}_\sigma \neq \emptyset$, since $\{f = 0\} \in \mathcal{I}_\sigma$.

The first step is to show for any $\sigma, r > 0$, there exists $\varepsilon \in (0, \sigma/2)$ such that whenever $\hat{\sigma} \in \mathcal{B}_\varepsilon(\sigma)$ and $\hat{f} \in \mathcal{I}_{\hat{\sigma}}$, we can always find a $f \in \mathcal{I}_\sigma$ satisfying $|f(x) - \hat{f}(x)| < r$ almost everywhere on $\mathbb{R}$. This can be proved as follows. If $\hat{f} \in \mathcal{I}_\sigma$, we can choose $f = \hat{f}$, which trivially satisfies $|f(x) - \hat{f}(x)| < r$; if $\hat{f} \notin \mathcal{I}_\sigma$, then $\mathbb{E}\hat{f}'(B + \sigma H) > \delta$. Since $\hat{f} \in \mathcal{I}_{\hat{\sigma}} \subseteq \mathcal{I}$, it follows that $|\hat{f}'| \leq 1$ almost everywhere on $\mathbb{R}$. Meanwhile, since $\sigma, \hat{\sigma} > 0$, by the properties

of Gaussian convolution we know both $B + \hat{\sigma}H$ and $B + \sigma H$ have smooth density functions supported on $\mathbb{R}$. Denote their density functions as q_1 and q_2 . Then we have

$$\begin{aligned}
|\mathbb{E}\hat{f}'(B + \hat{\sigma}H) - \mathbb{E}\hat{f}'(B + \sigma H)| &= \left| \int \hat{f}'(y)[q_1(y) - q_2(y)]dy \right| \\
&\leq \int |q_1(y) - q_2(y)|dy \\
&= 2\text{TV}(\mu_{B+\hat{\sigma}H}, \mu_{B+\sigma H}) \\
&\stackrel{(a)}{\leq} 2\text{TV}(\mu_{\hat{\sigma}H}, \mu_{\sigma H}) \\
&\stackrel{(b)}{\leq} \sqrt{2\text{KL}(\mu_{\hat{\sigma}H}, \mu_{\sigma H})} \\
&\stackrel{(c)}{\leq} C|\hat{\sigma} - \sigma|,
\end{aligned} \tag{191}$$

where $\text{TV}(\cdot, \cdot)$ and $\text{KL}(\cdot, \cdot)$ denote the total variation distance and KL-divergence between two probability measures and $C > 0$ is a fixed constant only depending on σ . In reaching (191), (b) follows from Pinsker's inequality and (c) follows from standard results of KL-divergence between Gaussian random variables and the fact that $\varepsilon \in (0, \sigma/2)$ and $\hat{\sigma} \in \mathcal{B}_\varepsilon(\sigma)$. Step (a) in (191) can be obtained as follows. Recall that $\Pi(\mu_1, \mu_2)$ denotes the set of all couplings between measures μ_1 and μ_2 . Then for any $(\hat{\sigma}H_1, \sigma H_2) \in \Pi(\mu_{\hat{\sigma}H}, \mu_{\sigma H})$ and $B_0 \sim \mu_B$ independent of $(\hat{\sigma}H_1, \sigma H_2)$, we have $(B_0 + \hat{\sigma}H_1, B_0 + \sigma H_2) \in \Pi(\mu_{B+\hat{\sigma}H}, \mu_{B+\sigma H})$. Therefore,

$$\begin{aligned}
\text{TV}(\mu_{B+\hat{\sigma}H}, \mu_{B+\sigma H}) &\stackrel{(a)}{=} \inf_{(Y_1, Y_2) \in \Pi(\mu_{B+\hat{\sigma}H}, \mu_{B+\sigma H})} \mathbb{P}(Y_1 \neq Y_2) \\
&\leq \inf_{\substack{(\hat{\sigma}H_1, \sigma H_2) \in \Pi(\mu_{\hat{\sigma}H}, \mu_{\sigma H}) \\ B_0 \sim \mu_B \text{ indep. of } (\hat{\sigma}H_1, \sigma H_2)}} \mathbb{P}(B_0 + \hat{\sigma}H_1 \neq B_0 + \sigma H_2) \\
&= \inf_{(\hat{\sigma}H_1, \sigma H_2) \in \Pi(\mu_{\hat{\sigma}H}, \mu_{\sigma H})} \mathbb{P}(\hat{\sigma}H_1 \neq \sigma H_2) \\
&\stackrel{(b)}{=} \text{TV}(\mu_{\hat{\sigma}H}, \mu_{\sigma H}),
\end{aligned} \tag{192}$$

where (a) and (b) follow from Strassen's Theorem [56, p.7]. From (191), when $\varepsilon \in (0, \sigma/2)$ and $\hat{\sigma} \in \mathcal{B}_\varepsilon(\sigma)$,

$$|\mathbb{E}\hat{f}'(B + \hat{\sigma}H) - \mathbb{E}\hat{f}'(B + \sigma H)| \leq C\varepsilon. \tag{193}$$

Since $\hat{f} \in \mathcal{I}_{\hat{\sigma}}$, from (190) and (193) we get

$$\mathbb{E}\hat{f}'(B + \sigma H) \leq \delta + C\varepsilon. \tag{194}$$

We now slightly shrink $\hat{f}$ to obtain the $f \in \mathcal{I}_\sigma$ satisfying $|f(x) - \hat{f}(x)| < r$ almost everywhere. On one hand, we know $B + \sigma H$ has a density function supported on $\mathbb{R}$ for $\sigma > 0$. On the other hand, since $\mathbb{E}\hat{f}'(B + \sigma H) > \delta$ (due to $\hat{f} \notin \mathcal{I}_\sigma$) and $|\hat{f}'| \leq 1$ almost everywhere, it is not hard to show if $\varepsilon \leq \frac{\delta}{2C}$, there exists some $\mathcal{A} \subset \mathbb{R}_{>0}$ such that $\mathbb{E}[\hat{f}'(Y)\mathbb{I}_{Y \in \mathcal{A}}] = C\varepsilon$. Accordingly, we set

$$f'(y) = \begin{cases} 0 & \pm y \in \mathcal{A}, \\ \hat{f}'(y) & \text{otherwise} \end{cases}$$

and choose $f(y)$ to be:

$$f(y) = \int_0^y f'(t)dt. \tag{195}$$

With this choice, from (194) we have $\mathbb{E}f'(Y) \leq \delta - C\varepsilon$ and thus $f \in \mathcal{I}_\sigma$. In addition, $0 \leq \widehat{f}(x) - f(x) \leq 2C\varepsilon$ almost everywhere. Hence we can choose $\varepsilon \leq \frac{r}{3C}$ so that $0 \leq \widehat{f}(x) - f(x) < r$ almost everywhere. Summing up, for any $\sigma, r > 0$, there exists C depending only on σ such that whenever $\varepsilon \leq \min\{\frac{r}{3C}, \frac{\sigma}{2}\}$, $\widehat{\sigma} \in \mathcal{B}_\varepsilon(\sigma)$ and $\widehat{f} \in \mathcal{I}_{\widehat{\sigma}}$, we can always find a $f \in \mathcal{I}_\sigma$ satisfying $0 \leq \widehat{f}(x) - f(x) < r$ almost everywhere on $\mathbb{R}$.

Define $\mathcal{L}(f, \sigma) \stackrel{\text{def}}{=} \mathbb{E}[f(B + \varsigma H) - B]^2$, which is the objective function in (28). For any compact interval $I_B \subseteq \mathbb{R}_{>0}$ and any $\sigma_1, \sigma_2 \in I_B$, $f \in \mathcal{I}$, we have

$$\begin{aligned} |\mathcal{L}(f, \sigma_1) - \mathcal{L}(f, \sigma_2)| &\leq \mathbb{E}|f^2(B + \sigma_1 H) - f^2(B + \sigma_2 H)| \\ &\quad + 2\mathbb{E}|B||f(B + \sigma_1 H) - f(B + \sigma_2 H)| \\ &\leq (\sqrt{\mathbb{E}(|B + \sigma_1 H| + |B + \sigma_2 H|)^2} + 2\sqrt{\mathbb{E}B^2})|\sigma_1 - \sigma_2| \\ &\leq C_1|\sigma_1 - \sigma_2|, \end{aligned} \tag{196}$$

where C_1 is a constant that only depends on I_B . Therefore, for any $f \in \mathcal{I}$, $\mathcal{L}(f, \sigma)$ is uniformly Lipschitz continuous w.r.t. σ on I_B . Consider $f_{\widehat{\sigma}}$ which is the optimal solution of (28) under $\widehat{\sigma}$. From the discussion in the last paragraph, for any $\sigma \in I_B$ and $r > 0$, if $\widehat{\sigma} \in \mathcal{B}_\varepsilon(\sigma) \cap I_B$ under small enough ε , then there exists some $f \in \mathcal{I}_\sigma$ satisfying $|f_{\widehat{\sigma}}(x) - f(x)| < r$ almost everywhere on $\mathbb{R}$. Therefore, for this f we have

$$\begin{aligned} |\mathcal{L}(f_{\widehat{\sigma}}, \widehat{\sigma}) - \mathcal{L}(f, \sigma)| &\leq |\mathcal{L}(f_{\widehat{\sigma}}, \widehat{\sigma}) - \mathcal{L}(f_{\widehat{\sigma}}, \sigma)| + |\mathcal{L}(f_{\widehat{\sigma}}, \sigma) - \mathcal{L}(f, \sigma)| \\ &\leq C_2 r, \end{aligned} \tag{197}$$

where C_2 is some constant that does not depend on ε, r and in the last step we use (196) and the fact that $|f_{\widehat{\sigma}}(x) - f(x)| < r$ almost everywhere. Since $\mathcal{L}(\widehat{\sigma}) = \mathcal{L}(f_{\widehat{\sigma}}, \widehat{\sigma})$, we have from (197)

$$\mathcal{L}(\widehat{\sigma}) \stackrel{(a)}{\geq} \mathcal{L}(f, \sigma) - C_2 r \stackrel{(b)}{\geq} \mathcal{L}(\sigma) - C_2 r,$$

where (a) follows from (197) and (b) follows from definition of $\mathcal{L}(\sigma)$ in (28). By exchanging σ and $\widehat{\sigma}$, we can also get $\mathcal{L}(\sigma) \geq \mathcal{L}(\widehat{\sigma}) - C_2 r$. In conclusion, for any compact interval $I_B \subseteq \mathbb{R}_{>0}$ and $r > 0$, there exists $\varepsilon > 0$ such that for any $\sigma, \widehat{\sigma} \in I_B$ satisfying $|\widehat{\sigma} - \sigma| \leq \varepsilon$, we have $|\mathcal{L}(\widehat{\sigma}) - \mathcal{L}(\sigma)| \leq r$. This proves the continuity of $\mathcal{L}(\sigma)$ over $\mathbb{R}_{>0}$. $\blacksquare$

Lemma 19: Equation (29) satisfies the following: (I) it always has a solution and the minimum solution $\sigma_0 \in [\sigma_w, (\sigma_w^2 + \delta^{-1}\mathbb{E}B^2)^{1/2}]$, (II) $\sigma_0 = \inf \mathcal{A}$, where set $\mathcal{A}$ is defined in (70).

Proof: (I) We first prove σ_0 always exists and $\sigma_0 \in I_B$, where $I_B := [\sigma_w, (\sigma_w^2 + \delta^{-1}\mathbb{E}B^2)^{1/2}]$. It is not hard to verify that at the boundary of I_B , $\mathcal{L}(\sigma)$ satisfies

$$\begin{cases} \mathcal{L}(\sigma) \leq \delta(\sigma^2 - \sigma_w^2) & \sigma = (\sigma_w^2 + \frac{\mathbb{E}B^2}{\delta})^{1/2}, \\ \mathcal{L}(\sigma) \geq \delta(\sigma^2 - \sigma_w^2) & \sigma = \sigma_w. \end{cases} \tag{198}$$

Indeed, when $\sigma = (\sigma_w^2 + \frac{\mathbb{E}B^2}{\delta})^{1/2}$, set $f = 0$ in (28) and $\mathbb{E}[f(B + \sigma H) - B]^2 = \mathbb{E}B^2 = \delta(\sigma^2 - \sigma_w^2)$, so $\mathcal{L}(\sigma) \leq \delta(\sigma^2 - \sigma_w^2)$. On the other hand, since $\mathcal{L}(\sigma) \geq 0$, the second inequality of (198) immediately follows. Then by (198), the continuity of $\mathcal{L}(\sigma)$ shown in Lemma 18 and the fact that $\sigma_0 \geq \sigma_w$, we know σ_0 always exists and $\sigma_0 \in I_B$.

(II) To prove $\inf \mathcal{A} = \sigma_0$, we proceed as follows:

(i) Show

$$\inf \mathcal{A} \in I_B. \quad (199)$$

(ii) Prove the following membership certificate of set $\mathcal{A}$ for those $\sigma \in I_B$:

$$\sigma \in \mathcal{A} \iff \mathcal{L}(\sigma) \leq \delta(\sigma^2 - \sigma_w^2). \quad (200)$$

It is not hard to see the above steps will imply $\inf \mathcal{A} = \sigma_0$. Indeed, combining (200) with (199) yields the following characterization of $\inf \mathcal{A}$:

$$\inf \mathcal{A} = \inf \{ \sigma \mid \sigma \in I_B \text{ and } \mathcal{L}(\sigma) \leq \delta(\sigma^2 - \sigma_w^2) \}. \quad (201)$$

By (198) and the continuity of $\mathcal{L}(\sigma)$, the infimum on the RHS of (201) is reached by σ_0 , which always exists. Therefore, $\inf \mathcal{A} = \sigma_0$. Step (i)-(ii) can be proved as follows:

[Proof of (i)] From the first equation of (69), we have $\inf \mathcal{A} \geq \sigma_w$. On the other hand, since $f = 0$, $\sigma = (\sigma_w^2 + \frac{\mathbb{E}B^2}{\delta})^{1/2}$, $\tau = 1$ is a solution of (66)-(67), from (68) we have $\sigma_{\text{opt}} \leq (\sigma_w^2 + \frac{\mathbb{E}B^2}{\delta})^{1/2}$. This together with the lower bound of σ_{opt} in (71) indicates: $\inf \mathcal{A} \leq \sigma_{\text{opt}} \leq (\sigma_w^2 + \frac{\mathbb{E}B^2}{\delta})^{1/2}$. Therefore, $\inf \mathcal{A} \in I_B$.

[Proof of (ii)] The “ $\Rightarrow$ ” direction of (200) immediately follows from the definition of $\mathcal{A}$ in (70). For the other direction, suppose we have a $\sigma \in I_B$ satisfying $\mathcal{L}(\sigma) \leq \delta(\sigma^2 - \sigma_w^2)$. Then

$$\mathbb{E}[f_\sigma(B + \sigma H) - B]^2 = \mathcal{L}(\sigma) \leq \delta(\sigma^2 - \sigma_w^2), \quad (202)$$

where the first equality is due to that $\mathcal{L}(\sigma)$ can be achieved by f_σ . Now consider the shrinkage of f_σ as: αf_σ , where $\alpha \in [0, 1]$. Clearly, for any $\alpha \in [0, 1]$, αf_σ still satisfies $\mathbb{E}[\alpha f'_\sigma(B + \sigma H)] \leq \delta$. Also we have:

$$\mathbb{E}[0 \cdot f_\sigma(B + \sigma H) - B]^2 = \mathbb{E}B^2 \geq \delta(\sigma^2 - \sigma_w^2), \quad (203)$$

where the last inequality follows from the condition that $\sigma \in I_B$. On the other hand, it can be easily checked that $\alpha \mapsto \mathbb{E}[\alpha f_\sigma(B + \sigma H) - B]^2$ is continuous, so from (203), (202), there exists $\alpha_0 \in [0, 1]$ such that $(\alpha_0 f_\sigma, \sigma)$ is a solution of (69), indicating $\sigma \in \mathcal{A}$. $\blacksquare$

N. Auxiliary Results for Proving Proposition 5

Lemma 20: For a probability measure $\mu_Y \in \mathcal{P}_2(\mathbb{R})$, define

$$\begin{aligned} \widetilde{\mathcal{M}}_{\mu_Y} &\stackrel{\text{def}}{=} \{ \eta(\cdot; \mu_Y, \mu_\Lambda) \mid \mu_\Lambda \in \mathcal{P}_2(\mathbb{R}) \text{ and if } q_0 > 0, \\ &\quad \int_t^{q_0} F_\Lambda^{-1}(u) du > \int_t^{q_0} F_{|Y|}^{-1}(u) du, \forall t \in [0, q_0) \} \end{aligned} \quad (204)$$

where $q_0 \stackrel{\text{def}}{=} \mathbb{P}(\eta(Y; \mu_Y, \mu_\Lambda) = 0)$. Then for any $\mu_Y \in \mathcal{P}_2(\mathbb{R})$, we have $\widetilde{\mathcal{M}}_{\mu_Y} = \mathcal{I}$. Correspondingly, for any $f(y) \in \mathcal{I}$, we can take μ_Λ as the law of $\max\{y_0, |Y| - f(|Y|)\}$ ($Y \sim \mu_Y$), so that $\eta(y; \mu_Y, \mu_\Lambda) = f(y)$. Here $y_0 \stackrel{\text{def}}{=} \sup_{y \geq 0} \{y \mid f(y) = 0\}$,

Proof: The proof is similar as Proposition 3. When μ_Λ is the law of $\max\{y_0, |Y| - f(|Y|)\}$, optimization (6) ($\tau = 1$) becomes:

$$\begin{aligned} \min_{g \in \mathcal{I}} & \frac{1}{2} \mathbb{E} \{ [f(|Y|) - g(|Y|)]^2 \mathbb{I}_{|Y| \geq y_0} \} + \frac{1}{2} \mathbb{E} \{ [(|Y| - y_0) - g(|Y|)]^2 \mathbb{I}_{|Y| < y_0} \} \\ & + \frac{1}{2} \mathbb{E} Y^2 - \mathbb{E} [|Y| - f(|Y|)]^2. \end{aligned} \quad (205)$$

Since the feasible set in (205) is $\mathcal{I}$, we know the optimal solution of (205) is exactly f . Recall that $\eta(y; \mu_Y, \mu_\Lambda)$ is the optimal solution of (6) ($\tau = 1$), so we have $\eta(y; \mu_Y, \mu_\Lambda) = f(y)$.

Finally, we show the law of $\max\{y_0, |Y| - f(|Y|)\}$ satisfies the constraint in (204). Since $\mu_Y \in \mathcal{P}_2(\mathbb{R})$ and $f \in \mathcal{I}$, we can easily get $\mu_\Lambda \in \mathcal{P}_2(\mathbb{R})$. On the other hand, suppose $q_0 > 0$. For any $t \in [0, q_0)$, we have

$$\begin{aligned} \int_t^{q_0} F_{|Y|}^{-1}(u) du &= \mathbb{E}(\mathbb{I}_{F_{|Y|}^{-1}(t) \leq |Y| \leq F_{|Y|}^{-1}(q_0)} \cdot |Y|) \\ &\stackrel{(a)}{\leq} \mathbb{E}(\mathbb{I}_{F_{|Y|}^{-1}(t) \leq |Y| \leq y_0} \cdot |Y|) \\ &\stackrel{(b)}{<} y_0 \mathbb{P}(F_{|Y|}^{-1}(t) \leq |Y| \leq y_0) \\ &= y_0(q_0 - t) \\ &\stackrel{(c)}{=} \int_t^{q_0} F_\Lambda^{-1}(u) du, \end{aligned}$$

where in (a) we use the fact that $F_{|Y|}^{-1}(q_0) \leq y_0$, since $\eta(y; \mu_Y, \mu_\Lambda) = f(y)$ and $q_0 = \mathbb{P}(\eta(Y; \mu_Y, \mu_\Lambda) = 0)$, (b) follows from $F_{|Y|}^{-1}(t) < F_{|Y|}^{-1}(q_0) = y_0$, since $t < q_0$ and (c) is due to our choice of Λ , which yields $F_\Lambda^{-1}(u) = y_0$ for any $0 \leq u \leq q_0$. ■

Lemma 21: For $\alpha \in [0, 1]$ and $\sigma > 0$, we have: (I) optimization problem (41) is convex and always has a unique optimal solution $f_{\alpha, \sigma} \in \mathcal{I}$, (II) $\mathcal{L}_\alpha(\sigma)$ defined in (41) is a continuous function over $\mathbb{R}_{>0}$, (III) equation $\mathcal{L}_\alpha(\sigma) = \delta(\sigma^2 - \sigma_w^2)$ always has a solution and the minimum solution $\sigma_{0, \alpha} \in [\sigma_w, \sqrt{\sigma_w^2 + \delta^{-1} \mathbb{E} B^2}]$.

Proof: (I) Comparing with (28) and (41), we can find the only difference is that in (41), we add a constraint $f \in \mathcal{F}_{\alpha, \sigma}$. It is not hard to check for any $\alpha \in [0, 1]$ and $\sigma > 0$, the set $\mathcal{F}_{\alpha, \sigma}$ is convex and closed in the L^2 space $\mathcal{H}_{\mu_Y}$ [definition can be found in (84)], so the uniqueness of optimal solution of (41) still holds using the same arguments.

(II) The case of $\alpha = 0$ or 1 is easy. When $\alpha = 1$, we have $\mathcal{I} \subseteq \mathcal{F}_{\alpha, \sigma}$, so $\mathcal{L}_\alpha(\sigma) = \mathcal{L}(\sigma)$ and its continuity is proved in the last part of Lemma 18; when $\alpha = 0$, $\mathcal{F}_{\alpha, \sigma}$ contains only one element: $f(x) = 0$ and $\mathcal{L}_\alpha(\sigma) = \mathbb{E} B^2$, which is trivially continuous. Therefore, it only remains to verify for the case when $\alpha \in (0, 1)$.

The proof for case $\alpha \in (0, 1)$ is similar to the proof of continuity of $\mathcal{L}(\sigma)$ in Lemma 18. For any $\alpha \in (0, 1)$ and $\sigma > 0$, define the following set

$$\mathcal{I}_{\alpha, \sigma} \stackrel{\text{def}}{=} \{f \mid f \in \mathcal{I} \cap \mathcal{F}_{\alpha, \sigma} \text{ and } \mathbb{E} f'(B + \sigma H) \leq \delta\}, \quad (206)$$

where $H \sim \mathcal{N}(0, 1)$ and $B \sim \mu_B$ are independent. We always have $\mathcal{I}_{\alpha, \sigma} \neq \emptyset$, since $\{f = 0\} \in \mathcal{I}_{\alpha, \sigma}$. Next we show that for any $\alpha \in (0, 1)$ and $\sigma, r > 0$, there exists $\varepsilon \in (0, \sigma/2)$ such that whenever $\hat{\sigma} \in \mathcal{B}_\varepsilon(\sigma)$ and $\hat{f} \in \mathcal{I}_{\alpha, \hat{\sigma}}$,

we can always find a $f \in \mathcal{I}_{\alpha, \sigma}$ satisfying $|f(x) - \hat{f}(x)| < r$ almost everywhere on $\mathbb{R}$. First, for any $\hat{\sigma} > 0$ and $\hat{f} \in \mathcal{I}_{\alpha, \hat{\sigma}}$, we have $|\hat{f}(y)| = 0$, when $|y| \leq \Phi^{-1}(1 - \frac{\alpha}{2})\hat{\sigma}$. We can then apply the following shrinkage to $\hat{f}$:

$$\hat{f}_T(y) = \text{sign}(y) \max\{0, \hat{f}(|y|) - \Phi^{-1}(1 - \frac{\alpha}{2})|\sigma - \hat{\sigma}|\}. \quad (207)$$

One can check $\hat{f}_T$ in (207) satisfies: $\hat{f}_T \in \mathcal{I}_{\alpha, \hat{\sigma}} \cap \mathcal{F}_{\alpha, \sigma}$ and

$$0 \leq \hat{f}(y) - \hat{f}_T(y) \leq \Phi^{-1}(1 - \frac{\alpha}{2})|\sigma - \hat{\sigma}| \quad (208)$$

almost everywhere on $\mathbb{R}$. For small enough $\varepsilon > 0$ and $\hat{\sigma} \in \mathcal{B}_\varepsilon(\sigma)$, we can then follow the same steps leading to (195) in the proof of continuity of $\mathcal{L}(\sigma)$ in Lemma 18 to obtain a $f \in \mathcal{I}$ satisfying

$$0 \leq \hat{f}_T(y) - f(y) \leq \frac{r}{2} \quad (209)$$

almost everywhere on $\mathbb{R}$ and $\mathbb{E}f'(B + \sigma H) \leq \delta$. Meanwhile, since $\hat{f}_T \in \mathcal{F}_{\alpha, \sigma}$, we also have $f \in \mathcal{F}_{\alpha, \sigma}$. As a result, $f \in \mathcal{I}_{\alpha, \sigma}$. Besides, combining (208) and (209) we have for small enough $\varepsilon > 0$ and $\hat{\sigma} \in \mathcal{B}_\varepsilon(\sigma)$, $0 \leq \hat{f}(y) - f(y) \leq r$ almost everywhere on $\mathbb{R}$. The remaining proof is completely same as the last part of the proof of $\mathcal{L}(\sigma)$'s continuity and we omit the details here.

(III) The proof is the same as Lemma 19. Using the same argument, it can be verified that $\mathcal{L}_\alpha(\sigma)$ also satisfies

$$\begin{cases} \mathcal{L}_\alpha(\sigma) \leq \delta(\sigma^2 - \sigma_w^2) & \sigma = (\sigma_w^2 + \frac{\mathbb{E}B^2}{\delta})^{1/2}, \\ \mathcal{L}_\alpha(\sigma) \geq \delta(\sigma^2 - \sigma_w^2) & \sigma = \sigma_w. \end{cases}$$

Then by the continuity of $\mathcal{L}_\alpha(\sigma)$ and the fact that $\sigma_{0, \alpha} \geq \sigma_w$, we get the desired result. $\blacksquare$

Lemma 22: For any given $\alpha \in [0, 1]$, we have the following.

- (a) It is always true that $\sigma_{\text{opt}, \alpha} \geq \sigma_{0, \alpha}$.
- (b) If $\delta^{-1} \mathbb{E}[f'_\alpha(Y_{0, \alpha})] < 1$, then $\sigma_{\text{opt}, \alpha} = \sigma_{0, \alpha}$ and the infimum of (73) can be achieved by choosing $\mu_\Lambda = \mu_{\text{opt}, \alpha}$.

Proof: The proof is similar to that of Proposition 4 (c). Recall that a key step in the proof is the conversion of the optimization over μ_Λ into an equivalent optimization over realizable limiting scalar functions f . We want to adopt the same strategy here, but since an additional constraint on μ_Λ is added, we need to determine the new realizable set of f , as we did in Proposition 3. It turns out that the new realizable set is still $\mathcal{I}$. This result is proved in Lemma 20 and it enables us to follow the same steps leading to (68) to show that (73) is equivalent to

$$\sigma_{\text{opt}, \alpha} = \inf\{\sigma \mid (\sigma, \tau) \in \mathcal{D}_F(\alpha), \text{ for some } \tau > 0\}, \quad (210)$$

where $\mathcal{D}_F(\alpha)$ is defined as:

$$\mathcal{D}_F(\alpha) \stackrel{\text{def}}{=} \{(\sigma, \tau) \in \mathbb{R}_{>0}^2 : \exists f \in \mathcal{I} \cap \mathcal{F}_{\alpha, \sigma} \text{ s.t. } (f, \sigma, \tau) \text{ satisfies (66)-(67)}\}.$$

Comparing (68) and (210), it can be seen that the only difference is that in (210) we require $f \in \mathcal{F}_{\alpha, \sigma}$, which is needed to control the type-I error level. Then similar as (70) we define

$$\mathcal{A}(\alpha) \stackrel{\text{def}}{=} \{\sigma \in \mathbb{R}_{>0} : \exists f \in \mathcal{I} \cap \mathcal{F}_{\alpha, \sigma}, \text{ s.t. } (f, \sigma) \text{ satisfies (69)}\}$$

and it holds that $\sigma_{\text{opt}, \alpha} \geq \inf \mathcal{A}(\alpha)$. Meanwhile, by the same reasoning in Lemma 19, we can also show $\inf \mathcal{A}(\alpha) = \sigma_{0, \alpha}$. This gives us $\sigma_{\text{opt}, \alpha} \geq \sigma_{0, \alpha}$.

Now consider the scenario where $\delta^{-1}\mathbb{E}[f'_\alpha(Y_{0,\alpha})] < 1$. In this case, we have $\tau_{0,\alpha} \in (0, \infty)$. Then it is not hard to check $(f_\alpha, \sigma_{0,\alpha}, \tau_{0,\alpha})$ satisfies equation (66)-(67). Therefore, $(\sigma_{0,\alpha}, \tau_{0,\alpha}) \in \mathcal{D}_F(\alpha)$. By (210), we have $\sigma_{0,\alpha} \geq \sigma_{\text{opt},\alpha}$. Since we already get $\sigma_{\text{opt},\alpha} \geq \sigma_{0,\alpha}$, we can conclude that $\sigma_{\text{opt},\alpha} = \sigma_{0,\alpha}$. Let us denote (σ_*, τ_*) as the solution of fixed-point equation (13)-(14), when $\mu_\Lambda = \mu_{\text{opt},\alpha}$. Using Lemma 20, it is not hard to check

$$(\sigma_*, \tau_*) = (\sigma_{0,\alpha}, \tau_{0,\alpha}) \quad (211)$$

and

$$\eta(y; \mu_{Y_*}, \mu_{\tau_*\Lambda}) = f_\alpha(y). \quad (212)$$

Meanwhile, we have $\mu_{\text{opt},\alpha} \in \tilde{\mathcal{P}}_\Lambda$. Recall that the infimum of σ_* in (73) is $\sigma_{0,\alpha}$ [c.f. (210)]. As a result, the infimum of σ_* in (73) is reached when $\mu_\Lambda = \mu_{\text{opt},\alpha}$. ■

Lemma 23: For $\alpha \in (0, 1]$, if $\delta^{-1}\mathbb{E}[f'_\alpha(Y_{0,\alpha})] < 1$ and $y_{0,\alpha} = \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{0,\alpha}$, then $\overline{\mathcal{P}}(\alpha) = \mathcal{P}(\alpha)$ and $\lim_{p \rightarrow \infty} \text{Power} = \mathcal{P}(\alpha)$, when $\mu_\Lambda = \mu_{\text{opt},\alpha}$.

Proof: In the proof, let (σ_*, τ_*) be the solution of fixed-point equation (13)-(14), when $\mu_\Lambda = \mu_{\text{opt},\alpha}$. Also denote $y_{\text{th}}^* := \sup_{y \geq 0} \{y \mid \eta(y; \mu_{Y_*}, \mu_{\tau_*\Lambda}) = 0\}$.

Assume $\delta^{-1}\mathbb{E}[f'_\alpha(Y_{0,\alpha})] < 1$ and $y_{0,\alpha} = \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{0,\alpha}$. Recall from the proof of Lemma 22, when $\mu_\Lambda = \mu_{\text{opt},\alpha}$, we have

$$y_{\text{th}}^* \stackrel{(a)}{=} y_{0,\alpha} \stackrel{(b)}{=} \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{0,\alpha} \stackrel{(c)}{=} \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_*,$$

where (a) follows from (212), (b) follows from assumption $y_{0,\alpha} = \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_{0,\alpha}$ and (c) follows from (211). Therefore,

$$\begin{aligned} \mathbb{P}(|B + \sigma_* H| \geq y_{\text{th}}^* \mid B \neq 0) &= \mathbb{P}(|B + \sigma_* H| \geq \Phi^{-1}(1 - \frac{\alpha}{2})\sigma_* \mid B \neq 0) \\ &= \overline{\mathcal{P}}(\alpha), \end{aligned} \quad (213)$$

where the last equality is due to that $\mu_{\text{opt},\alpha}$ is the optimal solution of (73), as is proved by Lemma 22. Therefore, from (213) we know when $\mu_\Lambda = \mu_{\text{opt},\alpha}$, the objective value of (40) equals to $\overline{\mathcal{P}}(\alpha)$. This implies $\mathcal{P}(\alpha) \geq \overline{\mathcal{P}}(\alpha)$, since $\mathcal{P}(\alpha)$ is the optimal value of (40). Combined with the fact that $\overline{\mathcal{P}}(\alpha)$ is the upper bound of $\mathcal{P}(\alpha)$ [c.f. (76)], it then follows that $\overline{\mathcal{P}}(\alpha) = \mathcal{P}(\alpha)$. Also when $\mu_\Lambda = \mu_{\text{opt},\alpha}$, $\lim_{p \rightarrow \infty} \text{Power} = \overline{\mathcal{P}}(\alpha) = \mathcal{P}(\alpha)$. ■

REFERENCES

- [1] H. Hu and Y. M. Lu, "Asymptotics and optimal designs of SLOPE for sparse linear regression," *2019 IEEE International Symposium on Information Theory (ISIT)*, pp. 375–379, 2019.
- [2] M. Bogdan, E. Van Den Berg, C. Sabatti, W. Su, and E. J. Candès, "SLOPE-adaptive variable selection via convex optimization," *The annals of applied statistics*, vol. 9, no. 3, p. 1103, 2015.
- [3] L. W. Zhong and J. T. Kwok, "Efficient sparse modeling with automatic feature grouping," *IEEE transactions on neural networks and learning systems*, vol. 23, no. 9, pp. 1436–1447, 2012.
- [4] X. Zeng and M. A. Figueiredo, "Decreasing weighted sorted ℓ_1 regularization," *IEEE Signal Processing Letters*, vol. 21, no. 10, pp. 1240–1244, 2014.
- [5] H. D. Bondell and B. J. Reich, "Simultaneous regression shrinkage, variable selection, and supervised clustering of predictors with OSCAR," *Biometrics*, vol. 64, no. 1, pp. 115–123, 2008.

- [6] M. Figueiredo and R. Nowak, “Ordered weighted ℓ_1 regularized regression with strongly correlated covariates: Theoretical aspects,” in *Artificial Intelligence and Statistics*, 2016, pp. 930–938.
- [7] W. Su, E. Candes *et al.*, “SLOPE is adaptive to unknown sparsity and asymptotically minimax,” *The Annals of Statistics*, vol. 44, no. 3, pp. 1038–1068, 2016.
- [8] P. C. Bellec, G. Lecué, A. B. Tsybakov *et al.*, “SLOPE meets LASSO: improved oracle bounds and optimality,” *The Annals of Statistics*, vol. 46, no. 6B, pp. 3603–3642, 2018.
- [9] M. Bogdan, E. v. d. Berg, W. Su, and E. Candes, “Statistical estimation and testing via the sorted ℓ_1 norm,” *arXiv preprint arXiv:1310.1969*, 2013.
- [10] M. Bayati and A. Montanari, “The LASSO risk for gaussian matrices,” *IEEE Transactions on Information Theory*, vol. 58, no. 4, pp. 1997–2017, 2012.
- [11] W. Su, M. Bogdan, E. Candes *et al.*, “False discoveries occur early on the LASSO path,” *The Annals of Statistics*, vol. 45, no. 5, pp. 2133–2150, 2017.
- [12] S. Oymak, C. Thrampoulidis, and B. Hassibi, “The squared-error of generalized lasso: A precise analysis,” in *2013 51st Annual Allerton Conference on Communication, Control, and Computing (Allerton)*. IEEE, 2013, pp. 1002–1009.
- [13] N. El Karoui, “On the impact of predictor geometry on the performance on high-dimensional ridge-regularized generalized robust regression estimators,” *Probability Theory and Related Fields*, vol. 170, no. 1-2, pp. 95–175, 2018.
- [14] C. Thrampoulidis, E. Abbasi, W. Xu, and B. Hassibi, “BER analysis of the box relaxation for BPSK signal recovery,” in *2016 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2016, pp. 3776–3780.
- [15] L. Zheng, A. Maleki, H. Weng, X. Wang, and T. Long, “Does ℓ_p -minimization outperform ℓ_1 -minimization?” *IEEE Transactions on Information Theory*, vol. 63, no. 11, pp. 6896–6935, 2017.
- [16] J. Barbier, M. Dia, N. Macris, and F. Krzakala, “The mutual information in random linear estimation,” in *2016 54th Annual Allerton Conference on Communication, Control, and Computing (Allerton)*. IEEE, 2016, pp. 625–632.
- [17] G. Reeves and H. D. Pfister, “The replica-symmetric prediction for compressed sensing with gaussian matrices is exact,” in *2016 IEEE International Symposium on Information Theory (ISIT)*. IEEE, 2016, pp. 665–669.
- [18] D. Donoho and J. Tanner, “Observed universality of phase transitions in high-dimensional geometry, with implications for modern data analysis and signal processing,” *Philosophical Transactions of the Royal Society of London A: Mathematical, Physical and Engineering Sciences*, vol. 367, no. 1906, pp. 4273–4293, 2009.
- [19] D. L. Donoho, A. Maleki, and A. Montanari, “Message-passing algorithms for compressed sensing,” *Proceedings of the National Academy of Sciences*, vol. 106, no. 45, pp. 18 914–18 919, 2009.
- [20] M. Bayati and A. Montanari, “The dynamics of message passing on dense graphs, with applications to compressed sensing,” *IEEE Transactions on Information Theory*, vol. 57, no. 2, pp. 764–785, 2011.
- [21] V. Chandrasekaran, B. Recht, P. A. Parrilo, and A. S. Willsky, “The convex geometry of linear inverse problems,” *Foundations of Computational mathematics*, vol. 12, no. 6, pp. 805–849, 2012.
- [22] F. Krzakala, M. Mézard, F. Sausset, Y. Sun, and L. Zdeborová, “Statistical-physics-based reconstruction in compressed sensing,” *Physical Review X*, vol. 2, no. 2, p. 021005, 2012.
- [23] A. Javanmard and A. Montanari, “State evolution for general approximate message passing algorithms, with applications to spatial coupling,” *Information and Inference: A Journal of the IMA*, vol. 2, no. 2, pp. 115–144, 2013.
- [24] N. El Karoui, D. Bean, P. J. Bickel, C. Lim, and B. Yu, “On robust regression with high-dimensional predictors,” *Proceedings of the National Academy of Sciences*, vol. 110, no. 36, pp. 14 557–14 562, 2013.
- [25] D. Amelunxen, M. Lotz, M. B. McCoy, and J. A. Tropp, “Living on the edge: Phase transitions in convex programs with random data,” *Information and Inference: A Journal of the IMA*, vol. 3, no. 3, pp. 224–294, 2014.
- [26] C. Thrampoulidis, S. Oymak, and B. Hassibi, “Regularized linear regression: A precise analysis of the estimation error,” in *Conference on Learning Theory*, 2015, pp. 1683–1709.
- [27] D. Donoho and A. Montanari, “High dimensional robust m-estimation: Asymptotic variance via approximate message passing,” *Probability Theory and Related Fields*, vol. 166, no. 3, pp. 935–969, 2016.
- [28] O. Dhifallah, C. Thrampoulidis, and Y. M. Lu, “Phase retrieval via linear programming: Fundamental limits and algorithmic improvements,” in *2017 55th Annual Allerton Conference on Communication, Control, and Computing (Allerton)*. IEEE, 2017, pp. 1071–1077.
- [29] C. Thrampoulidis, E. Abbasi, and B. Hassibi, “Precise error analysis of regularized M -estimators in high-dimensions,” *IEEE Transactions on Information Theory*, 2018.

- [30] P. Sur and E. J. Candès, “A modern maximum-likelihood theory for high-dimensional logistic regression,” *Proceedings of the National Academy of Sciences*, vol. 116, no. 29, pp. 14 516–14 525, 2019.
- [31] T. Tanaka, “A statistical-mechanics approach to large-system analysis of cdma multiuser detectors,” *IEEE Transactions on Information theory*, vol. 48, no. 11, pp. 2888–2910, 2002.
- [32] Y. Kabashima, T. Wadayama, and T. Tanaka, “A typical reconstruction limit for compressed sensing based on ℓ_p -norm minimization,” *Journal of Statistical Mechanics: Theory and Experiment*, vol. 2009, no. 09, p. L09003, 2009.
- [33] Y. Gordon, “Some inequalities for gaussian processes and applications,” *Israel Journal of Mathematics*, vol. 50, no. 4, pp. 265–289, 1985.
- [34] M. Stojnic, “A framework to characterize performance of lasso algorithms,” *arXiv preprint arXiv:1303.7291*, 2013.
- [35] F. Salehi, E. Abbasi, and B. Hassibi, “A precise analysis of phasemax in phase retrieval,” in *2018 IEEE International Symposium on Information Theory (ISIT)*. IEEE, 2018, pp. 976–980.
- [36] Z. Deng, A. Kammoun, and C. Thrampoulidis, “A model of double descent for high-dimensional binary linear classification,” *arXiv preprint arXiv:1911.05822*, 2019.
- [37] A. Montanari, F. Ruan, Y. Sohn, and J. Yan, “The generalization error of max-margin linear classifiers: High-dimensional asymptotics in the overparametrized regime,” *arXiv preprint arXiv:1911.01544*, 2019.
- [38] G. R. Kini and C. Thrampoulidis, “Analytic study of double descent in binary classification: The impact of loss,” in *2020 IEEE International Symposium on Information Theory (ISIT)*. IEEE, 2020, pp. 2527–2532.
- [39] D. Bean, P. J. Bickel, N. El Karoui, and B. Yu, “Optimal m-estimation in high-dimensional regression,” *Proceedings of the National Academy of Sciences*, vol. 110, no. 36, pp. 14 563–14 568, 2013.
- [40] H. Taheri, R. Pedarsani, and C. Thrampoulidis, “Sharp guarantees and optimal performance for inference in binary and gaussian-mixture models,” *Entropy*, vol. 23, no. 2, p. 178, 2021.
- [41] M. Advani and S. Ganguli, “Statistical mechanics of optimal convex inference in high dimensions,” *Physical Review X*, vol. 6, no. 3, p. 031034, 2016.
- [42] B. Aubin, F. Krzakala, Y. M. Lu, and L. Zdeborová, “Generalization error in high-dimensional perceptrons: Approaching bayes error with convex optimization,” *arXiv preprint arXiv:2006.06560*, 2020.
- [43] H. Taheri, R. Pedarsani, and C. Thrampoulidis, “Fundamental limits of ridge-regularized empirical risk minimization in high dimensions,” in *International Conference on Artificial Intelligence and Statistics*. PMLR, 2021, pp. 2773–2781.
- [44] H. Weng, A. Maleki, L. Zheng *et al.*, “Overcoming the limitations of phase transition by higher order analysis of regularization techniques,” *Annals of Statistics*, vol. 46, no. 6A, pp. 3099–3129, 2018.
- [45] S. Wang, H. Weng, A. Maleki *et al.*, “Which bridge estimator is the best for variable selection?” *Annals of Statistics*, vol. 48, no. 5, pp. 2791–2823, 2020.
- [46] M. Celentano and A. Montanari, “Fundamental barriers to high-dimensional regression with convex penalties,” *arXiv preprint arXiv:1903.10603*, 2019.
- [47] Z. Bu, J. M. Klusowski, C. Rush, and W. J. Su, “Algorithmic analysis and statistical estimation of slope via approximate message passing,” *IEEE Transactions on Information Theory*, vol. 67, no. 1, pp. 506–537, 2020.
- [48] S. Wang, H. Weng, and A. Maleki, “Does slope outperform bridge regression?” *arXiv preprint arXiv:1909.09345*, 2019.
- [49] M. Celentano, “Approximate separability of symmetrically penalized least squares in high dimensions: characterization and consequences,” *arXiv preprint arXiv:1906.10319*, 2019.
- [50] L. Miolane and A. Montanari, “The distribution of the Lasso: Uniform control over sparse balls and adaptive parameter tuning,” *arXiv preprint arXiv:1811.01212*, 2018.
- [51] A. Mousavi, A. Maleki, and R. G. Baraniuk, “Consistent parameter estimation for LASSO and approximate message passing,” *The Annals of Statistics*, vol. 46, no. 1, pp. 119–148, 2018.
- [52] D. L. Donoho, A. Maleki, and A. Montanari, “The noise-sensitivity phase transition in compressed sensing,” *IEEE Transactions on Information Theory*, vol. 57, no. 10, pp. 6920–6941, 2011.
- [53] A. W. Van der Vaart, *Asymptotic statistics*. Cambridge university press, 2000, vol. 3.
- [54] A. Montanari and P.-M. Nguyen, “Universality of the elastic net error,” in *2017 IEEE International Symposium on Information Theory (ISIT)*. IEEE, 2017, pp. 2338–2342.
- [55] A. Panahi and B. Hassibi, “A universal analysis of large-scale regularized least squares solutions,” in *Advances in Neural Information Processing Systems*, 2017, pp. 3381–3390.
- [56] C. Villani, *Topics in optimal transportation*. American Mathematical Soc., 2003, no. 58.

- [57] H. L. Royden, *Real analysis*. Prentice Hall, 2010.
- [58] T.-C. Hu and R. Taylor, “On the strong law for arrays and for the bootstrap mean and variance,” *International Journal of Mathematics and Mathematical Sciences*, vol. 20, no. 2, pp. 375–382, 1997.
- [59] M. Sion, “On general minimax theorems,” *Pacific Journal of mathematics*, vol. 8, no. 1, pp. 171–176, 1958.
- [60] D. P. Bertsekas, *Convex optimization theory*. Athena Scientific, 2009.
- [61] C. D. Aliprantis and K. C. Border, *Infinite Dimensional Analysis: a Hitchhiker's Guide*. Springer, 2006.
- [62] R. T. Rockafellar and R. J.-B. Wets, *Variational analysis*. Springer Science & Business Media, 2009, vol. 317.
- [63] W. Rudin *et al.*, *Principles of mathematical analysis*. McGraw-hill New York, 1976, vol. 3, no. 4.2.