

ON THE SPARSITY OF LASSO MINIMIZERS IN SPARSE DATA RECOVERY

SIMON FOUCART, EITAN TADMOR, AND MING ZHONG

Dedicated to Ron DeVore with friendship and admiration

ABSTRACT. We present a detailed analysis of the unconstrained ℓ_1 -weighted LASSO method for recovery of sparse data from its observation by randomly generated matrices, satisfying the Restricted Isometry Property (RIP) with constant $\delta < 1$, and subject to negligible measurement and compressibility errors. We prove that if the data is k -sparse, then the size of support of the LASSO minimizer, s , maintains a comparable sparsity, $s \leq C_\delta k$. For example, if $\delta = 0.7$ then $s < 11k$ and a slightly smaller $\delta = 0.4$ yields $s < 4k$. We also derive new ℓ_2/ℓ_1 error bounds which highlight precise dependence on k and on the LASSO parameter λ , before the error is driven below the scale of negligible measurement/ and compressibility errors.

CONTENTS

1. Introduction	1
1.1. Statement of main results	2
2. The Robust Null Space Property	4
3. On the sparsity of the unconstrained LASSO minimizer	5
3.1. Bounds of the sparsity	6
3.2. Numerical simulations	8
4. Error bounds	10
4.1. ℓ_2 -error bounds	10
4.2. ℓ_1 -error bound	13
4.3. Numerical simulations	15
References	16

1. INTRODUCTION

In 2006, the pioneering works of Candès, Romberg and Tao [13, 14] and of Donoho [21] suggested the framework of a constrained ℓ_1 -method to recover a sparse unknown $\mathbf{x}_* \in \mathbb{R}^N$ from its observation $\mathbf{y}_* = A\mathbf{x}_* \in \mathbb{R}^m$ ¹. The key point is that one can design observing matrices $A \in \mathbb{R}^{m \times N}$

Date: March 16, 2022.

2020 Mathematics Subject Classification. 94A12, 94A20.

Key words and phrases. Inverse problems, data recovery, compressive sensing, basis pursuit de-noising method, robust null space property.

Research was supported in part by NSF grants CCF-1934904, DMS-2053172 (SF), and DMS16-13911 (ET) and ONR grants N00014-2012787 (SF) and N00014-2112773 (ET).

¹Earlier announcement of the works was presented in the workshops, organized together with Ron DeVore, which can be found at home.cscamm.umd.edu/programs/srs05/candes_srs05.pdf and home.cscamm.umd.edu/programs/srs05/donoho_srs05.htm.

with a relatively small number of observations, $m \ll N$, such that a constrained ℓ_1 -method — also known as Basis Pursuit (BP) in [18, 16, 17] — finds a sparse solution as a minimizer of²

$$\mathbf{x}_{BP} := \arg \min_{\mathbf{x} \in \mathbb{R}^N} \{|\mathbf{x}|_1 \mid A\mathbf{x} = \mathbf{y}_*\}, \quad A \in \mathbb{R}^{m \times N}, \quad m \ll N.$$

This is closely related to the well-known LASSO algorithm introduced in 1996 in the statistics literature [39], $\arg \min_{|\mathbf{x}|_1 \leq \delta} \{|\mathbf{y}_*^\epsilon - A\mathbf{x}|_2^2\}$, which can be viewed as an ℓ_1 -penalty relaxation of a least squares subject to (possibly noisy) observation $\mathbf{y}_*^\epsilon$.

The BP minimizer, $\mathbf{x}_{BP}$, recovers the sparse $\mathbf{x}_*$ when the observing matrix A satisfies an appropriate recoverability condition; we mention here the Restricted Isometry Property (RIP) introduced in [13], the ℓ_1 -Coherence discussed in [40, 27, 22, 23], the restricted eigenvalue condition [6, §3], or the Null Space Property (NSP) of DeVore and his co-authors [19, 20], and related Robust Null Sparse Property (RNSP) of [26]. Important classes of such observing matrices with desired sparse recoverability conditions are randomly generated, e.g., [26, §9].

1.1. Statement of main results. Throughout the paper we will be using the two notions of sparsity and compressibility. A vector $\mathbf{x} \in \mathbb{R}^N$ is *sparse* if

$$s_{\mathbf{x}} := |\mathbf{x}|_0 \ll N.$$

In applications, sparsity is often difficult to acquire, and clean observations are not always available, since the observation process is inevitably and easily corrupted by errors — human and/or machine measurement errors. We turn our attention to the recovery of compressible unknown from its *noisy* observations. A vector $\mathbf{x} \in \mathbb{R}^N$ is *compressible* of order k , or simply k -compressible, if its content is faithfully captured by a k -sparse vector — specifically, if its ℓ_1 -distance to the set of all k -sparse vectors,

$$(1.1) \quad \sigma_k(\mathbf{x}) := \inf_{\mathbf{z} \in \mathbb{R}^N} \{|\mathbf{x} - \mathbf{z}|_1 : |\mathbf{z}|_0 \leq k\},$$

is small relative to $|\mathbf{x}|_1$. We note that $\sigma_k(\mathbf{x})$ is realized by a (not necessarily unique) vector, denoted $\mathbf{x}(k)$, whose non-zero entries are the k largest of $\mathbf{x}$ in absolute value.

Let $\mathbf{x}_*$ be a compressible unknown of order k so that $\sigma_k(\mathbf{x}_*) \ll |\mathbf{x}_*|_1$, and assume we only have access to its measured observation $\mathbf{y}_*^\epsilon = A\mathbf{x}_* + \epsilon$. The term ϵ is the measurement error caused by a number of factors which are assumed statistically independent of the unknown $\mathbf{x}_*$ and the observing operator A . The details of ϵ remain untraceable except for its size which is assumed to be negligibly small. In this case, one should not expect an exact recovery of a sparse $\mathbf{x}_*$, but instead, accept an approximate solution, $\mathbf{y}_*^\epsilon = A\mathbf{x}_*(k) + \epsilon'$, where $\epsilon' = A(\mathbf{x}_* - \mathbf{x}_*(k)) + \epsilon$ is adapted to the small scale built into the problem, which consists of two contributions — the small measurement error, $\epsilon := |\epsilon|_2$, and the small compressibility error,³ $|A(\mathbf{x}_* - \mathbf{x}_*(k))|_2 \leq \sigma_k(\mathbf{x}_*)$,

$$\mathbf{y}_*^\epsilon = A\mathbf{x}_*(k) + \epsilon', \quad |\epsilon'|_2 \leq \mu,$$

such that $\mu = \sigma_k(\mathbf{x}_*) + \epsilon$ is much smaller relative to the unknown data, $\mu \ll |\mathbf{x}_*|_1$. Although the observing operator A is linear, the recovery of $\mathbf{x}_*$ by a direct “solution” of the linear problem $A\mathbf{x} = \mathbf{y}_*^\epsilon$ is ill-posed, unless additional conditions on A and $\mathbf{x}_*$ are enforced so that the unknown

²Given $\mathbf{x} \in \mathbb{R}^N$ we let $|\mathbf{x}|_p$ denote its ℓ_p -norm, with the usual conventional limiting cases of $p = \infty$ and $p = 0$, where $|\mathbf{x}|_\infty := \max_{1 \leq i \leq N} |x_i|$, and respectively $|\mathbf{x}|_0 := |\text{supp}(\mathbf{x})|$ where $|\cdot|$ is the cardinality of a finite set.

³The columns of A are assumed ℓ_2 -normalized so that $|A|_{1 \rightarrow 2} = 1$.

object $\mathbf{x}_*$, or at least a faithful approximation of it, is recovered by solving an augmented *well-posed* regularized minimization problem. On the way, the original linear problem is replaced by a nonlinear procedure. To capture the compressible information of $\mathbf{x}_*$ from its noisy observation $\mathbf{y}_*^\epsilon$, we seek a minimizer of the unconstrained ℓ_1 -regularized Least Squares problem,

$$(1.2) \quad \mathbf{x}_\lambda := \arg \min_{\mathbf{x} \in \mathbb{R}^N} \left\{ \lambda |\mathbf{x}|_1 + \frac{1}{2} |\mathbf{y}_*^\epsilon - A\mathbf{x}|_2^2 \right\}, \quad A \in \mathbb{R}^{m \times N}, \quad m \ll N.$$

The unconstrained variational statement (1.2) falls under the general class of Tikhonov regularization. The distinctive feature is the ℓ_1 -regularization, leading to an approximate decomposition of the basis pursuit of Chen & Donoho [18], $\mathbf{y}_*^\epsilon = A\mathbf{x}_\lambda + \mathbf{r}_\lambda$ with (hopefully) small residual, $\mathbf{r}_\lambda = \mathbf{y}_*^\epsilon - A\mathbf{x}_\lambda$, depending on a parameter λ . This version of ℓ_1 -regularization, called “Basis Pursuit De-Noising” in [16], which became known as the unconstrained ℓ_1 -weighted LASSO, is the main focus of our work. As noted in the 1996 thesis [16], the work on this version of BP was motivated by a series of ideas using ℓ_0/ℓ_1 -based regularization that appeared in early 1990s, primarily the empirical atomic decomposition of Donoho and Johnstone [24], $\arg \min_{\mathbf{x} \in \mathbb{R}^N} \left\{ \lambda |\mathbf{x}|_0 + \frac{1}{2} |\mathbf{y}_*^\epsilon - A\mathbf{x}|_2^2 \right\}$, the multi-scale edge representation with wavelets of Hwang and Mallat [29] and the TV-based denoising method of ROF [31], $\arg \min_{\mathbf{x} \in \mathbb{R}^N} \left\{ \lambda |\mathbf{x}|_{TV} + \frac{1}{2} |\mathbf{y}_*^\epsilon - A\mathbf{x}|_2^2 \right\}$. These works were later further explored as the Lagrangian formulation of the quadratically constrained Basis Pursuit de-noising [17, 14] and the noise-aware ℓ_1 -minimization [26].

Since $\lambda > 0$ controls the distance between $A\mathbf{x}_\lambda$ and $\mathbf{y}_*$, the parameter λ can be interpreted as a regularization *scale*. In a subsequent work, [37], we pursue a *multi-scale* generalization based on a ladder of hierarchical scales constructed by the Hierarchical Decomposition (HD) method [34, 35, 36, 33]. The goal of this work is to analyze the sparsity behavior of the *mono-scale* LASSO (1.2), observed by a sub-class of RIP matrices satisfying the Robust Null Space Property (RNSP) which is discussed in section 2. Our main results, outlined and proved in section 3, are summarized in the following. Our results involve three main parameters: the Restricted Isometry Constant (RIC) in (2.2) below, $\delta = \delta_k < 1$, the related RNSP constant, $\beta_\delta = \frac{\sqrt{1+\delta}}{\sqrt{1-\delta^2} - \delta/4}$, depending on the RIC δ , and the small scale of compressibility+measurement, $\mu = \sigma_k(\mathbf{x}_*) + \epsilon$.

Theorem 1.1 (Main result). *Let $\mathbf{x}_*$ be k -compressible, and let $\mathbf{y}_*^\epsilon = A\mathbf{x}_* + \epsilon$ be its observation with observing matrix A satisfying the RIP (2.2) with constant δ large enough, $\delta > \delta_t$, such that (3.7) below holds. Let $\mathbf{x}_\lambda$ be the LASSO minimizer (1.2).*

(i) **(Sparsity).** *The sparsity of the LASSO minimizer, $s_\lambda = s_{\mathbf{x}_\lambda}$, does not exceed*

$$s_\lambda < (1 + \delta) \left(\beta_\delta \sqrt{k} + \frac{2\mu}{\lambda} \right)^2.$$

(ii) **(ℓ_2 -error bound).** *The following ℓ_2 -error bound holds*

$$\frac{1}{\sqrt{1+\delta}} \left(\frac{\sqrt{s_\lambda} \lambda}{\sqrt{1+\delta}} - \mu \right) \leq |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_2 \leq \frac{1}{\sqrt{1-\delta}} (\beta_\delta \sqrt{k} \lambda + 3\mu).$$

(iii) **(ℓ_1 -error bound).** *The following ℓ_1 -error bound holds*

$$|\mathbf{x}_\lambda - \mathbf{x}_*(k)|_1 < \frac{\sqrt{1+\delta}}{\sqrt{1-\delta}} \frac{1}{\lambda} ((\beta_\delta + 1/2) \sqrt{k} \lambda + 2\mu)^2.$$

We interpret these bounds as follows. Set $\theta = 2\mu/\sqrt{k}\lambda$, then (i) reads

$$s_\lambda \leq \chi^2 k, \quad \chi = \sqrt{1 + \delta}(\beta_\delta + \theta).$$

Thus, if $\theta \leq 1$ — namely, as long as λ does not get exceedingly small so that $\lambda \geq 2\mu/\sqrt{k}$, then the sparsity of $\mathbf{x}_\lambda$ is comparable to the sparsity of $\mathbf{x}_*$. Furthermore, in (ii) we have the ℓ_2 -error bound of order $\lesssim \sqrt{k}\lambda + \mu$ and in (iii), an ℓ_1 -error bound of order $\lesssim k\lambda + \sqrt{k}\mu$.

We conclude with a few comments on theorem 1.1. The sparsity bound in (i) with RIC $\delta = 0.7$ yields $s_\lambda < 11k$, while a slightly smaller RIC $\delta = 0.4$ yields $s_\lambda < 4k$. This should be compared with the sparsity bounds in [5, Theorem 3] and [32]. The ℓ_2 -upper bound on the right of (ii) is not new; here we recover the ℓ_2 -bound, derived under appropriate assumptions, in [11], [6, Theorem 7.1] and [30, 38, 28]. This should be contrasted with the ℓ_2 -error lower-bound on the left, derived in section 4.1 (see figure 4.1). Indeed, this ℓ_2 lower-bound is the essential ingredient in our proof of the sparsity bound in (i). Finally, the ℓ_1 -bound in (iii) with RIC $\delta = 0.7$ yields $|\mathbf{x}_\lambda - \mathbf{x}_*(k)|_1 < 16.97k\lambda + 24.04\sqrt{k}\mu$. Here, the linear decay with λ is not new and can be found for example, under various assumptions, in [6, Theorem 7.1] and [9, Theorem 6.1].

2. THE ROBUST NULL SPACE PROPERTY

Optimality of the minimizer. The variational problem (1.2) admits a minimizer, $\mathbf{x}_\lambda$, and at least for certain relevant classes of full row rank A 's, the minimizer is unique, [41]. The minimizer is completely characterized by its residual, $\mathbf{r}_\lambda := \mathbf{y}_*^\epsilon - A\mathbf{x}_\lambda$ (to simplify notations we suppress the dependence of $\mathbf{r}_\lambda$ on ϵ). We summarize the results from [35, §2.1], [33, Appendix] where we distinguish between two cases.

- (i) If $\lambda \geq \lambda_\infty := |A^\top \mathbf{y}_*^\epsilon|_\infty$ then (1.2) admits only the trivial minimizer $\mathbf{x}_\lambda \equiv \mathbf{0}$. In this case, λ is too large to extract the compressibility information in $\mathbf{y}_*^\epsilon$.
- (ii) If $\lambda < \lambda_\infty = |A^\top \mathbf{y}_*^\epsilon|_\infty$ then (1.2) admits a non-trivial minimizer, $\mathbf{x}_\lambda$, with the corresponding residual, $\mathbf{r}_\lambda = \mathbf{y}_*^\epsilon - A\mathbf{x}_\lambda$, such that $(\mathbf{x}_\lambda, \mathbf{r}_\lambda)$ forms an *extremal pair* in the sense that

$$(2.1) \quad \langle A\mathbf{x}_\lambda, \mathbf{r}_\lambda \rangle = \lambda |\mathbf{x}_\lambda|_1 \text{ and } |A^\top \mathbf{r}_\lambda|_\infty = \lambda.$$

To proceed we will need the following notations. The restriction of a vector $\mathbf{w} \in \mathbb{R}^N$ on an index set $\mathcal{K} \subset \{1, 2, \dots, N\}$ of size $k = |\mathcal{K}|$ is denoted $\mathbf{w}_\mathcal{K} := \{w_i, i \in \mathcal{K}\} \in \mathbb{R}^k$. Similarly, given a matrix $W \in \mathbb{R}^{m \times N}$ with columns $\mathbf{w}_1, \mathbf{w}_2, \dots$, its restriction on an index set $\mathcal{K}$ of size $k = |\mathcal{K}|$ consists of the k columns $W_\mathcal{K} := \text{col}\{\mathbf{w}_i, i \in \mathcal{K}\}$. The size of W can be measured by its induced matrix norm, $\|W\|_p = \sup_{|\mathbf{w}|_p=1} |W\mathbf{w}|_p$. The signum vector is defined component-wise,

$$\text{sgn}(\mathbf{w})_i = \text{sgn}(w_i), \text{ in terms of the usual signum function } \text{sgn}(w) = \begin{cases} -1, & w < 0 \\ 1, & w > 0 \end{cases} \text{ for } w \neq 0.$$

Restricted Isometry Property (RIP). A matrix A satisfies the Restricted Isometry Property (RIP) of order k with Restricted Isometry Constant (RIC) $\delta_k < 1$ if the following holds, [15, 21, 13, 7],

$$(2.2) \quad (1 - \delta_k) |\mathbf{x}|_2^2 \leq |A\mathbf{x}|_2^2 \leq (1 + \delta_k) |\mathbf{x}|_2^2, \quad \forall |\mathbf{x}|_0 \leq k.$$

Throughout the paper we adopt the usual assumption that δ_k is measured for A 's with ℓ^2 -normalized columns⁴. There are two classes of matrices $A \in \mathbb{R}^{m \times N}$ satisfying the RIP of order k : deterministic

⁴The RIP of A asserts that for any subset of its k columns, $\{\mathbf{a}_i\}_{i \in \mathcal{K}}$, the entries $|\langle \mathbf{a}_i, \mathbf{a}_j \rangle|_{i \neq j} \lesssim \delta_k$ while $|\mathbf{a}_i|_2^2 = 1 + \epsilon_i$ such that $|\epsilon_i| \lesssim \delta_k$. Therefore, one can always re-normalize the columns of A by a factor $\lesssim (1 - \delta_k)^{-1/2}$ yielding a new RIP matrix with ℓ_2 -normalized columns and with possibly slightly larger RIP constant $\delta'_k \lesssim \delta_k / (1 - \delta_k)$.

A 's with number of observations $m \gtrsim k^2$ (the quadratic bottleneck is lessened in [8]); and a large class of randomly generated A 's for which the restriction on the number of observations can be further lessened to having only m observations, [26, §9.4]

$$(2.3) \quad m \sim \text{Const} \cdot \delta^{-2} k \ln(eN/k).$$

Candès proved the exactness of the constrained BP for RIP matrices with $\delta < \sqrt{2} - 1$, [12]. Further refinements were reported in [25] before the definitive result of [10].

Robust Null Space Property (RNSP). A crucial step in quantifying the recovery error of $\mathbf{x}_*$ using (1.2) is to enforce a recoverability condition on the observing matrix A . This brings us to the Robust Null Sparse Property (RNSP) introduced in [26, §4.3]. A matrix $A \in \mathbb{R}^{m \times N}$ satisfies the RNSP of order k with constants $0 < \rho < 1$ and $\tau > 0$, if for all $\mathcal{K} \subset \{1, 2, \dots, N\}$ of size $|\mathcal{K}| \leq k$, there holds

$$(2.4) \quad |\mathbf{x}_{\mathcal{K}}|_1 \leq \rho |\mathbf{x}_{\mathcal{K}^c}|_1 + \tau |A\mathbf{x}|_2, \quad \forall \mathbf{x} \in \mathbb{R}^N.$$

We refer to the “RNSP $_{\rho, \tau}$ of order k ”, and unless needed, we suppress the dependence of (ρ, τ) on k . In particular, given a k -sparse $\mathbf{v}$ and any $\mathbf{u}$, we apply (2.4) to $\mathbf{x} = \mathbf{u} - \mathbf{v}$ with $\mathcal{K} = \text{supp}(\mathbf{v})$, where $|\mathbf{x}_{\mathcal{K}}|_1 - |\mathbf{x}_{\mathcal{K}^c}|_1 \geq |\mathbf{v}|_1 - |\mathbf{u}|_1$ yields the following useful consequence of RNSP.

Lemma 2.1. *If $A \in \mathbb{R}^{m \times N}$ satisfies the RNSP $_{\rho, \tau}$ of order k , then for all k -sparse $\mathbf{v}$'s and any $\mathbf{u}$,*

$$(2.5) \quad |\mathbf{v}|_1 - |\mathbf{u}|_1 \leq \tau |A(\mathbf{u} - \mathbf{v})|_2, \quad |\text{supp}(\mathbf{v})| \leq k.$$

As an example for the class of observation matrices satisfying the RNSP $_{\rho, \tau}$ of order k , we mention the class of randomly generated RIP matrices with RICs $\delta = \delta_{2k}$, [26, Theorem 6.13],

$$(2.6) \quad \rho = \frac{\delta}{\sqrt{1 - \delta^2} - \delta/4} \quad \text{and} \quad \tau = \beta \sqrt{k}, \quad \beta := \frac{\sqrt{1 + \delta}}{\sqrt{1 - \delta^2} - \delta/4}, \quad \delta = \delta_{2k}.$$

These RNSP parameters, (ρ, β) , are dictated as increasing functions of the RIC $\delta < 1$. A smaller δ requires an increased number of observations. All proofs invoke different classes of observing matrices which are randomly generated so that they satisfy a desirable observing properties—RIP, RNSP, or Constrained Minimal Singular Values (CMSV) property. Accordingly, the error statements are probabilistic in nature, referring to the ensemble of these observations.

3. ON THE SPARSITY OF THE UNCONSTRAINED LASSO MINIMIZER

We analyze the sparsity and ℓ_1/ℓ_2 -error bounds of the minimizer (1.2) in recovering $\mathbf{x}_*(k)$ from the observation $\mathbf{y}_*^\epsilon = A\mathbf{x}_* + \boldsymbol{\epsilon}$, with small measurement error, $\epsilon = |\boldsymbol{\epsilon}|_2$, and — assuming that $\mathbf{x}_*$ is k -compressible — with small compressibility error, $\sigma_k(\mathbf{x}_*) = |\mathbf{x}_* - \mathbf{x}_*(k)|_1$. Set

$$\mu := \sigma_k(\mathbf{x}_*) + \epsilon.$$

Clearly, since the exact solution is observed up to ℓ_2 residual error of order $|\mathbf{y}_*^\epsilon - A\mathbf{x}_*|_2 \leq \mu$, we do not have much to say when the computed residual error $|\mathbf{r}_\lambda|_2$ is of order μ and we will therefore limit ourselves to the parametric regime where $|\mathbf{r}_\lambda|_2 \gg \mu$. Below we show that $|\mathbf{r}_\lambda|_2 \sim \lambda \sqrt{k}$ and therefore throughout the paper we make the assumption

$$(3.1) \quad \theta := \frac{2\mu}{\lambda \sqrt{k}} \leq 1, \quad \mu = \sigma_k(\mathbf{x}_*) + \epsilon.$$

Thus, we assume the LASSO weight, λ , does not get exceedingly small, $\lambda \geq 2\mu/\sqrt{k}$. In concrete examples demonstrating the sparsity and error bounds reported below we use $\theta = 0.1$, corresponding to $\lambda \geq 20\mu/\sqrt{k}$.

Lemma 3.1 (The re-scaled residual — an upper-bound). Fix $\lambda < \lambda_\infty := |A^\top \mathbf{y}_*^\epsilon|_\infty$. Let $\mathbf{y}_*^\epsilon = A\mathbf{x}_* + \boldsymbol{\varepsilon}$ be the observation of a k -compressible unknown $\mathbf{x}_* \in \mathbb{R}^N$, observed by $A \in \mathbb{R}^{m \times N}$ satisfying the $\text{RNSP}_{\rho, \tau}$ of order k , (2.6). Let μ denote the small scale of k -compressibility and measurement errors, see (3.1). Then the residual of the LASSO (1.2), $\mathbf{r}_\lambda = \mathbf{y}_*^\epsilon - A\mathbf{x}_\lambda$, satisfies

$$(3.2) \quad \frac{|\mathbf{r}_\lambda|_2}{\lambda} \leq (\beta_\delta + \theta)\sqrt{k}, \quad \beta_\delta = \frac{\sqrt{1+\delta}}{\sqrt{1-\delta^2-\delta/4}}.$$

Proof. Clearly, $|A(\mathbf{x}_\lambda - \mathbf{x}_*(k))|_2 \leq |\mathbf{r}_\lambda|_2 + \mu$. Using (2.5) with the k -sparse $\mathbf{v} = \mathbf{x}_*(k)$ and $\mathbf{u} = \mathbf{x}_\lambda$ yields

$$(3.3) \quad |\mathbf{x}_*(k)|_1 - |\mathbf{x}_\lambda|_1 \leq \tau |A(\mathbf{x}_\lambda - \mathbf{x}_*(k))|_2 \leq \tau |\mathbf{r}_\lambda|_2 + \tau \mu.$$

Next, a lower-bound for the quantity on the left follows. Recall that $\mathbf{x}_\lambda$, being the LASSO minimizer (1.2), is characterized by the extremal property that its scaled residual $\mathbf{z} = \frac{\mathbf{r}_\lambda}{\lambda}$ satisfies (2.1),

$$|\mathbf{x}_\lambda|_1 = \langle A\mathbf{x}_\lambda, \mathbf{z} \rangle \text{ and } |A^\top \mathbf{z}|_\infty = 1, \quad \mathbf{z} := \frac{\mathbf{r}_\lambda}{\lambda}.$$

Hence

$$\begin{aligned} |\mathbf{x}_*(k)|_1 - |\mathbf{x}_\lambda|_1 &\geq \langle \mathbf{x}_*(k), A^\top \mathbf{z} \rangle - \langle A\mathbf{x}_\lambda, \mathbf{z} \rangle = \langle A\mathbf{x}_*(k) - A\mathbf{x}_\lambda, \mathbf{z} \rangle \\ &= \langle \mathbf{r}_\lambda, \mathbf{z} \rangle + \langle A\mathbf{x}_*(k) - \mathbf{y}_*^\epsilon, \mathbf{z} \rangle \\ &\geq \frac{|\mathbf{r}_\lambda|_2^2}{\lambda} - |A\mathbf{x}_*(k) - \mathbf{y}_*^\epsilon|_2 \frac{|\mathbf{r}_\lambda|_2}{\lambda}. \end{aligned}$$

Now, assumption (3.1) and the fact that $|A|_{1 \rightarrow 2} \leq 1$ imply,

$$|A\mathbf{x}_*(k) - \mathbf{y}_*^\epsilon|_2 \leq |A\mathbf{x}_* - \mathbf{y}_*^\epsilon|_2 + |A(\mathbf{x}_*(k) - \mathbf{x}_*)|_2 \leq \epsilon + \sigma_k(\mathbf{x}_*) = \mu,$$

and we end with the desired lower-bound

$$(3.4) \quad |\mathbf{x}_*(k)|_1 - |\mathbf{x}_\lambda|_1 \geq \frac{|\mathbf{r}_\lambda|_2^2}{\lambda} - \mu \frac{|\mathbf{r}_\lambda|_2}{\lambda}.$$

Combining (3.3) and (3.4) we conclude that $|\mathbf{z}|_2 = \frac{|\mathbf{r}_\lambda|_2}{\lambda}$ satisfies the quadratic inequality, $|\mathbf{z}|_2^2 \leq \left(\tau + \frac{\mu}{\lambda}\right)|\mathbf{z}|_2 + \tau \frac{\mu}{\lambda}$, and therefore

$$(3.5) \quad \frac{|\mathbf{r}_\lambda|_2}{\lambda} = |\mathbf{z}|_2 < \tau + \frac{2\mu}{\lambda} = (\beta_\delta + \theta)\sqrt{k},$$

proving (3.2). $\square$

3.1. Bounds of the sparsity. We now come to the main point of the lower-bound on the scaled residual in terms of the size of the support of $\mathbf{x}_\lambda$, $\frac{|\mathbf{r}_\lambda|_2}{\lambda} \gtrsim \sqrt{s_\lambda}$. Fix $\lambda < \lambda_\infty := |A^\top \mathbf{y}_*^\epsilon|_\infty$. Recall that if $\mathbf{x}_\lambda$ is the LASSO minimizer (1.2) then by the extremal property (2.1), the scaled residual $\mathbf{z} = \frac{\mathbf{r}_\lambda}{\lambda}$ satisfies the two properties $\langle A\mathbf{x}_\lambda, \mathbf{z} \rangle = |\mathbf{x}_\lambda|_1$ and $|A^\top \mathbf{z}|_\infty = 1$. Thus, the extremal $\mathbf{x}_\lambda$ with support $\mathcal{S} = \text{supp}(\mathbf{x}_\lambda)$ of size $s_\lambda = |\mathbf{x}_\lambda|_0$, is identified by a re-scaled residual satisfying

$$(3.6) \quad (A^\top \mathbf{z})_{\mathcal{S}} = \mathbf{sgn}(\mathbf{x}_{\lambda, \mathcal{S}}), \quad \mathbf{z} = \frac{\mathbf{r}_\lambda}{\lambda}, \quad \mathcal{S} = \text{supp}(\mathbf{x}_\lambda).$$

Fix the integer t ,

$$(3.7) \quad t := [(1 + \delta)(\beta_\delta + \theta)^2 k] + 1 \text{ with constant } \delta > \delta_t.$$

Since the RIC δ_t is increasing with the order t , there is no need to trace a precise fixed point associated with (3.7), $t = [(1 + \delta_t)(\beta_{\delta_t} + \theta)^2 k] + 1$. Instead, we can use a priori bounds of δ_t ; for example, if we restrict ourselves to the range $\delta < 0.7$, we can set the integer upper bound $t = 11k$. Below, we demonstrate refined versions of this bound.

We claim that

$$(3.8) \quad s_\lambda < t = [(1 + \delta)(\beta_\delta + \theta)^2 k] + 1.$$

To this end we proceed by contradiction. Assume $s_\lambda \geq t$. Then the support of $\mathbf{x}_\lambda$ has a subset $\mathcal{T}$ of size t for which the extremal property (3.6) reads $(A^\top \mathbf{z})_{\mathcal{T}} = \mathbf{sgn}(\mathbf{x}_{\lambda, \mathcal{T}})$, and the RIP (2.2) for such set $\mathcal{T}$ implies

$$(3.9) \quad |A(A^\top \mathbf{z})_{\mathcal{T}}|_2^2 \leq (1 + \delta_t) |(A^\top \mathbf{z})_{\mathcal{T}}|_2^2 = (1 + \delta_t) |\mathbf{sgn}(\mathbf{x}_{\lambda, \mathcal{T}})|_2^2 = (1 + \delta_t)t.$$

On the other hand, we have

$$|A(A^\top \mathbf{z})_{\mathcal{T}}|_2^2 \geq \frac{1}{|\mathbf{z}|_2^2} \langle A(A^\top \mathbf{z})_{\mathcal{T}}, \mathbf{z} \rangle^2 = \frac{1}{|\mathbf{z}|_2^2} \langle (A^\top \mathbf{z})_{\mathcal{T}}, A^\top \mathbf{z} \rangle^2 = \frac{1}{|\mathbf{z}|_2^2} |(A^\top \mathbf{z})_{\mathcal{T}}|_2^4 = \frac{t^2}{|\mathbf{z}|_2^2}.$$

The last two inequalities followed by Lemma 3.1 imply $t \leq (1 + \delta_t) |\mathbf{z}|_2^2 < (1 + \delta)(\beta_\delta + \theta)^2 k$, which contradicts the definition of t ,

$$t = [(1 + \delta)(\beta_\delta + \theta)^2 k] + 1 \geq (1 + \delta_t)(\beta_\delta + \theta)^2 k.$$

Thus, (3.8) holds.

In fact, a refined statement follows. Now that we know $|\mathcal{S}| \leq [(1 + \delta)(\beta_\delta + \theta)^2 k]$ we can argue along the same line as above with $\mathcal{T} = \mathcal{S}$, obtaining $s_\lambda \leq (1 + \delta) |\mathbf{z}|_2^2$.

Lemma 3.2 (The re-scaled residual — a lower-bound). Fix $\lambda < \lambda_\infty := |A^\top \mathbf{y}_*^\epsilon|_\infty$ and let $\mathbf{x}_\lambda$ be the s_λ -sparse minimizer of the corresponding LASSO (1.2), observed with RIC δ such that (3.7) holds. Then the residual, $\mathbf{r}_\lambda = \mathbf{y}_*^\epsilon - A\mathbf{x}_\lambda$, satisfies

$$(3.10) \quad \frac{|\mathbf{r}_\lambda|_2^2}{\lambda^2} \geq \frac{s_\lambda}{1 + \delta}.$$

Combining the lower- and upper-bounds of $\frac{|\mathbf{r}_\lambda|_2^2}{\lambda^2}$ we conclude the following.

Theorem 3.3 (Sparsity bound). Fix $\lambda < \lambda_\infty := |A^\top \mathbf{y}_*^\epsilon|_\infty$ and let $\mathbf{x}_\lambda$ be the s_λ -sparse minimizer of the corresponding LASSO (1.2), observed with RIC δ such that (3.7) holds. Then

$$(3.11) \quad \frac{s_\lambda}{1 + \delta} \leq \frac{|\mathbf{r}_\lambda|_2^2}{\lambda^2} \leq (\beta_\delta + \theta)^2 k, \quad \delta > \delta_t, \quad t = [(1 + \delta_t)(\beta_{\delta_t} + \theta)^2 k] + 1.$$

In particular, we recover (3.8), $s_\lambda \leq [(1 + \delta)(\beta_\delta + \theta)^2 k]$.

We demonstrate the application of corollary 3.3 for different choices of RICs. In all cases, we set $\theta = 0.1$. We begin with the RIC $\delta = 0.7$, obtaining $(\beta_\delta + \theta) = 2.52 \leadsto s_\lambda \leq (1 + \delta)(\beta_\delta + \theta)^2 k < 11k$. Thus, with $t = 11k$ we require $\delta_{11k} < 0.7$ which in turn, by (2.3), set the number of required observations

$$s_\lambda < 11k : \quad m \approx \text{Const.} \frac{11k}{0.7^2} \ln(eN/k) \approx \text{Const.} 22.4 k \ln(eN/k).$$

For a second example we choose a smaller RIC $\delta = 0.4$ and $\theta = 0.1$. Recall, that a smaller δ requires more observations yet in the number of observations in the present context depends on δ_t .

In this case $(\beta_\delta + \theta) = 1.55 \leadsto s_\lambda \leq (1 + \delta)(\beta_\delta + \theta)^2 k = 3.36k < 4k$. This requires a slightly larger number of observations (or at least a smaller bound (2.3))

$$s_\lambda < 4k : \quad m \approx \text{Const.} \frac{4k}{0.4^2} \ln(eN/k) \approx \text{Const.} 25 k \ln(eN/k).$$

Finally, as a third example we choose an even smaller the RIC $\delta = 0.26$ and the same $\theta = 0.1$. In this case $(\beta_\delta + \theta) = 1.35 \leadsto s_\lambda \leq (1 + \delta)(\beta_\delta + \theta)^2 k = 2.28k < 3k$, and this yields the number of required observations

$$s_\lambda < 3k : \quad m \approx \text{Const.} \frac{3k}{0.26^2} \ln(eN/k) \approx \text{Const.} 44.4 k \ln(eN/k).$$

Remark 3.4 (On the threshold parameter χ). Observe that the sparsity bound s_λ is uniform in the small scale μ throughout the parametric regime assumed in (3.1). Thus, in the range of $\lambda \gg 2\mu/\sqrt{k}$, the support of the computed solution $|\mathbf{x}_\lambda|_0$ can grow at most by a fixed factor relative to the k -support of underlying unknown $\mathbf{x}_*$, [38, Appendix A]. We write

$$(3.12) \quad s_\lambda < ([\chi^2] + 1)k, \quad \chi := \sqrt{1 + \delta}(\beta_\delta + \theta) = \frac{1 + \delta}{\sqrt{1 - \delta^2} - \delta/4} + \sqrt{1 + \delta}\theta.$$

We have the theoretical bounds $[\chi^2] + 1 = 11$ corresponding to $\delta = 0.7$ and $[\chi^2] + 1 = 4$ corresponding to $\delta \approx 0.4$.

3.2. Numerical simulations. We report here on our simulations of the unconstrained LASSO (1.2), applied to the recovery of k -sparse data, $\sigma_k = 0$, that is $\mu = \epsilon$, with $(k, m, N) = (160, 1024, 4096)$. We consider different levels of noise $\epsilon = 10^{-3}, 10^{-2}, 10^{-1}$, in the corresponding parametric regime (3.1), $\lambda > 2\mu/\sqrt{k} = 0.16\epsilon$. The results are obtained by averaging 100 observations using randomly generated $\text{RNSP}_{\rho, \tau}$ matrices based on Gaussian distributions. A simple proof of the RIP for such matrices can be found in [4]. The results are compared with the sparsity bound of theorem 3.3. We note that our sparsity bound depends in an essential manner on the the RICs, $1 \pm \delta$, in (2.2). The parametric regime in (2.3) provides only a rough estimate on the range of allowable RICs, and in particular, does not cover the parameters used in the simulations below, [26, §9.4]. A detailed study which traces the sharp RICs can be found in [2, 3], but is beyond the scope of our work. We compare the simulations with our sparsity bound based on the RIC $\delta = 0.7$. This is partly motivated by the result of [10] in which the authors prove an exact BP recovery of k -sparse data from the RIP with $\delta_{tk} < \sqrt{(t-1)/t}$. In our case, the computation reported in figure 3.1 indicates the *actual* sparsity $s_\lambda < tk$ with $t = 2$ which is consistent with $\delta < \sqrt{1/2} \approx 0.7$. Although the RIC $\delta = 0.7$ does not provide a tight bound, $s_\lambda < 11k$, it suffices to detect the correct *behavior* of the LASSO minimizer, reported in figures 3.1–3.3 and 4.1–4.2.

We record here the corresponding parameters involved in our bounds:

$$\beta_{\delta=0.7} = \frac{\sqrt{1 + \delta}}{\sqrt{1 - \delta^2} - \delta/4} = 2.42, \quad \chi_{\delta=0.7} = \sqrt{1 + \delta}(\beta_\delta + \theta) = 3.16, \quad \eta_{\delta=0.7} = \frac{1}{1 + \delta} = 0.59$$

Our main result on the sparsity of the LASSO minimizer in theorem 3.3 provides a reasonably accurate information about the behavior of the unconstrained LASSO minimization (1.2). Figure 3.1 shows the behavior of the support, $s_\lambda = |\mathbf{x}_\lambda|_0$, starting with $s_\lambda = 0$ for $\lambda > \lambda_\infty$ and monotonically increasing as λ decreases all the way to a critical value, $\lambda_c \sim 0.11$, at which point s_{λ_c} reaches its maximal value of 215. This should be compared with our bound $s_\lambda \leq (1 + \delta)(\beta_\delta + \theta)^2 k$. For $\delta = 0.7$ we have $s_\lambda \leq 11k$, which is a rough sparsity bound, relative to the actual $s_\lambda \sim 215$. A smaller RIC $\delta \sim 0.2$ yields a tighter sparsity bound $1.66k \sim 313$.

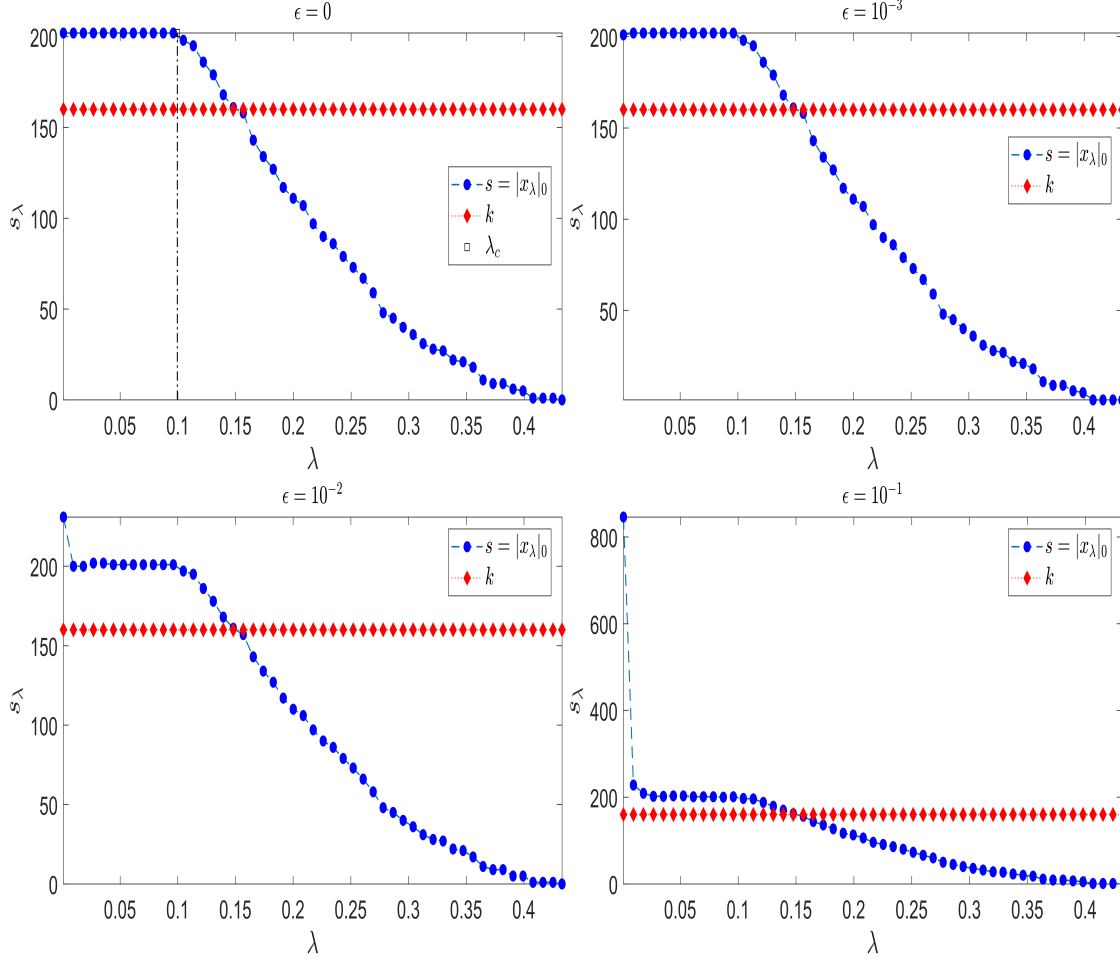


FIGURE 3.1. The support for computed minimizer $s_\lambda = |x_\lambda|_0$ of k -sparse data, $k = 160$, peaks at the threshold value of $k_{\max} \sim 215$ when λ reaches $\lambda_c \sim 0.11$. This should be compared with the rough upper bound $(1 + \delta)(\beta_\delta + \theta)^2 k \leq 11k$ corresponding to the RIC $\delta = 0.7$, and the more realistic bound $4k$ corresponding to $\delta = 0.4$. Observe (lower figures) that for exceedingly small $\lambda \ll \epsilon$, there is an additional growth of order $\frac{\epsilon}{\lambda}$.

Observe that according to (3.4), the ℓ_1 -size of the LASSO minimizer x_λ remains smaller than the target $|x_*(k)|_1$. Indeed, as long as the residual $|r_\lambda|_2 > \mu$, then

$$(3.13) \quad |x_*(k)|_1 - |x_\lambda|_1 \geq (|r_\lambda|_2 - \mu) \frac{|r_\lambda|_2}{\lambda}.$$

This is depicted in figure 3.2: as λ decreases, the ratio $\frac{|r_\lambda|_2}{\lambda} \gtrsim \sqrt{s_\lambda}$ is increasing until $|x_\lambda|_1$ reaches its upper bound of $|x_*(k)|_1$.

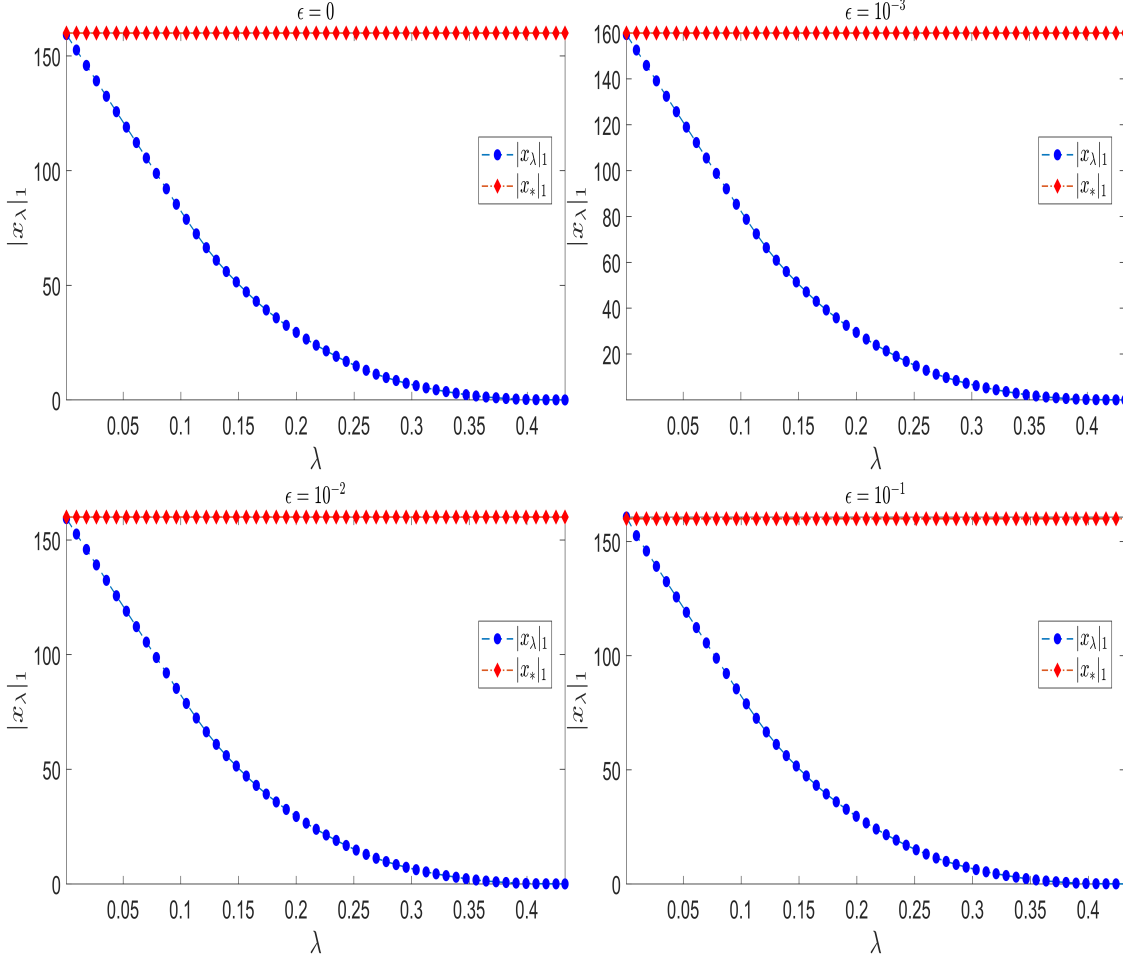


FIGURE 3.2. ℓ_1 norm of x_λ approaches its upper-bound $|x_*(k)|_1$ as λ decreases.

Figure 3.3 shows the aptitude of the lower- and upper-bounds of the re-scaled residual (3.11), in capturing the re-scaled residual $\frac{|r_\lambda|_2}{\lambda}$. Again, the three quantities increase with decreasing λ , until λ reaches the threshold λ_c at which point the re-scaled residual, $\frac{|r_\lambda|_2}{\lambda}$, peaks at its maximal value ~ 27 , in agreement with the upper-bound (3.2), $\frac{|r_\lambda|_2}{\lambda} < \beta_\delta \sqrt{k} + \frac{2\epsilon}{\lambda} < 30.61 + \frac{2\epsilon}{\lambda}$.

4. ERROR BOUNDS

4.1. ℓ_2 -error bounds. The sparsity bound (3.11) was derived based on a two-sided ℓ_2 -bound of the scaled residual. The latter can be converted into a two-sided ℓ_2 error bound of $|x_\lambda - x_*(k)|_2$. Note that since $||r_\lambda|_2 - |A(x_*(k) - x_\lambda)|_2| \leq \mu$, then the upper-bound on $|r_\lambda|_2$, see (3.5), also bounds the ‘observed error’ $A(x_\lambda - x_*(k))$,

$$(4.1) \quad |A(x_\lambda - x_*(k))|_2 \leq (\beta_\delta + \theta)\sqrt{k}\lambda + \mu \leq \beta_\delta\sqrt{k}\lambda + 3\mu.$$

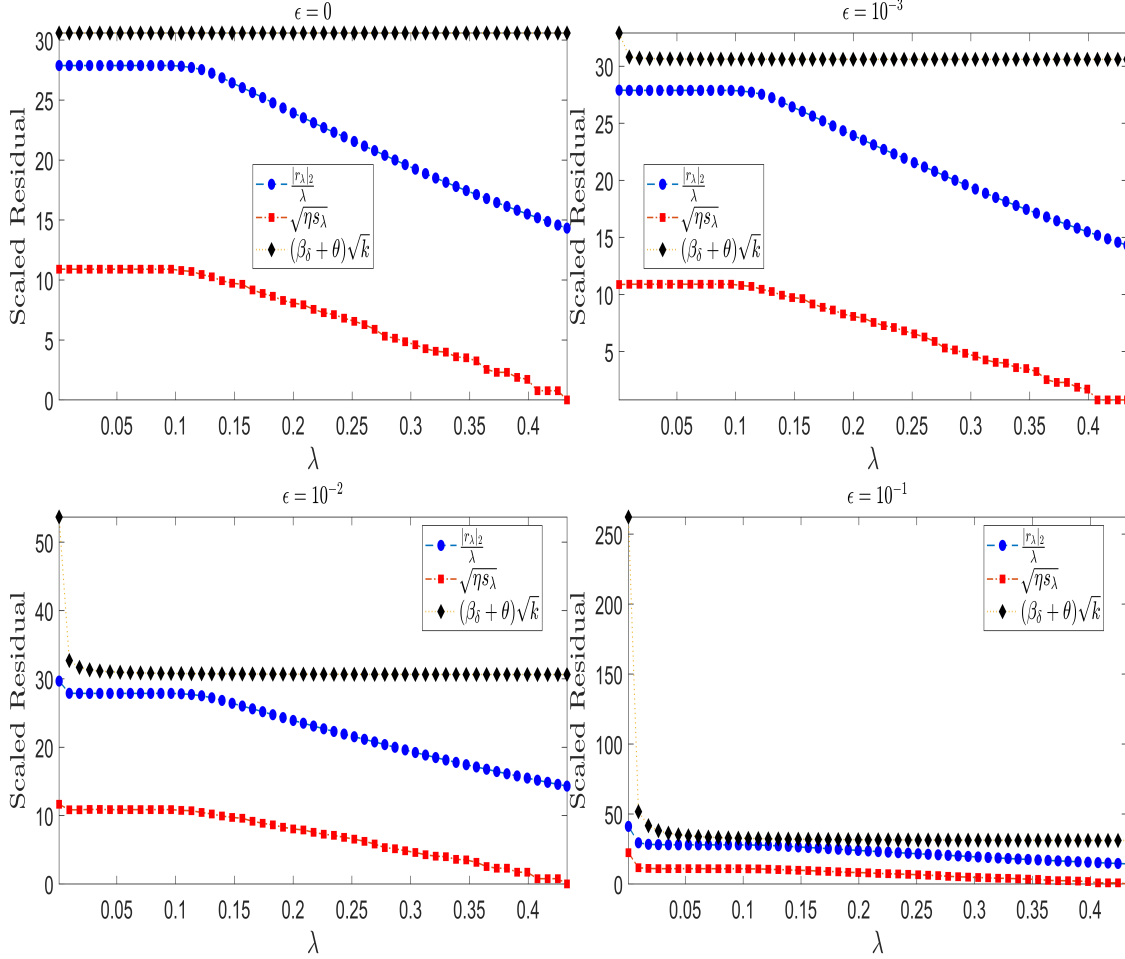


FIGURE 3.3. Re-scaled residual $\frac{|r_\lambda|_2}{\lambda}$ captured between its lower- and upper-bounds (3.10) and (3.2), $\sqrt{\eta s_\lambda} \leq \frac{|r_\lambda|_2}{\lambda} \leq \beta_\delta \sqrt{k} + \frac{2\epsilon}{\lambda} \approx 30.61$ with $(\eta, \beta, \theta) = (0.59, 2.42, 0.1)$ corresponding to $\delta = 0.7$. It peaks at a threshold value of 27, independent of the level of noise. When $\lambda \ll \epsilon$, there is an additional large term of order $\frac{2\epsilon}{\lambda}$.

The sparsity of $\mathbf{x}_\lambda$ does not exceed $s_\lambda \leq ([\chi^2] + 1)k$ hence $\mathbf{x}_\lambda - \mathbf{x}_*(k)$ has sparsity $([\chi^2] + 2)k$, and the RIP (2.2) implies the ℓ_2 -error upper-bound

$$(4.2) \quad |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_2 \leq \frac{1}{\sqrt{1-\delta}} (\beta_\delta \sqrt{k} \lambda + 3\mu), \quad \delta = \delta_{([\chi^2]+2)k}.$$

In particular, (4.2) with $\frac{1}{\sqrt{1-\delta}} \leq 1.83$ and $\beta_\delta \leq 2.42$ corresponding to $\delta = 0.7$ yields

$$(4.3) \quad |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_2 \lesssim 4.43 \sqrt{k} \lambda + 5.48 \mu.$$

This recovers a quantitative version of the ℓ_2 upper bound proved under additional condition of an incoherence design assumption in [30, Theorem 1], an ℓ_1 -CMSV assumption⁵ [38], or restricted eigenvalue bound in [28, Theorem 11.1].

⁵In fact, we slightly improve the quadratic dependence of the bound in [38, (23)] on the ℓ_1 -CMSV constant $\sim \rho_{4k}^{-2}$, mentioned in (4.6) below.

The upper-bound (4.2) is sharp in the sense of having a tight ℓ_2 -lower bound: since the error $\mathbf{x}_\lambda - \mathbf{x}_*(k)$ is at most $(\lceil \chi^2 \rceil + 2)k$ -sparse, we can use the RIP to translate the lower bound (3.10) into an ℓ_2 lower-bound,

$$|\mathbf{x}_\lambda - \mathbf{x}_*(k)|_2 \geq \frac{1}{\sqrt{1+\delta}} |A(\mathbf{x}_\lambda - \mathbf{x}_*(k))|_2 \geq \frac{1}{\sqrt{1+\delta}} (|\mathbf{r}_\lambda|_2 - \mu) \geq \frac{\sqrt{s_\lambda} \lambda}{1+\delta} - \frac{\mu}{\sqrt{1+\delta}}.$$

We summarize these bounds in the following form.

Theorem 4.1 (ℓ_2 -bound). *Fix $\lambda < \lambda_\infty := |A^\top \mathbf{y}_*^\epsilon|_\infty$ and let $\mathbf{x}_\lambda$ be the s_λ -sparse minimizer of the corresponding LASSO (1.2), observed with RIP matrix A . Then*

$$(4.4) \quad \frac{1}{\sqrt{1+\delta}} \left(\frac{\sqrt{s_\lambda} \lambda}{\sqrt{1+\delta}} - \mu \right) \leq |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_2 \leq \frac{1}{\sqrt{1-\delta}} (\beta_\delta \sqrt{k} \lambda + 3\mu), \quad \delta = \delta_{(\lceil \chi^2 \rceil + 2)k}.$$

Remark 4.2 (Compared with the ℓ_1 -entropy bound). *The extremal relation $\langle A_S \mathbf{x}_{\lambda,S}, \mathbf{r}_\lambda \rangle = \lambda |\mathbf{x}_{\lambda,S}|_1$ and the RIP (2.2) yield $\lambda |\mathbf{x}_{\lambda,S}|_1 \leq \sqrt{1+\delta} |\mathbf{x}_{\lambda,S}|_2 |\mathbf{r}_\lambda|_2$, and hence we end up with a lower-bound involving the ℓ_1 -entropy of $\{\mathbf{x}_{\lambda,S}\}$,*

$$(4.5) \quad \frac{|\mathbf{r}_\lambda|_2^2}{\lambda^2} \geq \frac{1}{1+\delta} \text{Ent}(\mathbf{x}_{\lambda,S}) \quad \text{Ent}(\mathbf{x}) := \frac{|\mathbf{x}|_1^2}{|\mathbf{x}|_2^2}.$$

This bound is tied to a Null Entropy Property of A [1, §3.2] or the ℓ_1 -CMSV constant $\rho_s(A)$ introduced in [38]⁶

$$(4.6) \quad \frac{|\mathbf{r}_\lambda|_2^2}{\lambda^2} \geq \frac{\text{Ent}(\mathbf{x}_{\lambda,S})}{\rho_s(A)}, \quad \rho_s(A) := \min_{|\mathbf{x}|_2=1} \left\{ |A\mathbf{x}|_2 : \text{Ent}(\mathbf{x}) \leq s \right\}.$$

⁶Which is not to be confused with the RNSP parameter in (2.6)

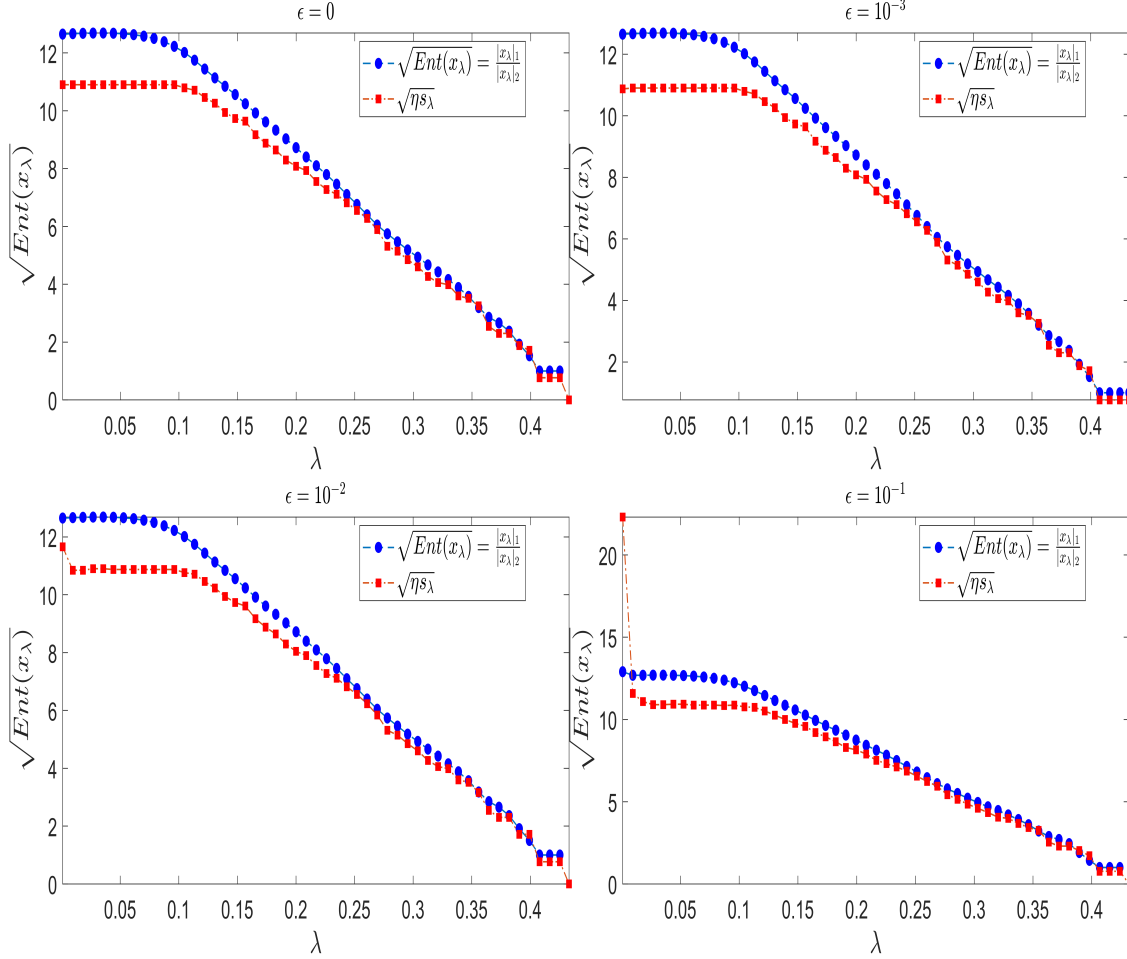


FIGURE 4.1. Lower bounds of the re-scaled residual: (3.10) with $\eta := \frac{1}{1+\delta} = 0.592$ vs. the ℓ_1 -entropy based (4.5).

Clearly, if $\mathbf{x}_\lambda$ has the sparsity s_λ then $\text{Ent}(\mathbf{x}_{\lambda,S}) \leq s_\lambda$. Here we note about the reverse implication, namely — if the reverse inequality holds, $\text{Ent}(\mathbf{x}_{\lambda,S}) \gtrsim s_\lambda$, then it would yield our sparsity result of lemma 3.2, based on the lower bound $\frac{|\mathbf{r}_\lambda|_2}{\lambda} \gtrsim \frac{\sqrt{s_\lambda}}{\rho_{s_\lambda}}$. Theorem 3.3 suggests the lower-entropy bound for the minimizers $\mathbf{x}_\lambda$. Indeed, figure 4.1 shows a remarkable agreement between the lower bound (3.10) with $\delta = 0.7$ and the ℓ_1 -entropy bound (4.5), $\text{Ent}(\mathbf{x}_\lambda)$, at least before the support of $\mathbf{x}_\lambda$ reaches its peak at $k_{\max}$.

4.2. ℓ_1 -error bound. We recall the ℓ_2 -bound (4.2) which we express in the form $|\mathbf{x}_\lambda - \mathbf{x}_*(k)|_2 \leq \frac{1}{\sqrt{1-\delta}}(\beta_\delta + 3/2\theta)\sqrt{k}\lambda$. Since $\mathbf{x}_\lambda - \mathbf{x}_*(k)$ has sparsity of order $\leq k + \chi^2 k$, we derive the following ℓ_1 -bound.

Theorem 4.3 (ℓ_1 -error bound). Fix $\lambda < \lambda_\infty := |A^\top \mathbf{y}_*^\epsilon|_\infty$ and let $\mathbf{x}_\lambda$ be the LASSO minimizer of (1.2), observed with RIP matrix A with RIC δ such that (3.7) holds. Then the following ℓ_1 -error

bound holds,

$$\begin{aligned}
 |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_1 &\leq \sqrt{(1 + (1 + \delta)(\beta_\delta + \theta)^2)k} |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_2 \\
 (4.7) \quad &< \sqrt{1 + \delta}(\beta_\delta + 1/2 + \theta) \sqrt{k} \frac{1}{\sqrt{1 - \delta}} (\beta_\delta + 3/2\theta) \sqrt{k} \lambda \\
 &< \frac{\sqrt{1 + \delta}}{\sqrt{1 - \delta}} \frac{1}{\lambda} \left((\beta_\delta + 1/2) \sqrt{k} \lambda + 2\mu \right)^2.
 \end{aligned}$$

The amplitude of $k\lambda$ in the ℓ_1 -error bound (4.7) is not sharp. For example, with RIC $\delta = \delta_{11k} < 0.7$ we have $\beta_\delta > 2$ in which case, omitting the negligibly small μ^2/λ terms, one ends up with the improved bound

$$(4.8) \quad |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_1 < \frac{\sqrt{1 + \delta}}{\sqrt{1 - \delta}} \left((\beta_\delta + 1/4)^2 k \lambda + (4\beta_\delta + 1) \sqrt{k} \mu \right) < 16.97 k \lambda + 24.04 \sqrt{k} \mu.$$

We conclude with an alternative derivation of an ℓ_1 -error bound. To this end, we recall the RNSP bound [26, Theorem 4.20], which states that for all $\mathcal{K} \subset \{1, 2, \dots, N\}$ of size $\leq k$ and for any $\mathbf{u}, \mathbf{v} \in \mathbb{R}^N$, the following holds,

$$|\mathbf{u} - \mathbf{v}|_1 \leq \frac{1 + \rho}{1 - \rho} (|\mathbf{u}|_1 - |\mathbf{v}|_1 + 2|\mathbf{v}_{\mathcal{K}^c}|_1) + \frac{2\tau}{1 - \rho} |A(\mathbf{u} - \mathbf{v})|_2, \quad |\mathcal{K}| \leq k.$$

Using it with $(\mathbf{u}, \mathbf{v}) = (\mathbf{x}_\lambda, \mathbf{x}_*(k))$ and $\mathcal{K} = \text{supp}(\mathbf{x}_*(k))$ yields

$$(4.9) \quad |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_1 \leq \frac{1 + \rho}{1 - \rho} \left(|\mathbf{x}_\lambda|_1 - |\mathbf{x}_*(k)|_1 \right) + \frac{2\tau}{1 - \rho} |A(\mathbf{x}_\lambda - \mathbf{x}_*(k))|_2.$$

Now, using (3.4) to bound the term inside the first parenthesis on the right, and as before, noting that the second term does not exceed $|A(\mathbf{x}_\lambda - \mathbf{x}_*(k))|_2 \leq |\mathbf{r}_\lambda|_2 + \mu$, we find

$$|\mathbf{x}_\lambda - \mathbf{x}_*(k)|_1 \leq \frac{1 + \rho}{1 - \rho} \left\{ |\mathbf{r}_\lambda|_2 \left(\frac{2\tau}{1 + \rho} + \frac{\mu}{\lambda} - \frac{|\mathbf{r}_\lambda|_2}{\lambda} \right) + \frac{2\tau}{1 + \rho} \mu \right\}.$$

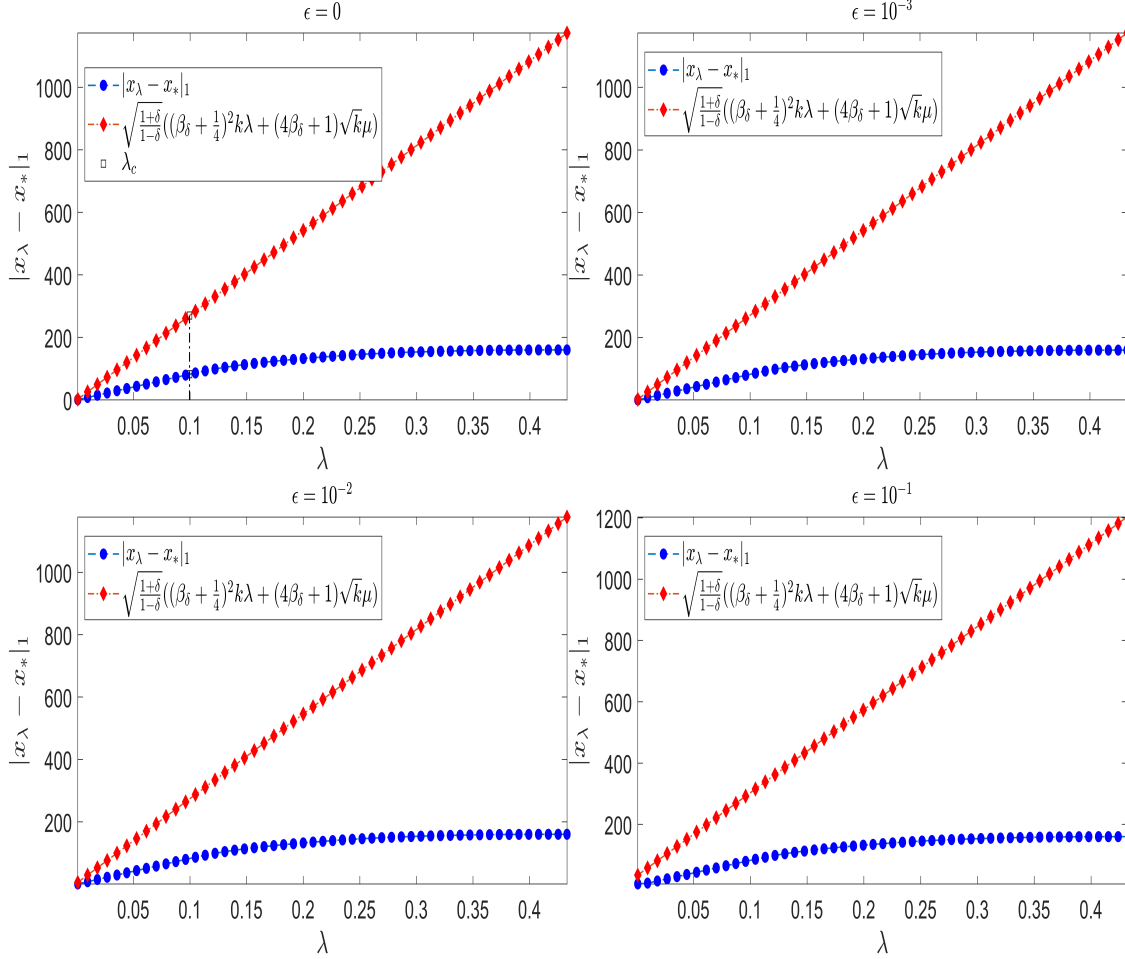
Given the RNSP parameters (2.6), $2\tau = 2\beta\sqrt{k}$ and $\frac{\mu}{\lambda} < \frac{\theta}{1 + \rho}\sqrt{k}$, the last bound yields

$$(4.10) \quad |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_1 \leq \frac{1}{1 - \rho} \left\{ |\mathbf{r}_\lambda|_2 \left(2(\beta_\delta + \theta)\sqrt{k} - (1 + \rho)\frac{|\mathbf{r}_\lambda|_2}{\lambda} \right) + \beta_\delta \theta k \lambda \right\}.$$

Viewed as quadratic in $\frac{|\mathbf{r}_\lambda|_2}{\lambda}$, the first expression on the right admits a maximal value $\frac{(\beta_\delta + \theta)^2}{1 + \rho} k \lambda$, and we finally end up with

$$(4.11) \quad |\mathbf{x}_\lambda - \mathbf{x}_*(k)|_1 \leq \frac{1}{1 - \rho^2} \left((\beta_\delta + \theta)^2 k \lambda + (1 + \rho) \beta_\delta \theta k \lambda \right) \leq \frac{(\beta_\delta + 2\theta)^2}{1 - \rho^2} k \lambda.$$

This recovers the ℓ_1 -bound of order $\mathcal{O}(k\lambda)$ as in (4.7). However, since the ℓ_1 bound (4.11) involves the value of $1/(1 - \rho)^2$, it is therefore limited to the RIC $\delta < 4/\sqrt{41}$ where ρ approaches 1.

FIGURE 4.2. ℓ_1 -error for recovery of sparse data compared with the upper-bound (4.12).

4.3. Numerical simulations. We report on the error behavior in our simulations of the unconstrained LASSO (1.2), applied to the recovery of k -sparse data, $\sigma_k = 0$, that is $\mu = \epsilon$, with $(k, m, N) = (160, 1024, 4096)$. The results are obtained by averaging 100 observations using randomly generated $\text{RNSP}_{\rho, \tau}$ matrices based on Gaussian distributions. We compare the ℓ_1 -error with the error bound (4.8)

$$(4.12) \quad |x_\lambda - x_*(k)|_1 \leq 16.97 * 160\lambda + 24.04\sqrt{160} \epsilon.$$

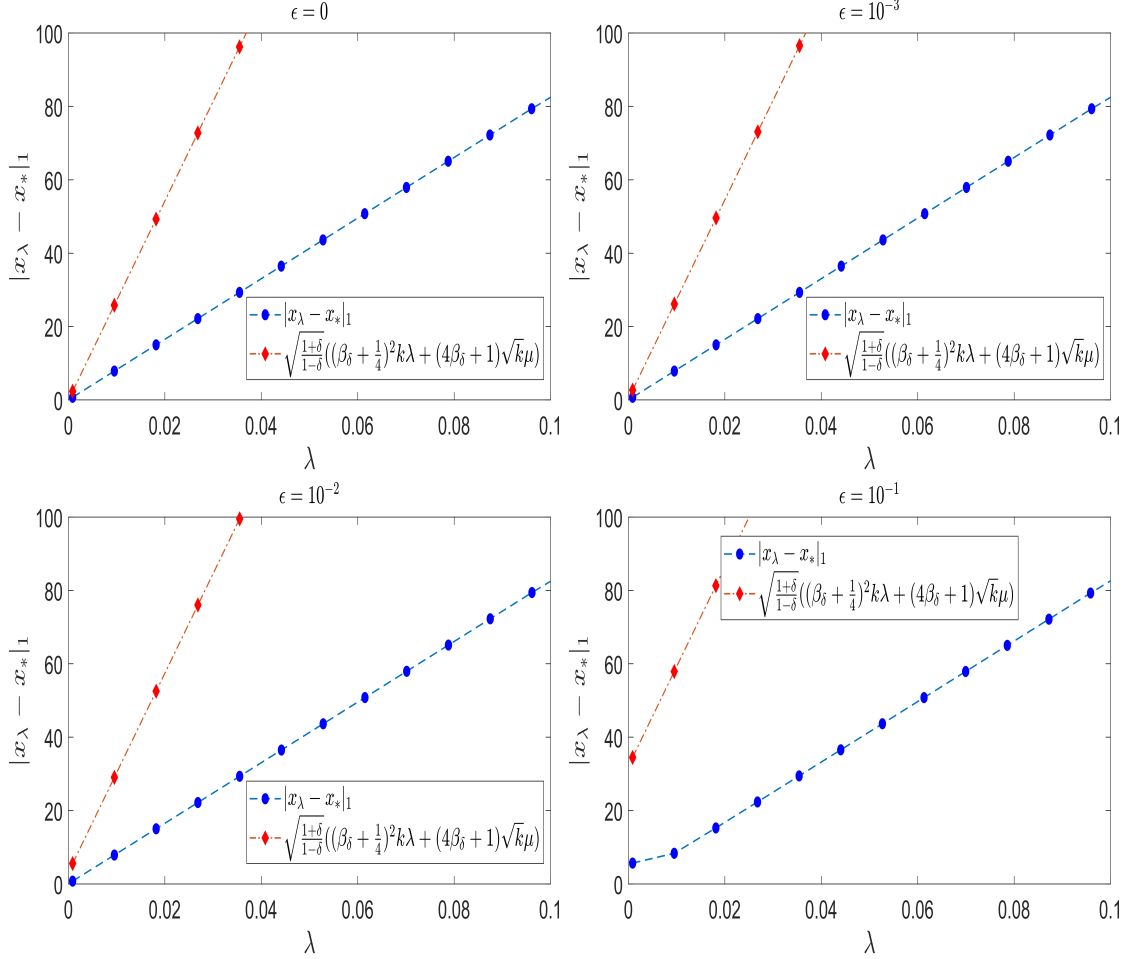


FIGURE 4.3. The ℓ_1 -error compared with the upper-bound (4.12) zoomed near $\lambda = 0$.

Acknowledgment. We are indebted to Wolfgang Dahmen for comments which greatly improved the presentation of material of this paper.

REFERENCES

- [1] J. ANDERSSON AND J.-O. STRÖMBERG, *On the theorem of uniform recovery of random sampling matrices*, IEEE Transactions on Information Theory, 60 (2014), pp. 1700–1710.
- [2] B. BAH AND J. TANNER, *Improved bounds on restricted isometry constants for Gaussian matrices*, SIAM Journal on Matrix Analysis and Applications, 31 (2010), pp. 2882–2898.
- [3] —, *Bounds of restricted isometry constants in extreme asymptotics: formulae for Gaussian matrices*, Linear Algebra and its Applications, 441 (2014), pp. 88–109.
- [4] R. BARANIUK, M. DAVENPORT, R. DEVORE, AND M. WAKIN, *A simple proof of the restricted isometry property for random matrices*, Constructive Approximation, 28 (2008), pp. 253–263.
- [5] A. BELLONI AND V. CHERNOZHUKOV, *Least squares after model selection in high-dimensional sparse models*, Bernoulli, 19 (2013), pp. 521–547.
- [6] P. J. BICKEL, Y. RITOV, AND A. B. TSYBAKOV, *Simultaneous analysis of Lasso and Dantzig selector*, The Annals of Statistics, 37 (2009), pp. 1705–1732.
- [7] J. D. BLANCHARD, C. CARTIS, AND J. TANNER, *Compressed sensing: How sharp is the restricted isometry property?*, SIAM Review, 53 (2011), pp. 105–125.

- [8] J. BOURGAIN, S. DILWORTH, K. FORD, S. KONYAGIN, AND D. KUTZAROVA, *Explicit constructions of RIP matrices and related problems*, Duke Mathematical Journal, 159 (2011), pp. 145–185.
- [9] P. BÜHLMANN AND S. VAN DE GEER, *Statistics for High-Dimensional Data: Methods, Theory and Applications*, Springer Science & Business Media, 2011.
- [10] T. T. CAI AND A. ZHANG, *Sparse representation of a polytope and recovery of sparse signals and low-rank matrices*, IEEE transactions on information theory, 60 (2013), pp. 122–132.
- [11] E. CANDÈS AND T. TAO, *The Dantzig selector: Statistical estimation when p is much larger than n* , The Annals of Statistics, 35 (2007), pp. 2313–2351.
- [12] E. J. CANDÈS, *The restricted isometry property and its implications for compressed sensing*, Comptes Rendus Mathematique, 346 (2008), pp. 589–592.
- [13] E. J. CANDÈS, J. K. ROMBERG, AND T. TAO, *Robust uncertainty principles: exact signal reconstruction from high incomplete frequency information*, IEEE Trans. Inform. Theor., 52 (2006), pp. 489 – 509.
- [14] ———, *Stable signal recovery from incomplete and inaccurate measurements*, Commun. Pure Appl. Math., 59 (2006), pp. 1207 – 1233.
- [15] E. J. CANDÈS AND T. TAO, *The Dantzig selector: statistical estimation when p is much larger than n* , Ann. Stat., 35 (2007), pp. 2313 – 2351.
- [16] S. S. CHEN, *Basis pursuit*, PhD thesis, Stanford University, 1996.
- [17] S. S. CHEN, D. L. DONOHO, AND M. A. SAUNDERS, *Atomic decomposition by Basis Pursuit*, SIAM J. Sci. Comput., 20 (1999), pp. 33 – 61.
- [18] S. S. CHEN AND D. L. DONOHO, *Basis pursuit*, in Proceedings of 1994 28th Asilomar Conference on Signals, Systems and Computers, vol. 1, IEEE, 1994, pp. 41–44.
- [19] A. COHEN, W. DAHMEN, AND R. A. DEVORE, *Compressed sensing and the best k -term approximation*, J. Amer. Math. Soc., 22 (2009), pp. 211 – 231.
- [20] R. DEVORE, G. PETROVA, AND P. WOJTASZCZYK, *Instance-optimality in probability with an ℓ_1 -minimization decoder*, Appl. Comput. Harmon. Anal., 27 (2009), pp. 275–288.
- [21] D. L. DONOHO, *Compressed sensing*, IEEE Trans. Inform. Theor., 52 (2006), pp. 1289 – 1306.
- [22] D. L. DONOHO AND M. ELAD, *Optimally sparse representations in general (non-orthogonal) dictionaries via ℓ^1 minimization*, Proc. Nat. Acad. Sci., 100 (2003), pp. 2197 – 2202.
- [23] D. L. DONOHO AND X. HUO, *Uncertainty principles and ideal atomic decompositions*, IEEE Trans. Inform. Theor., 47 (2001), pp. 2845 – 2862.
- [24] D. L. DONOHO AND I. M. JOHNSTONE, *Empirical atomic decomposition*, Unpublished manuscript, (1995).
- [25] S. FOUCART, *Sparse recovery algorithms: sufficient conditions in terms of restricted isometry constants*, in Approximation Theory XIII: San Antonio 2010, Springer, 2012, pp. 65–77.
- [26] S. FOUCART AND H. RAUHUT, *A Mathematical Introduction to Compressive Sensing*, Applied and Numerical Harmonic Analysis, Birkhäuser, 2013.
- [27] R. GRIBONVAL AND M. NIELSEN, *Sparse representations in unions of bases*, IEEE Trans. Inform. Theor., 49 (2003), pp. 3320 – 3325.
- [28] T. HASTIE, R. TIBSHIRANI, AND M. WAINWRIGHT, *Statistical Learning with Sparsity: the LASSO and generalizations*, CRC Press, 2015.
- [29] S. MALLAT AND W. L. HWANG, *Singularity detection and processing with wavelets*, IEEE Transactions on Information Theory, 38 (1992), pp. 617–643.
- [30] N. MEINSHAUSEN, B. YU, ET AL., *Lasso-type recovery of sparse representations for high-dimensional data*, The Annals of Statistics, 37 (2009), pp. 246–270.
- [31] L. I. RUDIN, S. OSHER, AND E. FATEMI, *Nonlinear total variation based noise removal algorithms*, Physica D: nonlinear phenomena, 60 (1992), pp. 259–268.
- [32] W. SU, M. BOGDAN, AND E. CANDÈS, *False discoveries occur early on the Lasso path*, The Annals of Statistics, (2017), pp. 2133–2150.
- [33] E. TADMOR, *Hierarchical construction of bounded solutions in regularity spaces*, Communications in Pure & Applied Mathematics, 69(6) (2015), pp. 1087–1109.
- [34] E. TADMOR, S. NEZZAR, AND L. VESSE, *A multiscale image representation using hierarchical (BV, L^2) decompositions*, Multiscale Model. Simul., 2 (2004), pp. 554 – 579.
- [35] ———, *Multiscale hierarchical decomposition of images with applications to deblurring, denoising and segmentation*, Commun. Math. Sci., 6 (2008), pp. 281 – 307.

- [36] E. TADMOR AND C. TAN, *Hierarchical construction of bounded solutions of $\operatorname{div} U = F$ in critical regularity spaces*, in Nonlinear Partial Differential Equations, H. Holden and K. Karlsen, eds., Abel Symposia 7, Oslo, Sep. 2010, pp. 255 – 269. 2010 Abel Symposium.
- [37] E. TADMOR AND M. ZHONG, *The multi-scale hierarchical reconstruction of ill-posed inverse problems*, In preparation, (2021).
- [38] G. TANG AND A. NEHORAI, *Performance analysis of sparse recovery based on constrained minimal singular values*, IEEE Transactions on Signal Processing, 59 (2011), pp. 5734–5745.
- [39] R. TIBSHIRANI, *Regression shrinkage and selection via the lasso*, J. Roy. Stat. Soc. B, 58 (1996), pp. 267 – 288.
- [40] J. A. TROPP, *Greed is good: Algorithmic results for sparse approximation*, IEEE Trans. Inform. Theor., 50 (2004), pp. 2231 – 2242.
- [41] H. ZHANG, W. YIN, AND L. CHENG, *Necessary and sufficient conditions on solution uniqueness in ℓ_1 minimization*, Journal of Optimization Theory and Applications, 164 (2015), pp. 109 – 122.

DEPARTMENT OF MATHEMATICS AND TEXAS A&M INSTITUTE OF DATA SCIENCE
TEXAS A&M UNIVERSITY, COLLEGE STATION, TX 77843
Email address: focart@tamu.edu

DEPARTMENT OF MATHEMATICS AND INSTITUTE FOR PHYSICAL SCIENCES & TECHNOLOGY (IPST)
UNIVERSITY OF MARYLAND, COLLEGE PARK, MD 20742
Email address: tadmor@umd.edu

TEXAS A&M INSTITUTE OF DATA SCIENCE
TEXAS A&M UNIVERSITY, COLLEGE STATION, TX 77843
Email address: mingzhong@tamu.edu