

Analysis of The Ratio of ℓ_1 and ℓ_2 Norms in Compressed Sensing

Yiming Xu¹, Akil Narayan^{1,2}, Hoang Tran³ and Clayton G. Webster^{4,5}

¹Department of Mathematics, University of Utah, Salt Lake City, 84112, yxu@math.utah.edu

²Scientific Computing and Imaging Institute, University of Utah, Salt Lake City, 84112, akil@sci.utah.edu

³Computer Science and Mathematics Division, Oak Ridge National Laboratory, Oak Ridge, TN, 37831, tranha@ornl.gov

⁴Oden Institute for Computational Engineering & Sciences, The University of Texas at Austin, Austin, TX, 78712, claytonwebster@utexas.edu

⁵Behavioral Reinforcement Learning Lab, Lirio LLC., Knoxville, TN, 37923

Abstract

We study the ratio of ℓ_1 and ℓ_2 norms (ℓ_1/ℓ_2) as a sparsity-promoting objective in compressed sensing. We first propose a novel criterion that guarantees that an s -sparse signal is the local minimizer of the ℓ_1/ℓ_2 objective; our criterion is interpretable and useful in practice. We also give the first uniform recovery condition using a geometric characterization of the null space of the measurement matrix, and show that this condition is easily satisfied for a class of random matrices. We also present analysis on the robustness of the procedure when noise pollutes data. Numerical experiments are provided that compare ℓ_1/ℓ_2 with some other popular non-convex methods in compressed sensing. Finally, we propose a novel initialization approach to accelerate the numerical optimization procedure. We call this initialization approach *support selection*, and we demonstrate that it empirically improves the performance of existing ℓ_1/ℓ_2 algorithms.

Keywords: Compressed sensing, High-dimensional geometry, Random matrices, Non-convex optimization

1 Introduction

The goal of compressed sensing (CS) problem is to seek the sparsest solution of an underdetermined linear system:

$$\min \|\mathbf{x}\|_0 \quad \text{subject to} \quad \mathbf{Ax} = \mathbf{b}, \quad (1.1)$$

where $\mathbf{x} \in \mathbb{R}^n$, $\mathbf{b} \in \mathbb{R}^m$ and $\mathbf{A} \in \mathbb{R}^{m \times n}$ with $m \ll n$. The quasinorm $\|\mathbf{x}\|_0$ measures the number of nonzero components in $\mathbf{x}$. In CS applications, one typically considers $\mathbf{x}$ as the frame/basis coordinates of an unknown signal, and it is typically assumed that the coordinate representation is *sparse*, i.e., that $\|\mathbf{x}\|_0$ is “small”. $\mathbf{A}$ is the measurement matrix that encodes linear measurements of the signal $\mathbf{x}$, and $\mathbf{b}$ contains the corresponding measured values. In the language of signal processing, (1.1) is equivalent to applying the sparsity decoder to reconstruct a signal from the undersampled measurement pair $(\mathbf{A}, \mathbf{b})$. A naive, empirical counting argument suggests that if $m \ll n$ measurements $\mathbf{b}$ of an unknown signal $\mathbf{x}$ are available, then we can perhaps compute the original signal coordinates $\mathbf{x}$, assuming $\mathbf{x}$ is approximately m -sparse. The optimization (1.1) is the quantitative manifestation of this argument.

It was established in [11] that under mild conditions, (1.1) has a unique minimizer. In the rest of the paper we assume that the minimizer is unique and denote it by $\mathbf{x}_0$. One of the central problems in compressed sensing is to design an effective algorithm to find $\mathbf{x}_0$: Directly solving

(1.1) via combinatorial search is NP-hard [26]. A more practical approach, which was proposed in the seminal work [12], is to relax the sparsity measure $\|\cdot\|_0$ to the convex ℓ_1 norm $\|\cdot\|_1$:

$$\min \|\mathbf{x}\|_1 \quad \text{subject to } \mathbf{Ax} = \mathbf{b}. \quad (1.2)$$

The convexity of the problem (1.2) ensures that efficient algorithms can be leveraged to compute solutions. Many pioneering works in compressed sensing have focused on understanding the equivalence between (1.1) and (1.2), see [12, 6, 13]. A major theoretical cornerstone of such equivalence results is the *Null Space Property* (NSP), which was first introduced in [9] and plays a crucial role in establishing sufficient and necessary conditions for the equivalence between (1.1) and (1.2). A sufficient condition for such an equivalence is called an *exact recovery condition*. A closely related but stronger condition is the *Restricted Isometry Property* (RIP), see [6]. The RIP is more flexible than the NSP for practical usage, yet conditions given by both the NSP and RIP are hard to verify in the case when measurements (i.e., the matrix $\mathbf{A}$) are deterministically sampled. An alternative approach based on analyzing the *mutual coherence* of $\mathbf{A}$ produces a practically computable but suboptimal condition, see [11]. We will use a slightly more general definition of the NSP that was introduced in [15]:

Definition 1.1 (Null Space Property). Given $s \in \mathbb{N}$ and $c > 0$, a matrix $\mathbf{A} \in \mathbb{R}^{m \times n}$ satisfies the (s, c) -NSP in the quasi-norm ℓ_q ($0 < q \leq 1$) if for all $\mathbf{h} \in \ker(\mathbf{A})$ and $T \in [n]_s$, we have

$$\|\mathbf{h}_T\|_q^q < c \|\mathbf{h}_{T^c}\|_q^q. \quad (1.3)$$

Here, $[n]_s$ is the collection of all subsets of $\{1, \dots, n\}$ with cardinality at most s ,

$$[n]_s := \{T \subset [n] \mid |T| \leq s\}, \quad [n] := \{1, \dots, n\},$$

$\mathbf{h}_T$ is the restriction of $\mathbf{h}$ to the index set T , and $T^c := [n] \setminus T$.

Nearly all exact recovery conditions based on the RIP are probabilistic in nature. This means that such analysis typically is split into two major thrusts: (i) the first one establishes that (1.1) and (1.2) are equivalent for a class of sparse signals if $\mathbf{A}$ satisfies an RIP condition, and (ii) the second one proves that the RIP condition for a suitable random matrix $\mathbf{A}$ is achievable with high probability. Such random arguments appear to be necessary in practice for the RIP analysis in order to mitigate pathological measurement configurations.

Under proper randomness assumptions, an alternative approach that circumvents the RIP also yields fruitful results in the study of (1.2), see [47, 39]. This approach is more reliant on a geometric characterization of the nullspace of the measurements, and therefore could be potentially adapted to analyzing non-convex objectives with similar geometric interpretations. We take this approach for analysis in this paper.

Although (1.2) has attracted a lot of interest in the past decades, the community realized that ℓ_1 minimization is not as robust for computing sparsity-promoting solutions compared to other objective functions, in particular compared to other non-convex objectives. This motivates the study of non-convex relaxation methods (use non-convex objectives to approximate $\|\cdot\|_0$), which are believed to be more sparsity-aware. Many non-convex objective functions, such as ℓ_q ($0 < q < 1$) [18, 8, 15], reweighted ℓ_1 [7], CoSaMP [27], IHT [2], $\ell_1 - \ell_2$ [44, 45], and ℓ_1/ℓ_2 [20, 21, 44, 30], are empirically shown to outperform ℓ_1 in certain contexts. However, relatively few such approaches have been successfully analyzed for theoretical recovery. In fact, obtaining exact recovery conditions and robustness analysis for general non-convex optimization problems is difficult, unless the objective possesses certain exploitable structure, see [25, 37].

We aim to investigate exact recovery conditions as well as the robustness for the objective ℓ_1/ℓ_2 in this paper. We are interested in providing conditions under which (1.1) is equivalent to the following problem:

$$\min \frac{\|\mathbf{x}\|_1}{\|\mathbf{x}\|_2} \quad \text{subject to } \mathbf{Ax} = \mathbf{b}. \quad (1.4)$$

To our knowledge, ℓ_1/ℓ_2 does not belong to any particular class of non-convex functions that has a systematic theory behind it. This is mainly because ℓ_1/ℓ_2 is neither convex nor concave, and is

not even globally continuous. However, there are a few observations that make this non-convex objective worth investigating. First, in [28] it was shown numerically that the ℓ_1/ℓ_2 outperforms ℓ_1 by a notable margin in jointly sparse recovery problems (in the sense that many fewer measurements are required to achieve the same recovery rate); particularly, ℓ_1/ℓ_2 admits a high-dimensional generalization called *orthogonal factor* and the corresponding minimization problem can be effectively solved using modern methods of manifold optimization. Understanding ℓ_1/ℓ_2 in one dimension would offer a baseline for its higher-dimensional counterparts. Secondly, in the matrix completion problem [5], one desires a matrix with minimal rank under the component constraints. Note that the rank of a matrix is the ℓ_0 measure of its singular value vector. A natural relaxation of rank to a more regular objectives include the so-called numerical intrinsic rank, which is defined by the ratio ℓ_1/ℓ_∞ of the singular value vector, and the numerical/stable rank, which is defined by the ratio ℓ_2/ℓ_∞ of the singular value vector. This suggests that the ratio between different norms might be a useful function to measure sparsity (complexity) of an object, and therefore leads us to study the objective ℓ_1/ℓ_2 in compressed sensing.

A few attempts have been made recently to reveal both the theory and applications behind the ℓ_1/ℓ_2 problem [23, 14, 44, 30, 42]. However, the existing analysis is either applicable only for non-negative signals, or yields a local optimality condition which is often too strict in practice. The investigation of efficient algorithms for solving the ℓ_1/ℓ_2 minimization is also an active area of research [30, 43, 3, 46, 35, 41].

Our contributions in this paper are two-fold. First we propose a new local optimality criteria which provides evidence that a large “dynamic range” may lead to better performance of an ℓ_1/ℓ_2 procedure, as was observed in [43]. We also conduct a first attempt at analyzing the exact recovery condition (global optimality) of ℓ_1/ℓ_2 ; a sufficient condition for uniform recoverability as well as some analysis of the robustness to noise are also given. We also provide numerical demonstrations, in which a novel initialization step for the optimization is proposed and explored to improve the performance of existing algorithms. We remark that since this problem is non-convex, none of the results in this paper are tight; they only serve as the initial insight into certain aspects of the method that have been observed in practice.

The rest of the paper is organized as follows. In Section 2 we briefly introduce the results in [47] and [39] obtained by a high-dimensional geometry approach, which are relevant to our analysis. In Section 3, we give a new local optimality condition ensuring that an s -sparse vector is the local minimizer of the ℓ_1/ℓ_2 problem. In Section 4 we investigate the uniform recoverability of ℓ_1/ℓ_2 and propose a new sufficient condition for this recoverability. We also show that this condition is easily satisfied for a large class of random matrices. In Section 5, we give some elementary analysis on how the solution to ℓ_1/ℓ_2 minimization problem is affected in the presence of noise. In Section 6, we provide some numerical experiments to support our findings and propose a novel initialization technique that further improves an existing ℓ_1/ℓ_2 algorithm from [43]. In Section 7, we summarize our findings as well as point out some possible directions for future investigation.

2 A geometric perspective on ℓ_1 minimization

Geometric interpretation of compressed sensing first appeared in an abstract formulation of the problem in Donoho’s original work [12]. In this section, we will take a selection of geometric views on ℓ_1 minimization based on the discussions in [39] and [47], which do not hinge on RIP analysis. We will see that they provide valuable insight for our analysis of (1.1) in the case of non-convex relaxation.

To interpret (1.2) geometrically, we assume that entries of $\mathbf{A}$ are iid standard normal, i.e., $(\mathbf{A})_{i,j} \sim \mathcal{N}(0,1)$. In this case, an $\mathbf{x}$ solving (1.2) belongs to the translate of a subspace that is uniformly drawn from the Grassmanian $G_{n-m,n}$, where,

$$G_{r,n} = \{A \subset \mathbb{R}^n \mid A \text{ is an } r \text{ dimensional subspace}\}.$$

The objective function, on the other hand, can be considered as an origin-centered symmetric convex body, i.e., a scaled ℓ_1 ball in $\mathbb{R}^n$. Therefore, (1.2) is associated to a problem of under-

standing the random section of a convex set in $\mathbb{R}^n$. We thus seek to understand random sections of convex sets.

We began by introducing the approach in [39]. Visualization of convex sets in high dimensions often depends on two parts: the bulk and the outliers. The bulk is the largest inscribed part of a convex set that resembles an ellipsoid, and the outliers are those points outside the bulk contributing to the diameter of the set. As the name outliers suggests, a random low-dimensional section of a convex set tends to avoid outliers, and the resulting shape is close to the section of the bulk, i.e., an ellipsoid. As the dimension increases, the random section is more likely to capture the outliers, and the diameter of the intersected region will grow in a manner determined by the geometry of the convex set. This is made precise by the following theorems:

Theorem 2.1. [39, Theorem 3.3, Dvoretzky] *Let $0 < \varepsilon, \delta < 1$ be two fixed numbers and $d \in \mathbb{N}$. Let K be an origin-centered symmetric convex body in $\mathbb{R}^n$ such that the largest ellipsoid inscribed in it is the unit Euclidean ball. Let E be a random subspace drawn uniformly from $G_{d,n}$. There exists $R > 0$ which only depends on K such that with probability at least $1 - \delta$,*

$$(1 - \varepsilon)B(R) \subset K \cap E \subset (1 + \varepsilon)B(R), \quad (2.1)$$

provided that $d \leq C(\varepsilon, \delta) \log n$, where $B(R)$ is the Euclidean ball of radius R in E and $C(\varepsilon, \delta)$ is a constant depending only on ε and δ . The condition on d can be improved to $d \leq C(\varepsilon, \delta)n$ when K is the ℓ_1 ball with radius $\sqrt{n}$ and $R = 1$.

Theorem 2.2. [39, Theorem 5.1, M^* bound] *Let K be a bounded subset of $\mathbb{R}^n$. Let E be a random subspace drawn uniformly from $G_{n-m,n}$. Then,*

$$\mathbb{E} \sup_{\mathbf{u} \in K \cap E} \|\mathbf{u}\|_2 \leq \sqrt{\frac{8\pi}{m}} \cdot \mathbb{E} \sup_{\mathbf{u} \in K} |\langle \mathbf{g}, \mathbf{u} \rangle|, \quad (2.2)$$

where $\mathbf{g} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$.

Theorem 2.1 and 2.2 describe the low-dimensional and high-dimensional section of an origin-centered symmetric convex body, respectively. Their proofs can be found in [40], and the idea of the latter is crucial for the practicality part of our result. The quantity $\mathbb{E} \sup_{\mathbf{u} \in K} |\langle \mathbf{g}, \mathbf{u} \rangle|$ on the right-hand side of (2.2) is closely related to the concept of Gaussian width or Gaussian complexity of a set K :

Definition 2.3 (Gaussian complexity). Let K be a bounded set in $\mathbb{R}^n$. The Gaussian width of K is defined by $w(K) = \mathbb{E} \sup_{\mathbf{x} \in K} \langle \mathbf{g}, \mathbf{x} \rangle$, where $\mathbf{g} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$. The Gaussian complexity of K is defined as $w'(K) = w(K - K)$, where $K - K$ is the Minkowski difference between K and itself.

It is easy to check that the Gaussian width of a set remains unchanged after taking the convex hull, i.e., $w(K) = w(\text{conv}(K))$, implying an immediate upper bound for the Gaussian width of the unit ℓ_1 ball B_1^n in $\mathbb{R}^n$ (B_1^n is the convex hull of the set $\{\pm \mathbf{e}_i, i \leq n\}$, where $\mathbf{e}_i$ is the i -th unit vector in $\mathbb{R}^n$):

$$w(B_1^n) = \mathbb{E} \max_{i \leq n} |(\mathbf{g})_i| \leq \sqrt{8 \log n}, \quad (2.3)$$

where $\mathbf{g} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$.

Indeed, when K is a symmetric convex body centered at origin, then $K - K = 2K$, so that (2.2) implies

$$\mathbb{E} \text{diam}(K \cap E) = 2 \mathbb{E} \sup_{\mathbf{u} \in K \cap E} \|\mathbf{u}\|_2 \leq \sqrt{\frac{8\pi}{m}} \cdot 2 \mathbb{E} \sup_{\mathbf{u} \in K} \langle \mathbf{g}, \mathbf{u} \rangle = \sqrt{\frac{8\pi}{m}} \cdot w'(K). \quad (2.4)$$

Note that $\sup_{\mathbf{u} \in K - K} \langle \mathbf{g}, \mathbf{u} \rangle$ is the distance between two hyperplanes (with normal direction $\mathbf{g}$) that exactly sandwich K . $w'(K)$ can therefore be interpreted as the average width of K under the Gaussian measure, which is a geometric attribute of K measuring its complexity. As was observed in [39], Theorem 2.2 implies the following average relative recovery error estimate in ℓ_1 minimization:

Theorem 2.4. Let $\mathbf{x}^*$ be the solution to (1.2), and $\mathbf{x}_0$ be an s -sparse signal satisfying $\mathbf{A}\mathbf{x}_0 = \mathbf{b}$ and $\|\mathbf{x}_0\|_0 = s$. Then,

$$\mathbb{E} \frac{\|\mathbf{x}^* - \mathbf{x}_0\|_2}{\|\mathbf{x}_0\|_2} \lesssim \sqrt{\frac{s \log n}{m}}, \quad (2.5)$$

where $a \lesssim b$ means that $a \leq Cb$ for a universal constant C .

Proof. Let $K_1 = \|\mathbf{x}_0\|_1 \cdot B_1^n$, where B_1^n is the unit ℓ_1 ball in $\mathbb{R}^n$. By definition, $\|\mathbf{x}^*\|_1 \leq \|\mathbf{x}_0\|_1$ so that $\mathbf{x}^* \in K_1$. Therefore, $\mathbf{x}^* - \mathbf{x}_0 \in (K_1 - K_1) \cap \ker(\mathbf{A})$. It follows immediately from (2.4) with $K = K_1 - K_1$ and $E = \ker(\mathbf{A})$ that

$$\begin{aligned} \mathbb{E} \|\mathbf{x}^* - \mathbf{x}_0\|_2 &\leq \frac{1}{2} \mathbb{E} \text{diam}(K \cap E) \leq \sqrt{\frac{2\pi}{m}} \cdot w'(K) \\ &\leq \sqrt{\frac{8\pi}{m}} \cdot w'(K_1) \lesssim \|\mathbf{x}_0\|_1 \cdot \sqrt{\frac{\log n}{m}} \lesssim \|\mathbf{x}_0\|_2 \cdot \sqrt{\frac{s \log n}{m}}, \end{aligned}$$

where the penultimate inequality uses $w'(B_1^n) = 2w(B_1^n)$ and (2.3), and the last inequality follows from $\|\mathbf{x}_0\|_0 \leq s$ and the Cauchy-Schwarz inequality. Dividing $\|\mathbf{x}_0\|_2$ on both sides finishes the proof. $\square$

Remark 2.5. Taking $m \gtrsim s \log n$ in (2.5) results in a bound on the right-hand side of (2.5), which, up to logarithmic factors, achieves the desired statement that m measurements can recover s -sparse signals. Note that the statement in (2.5) is only concerned with the *average* relative error. One can go further to obtain a (pathwise) exact recovery result using Gordon Escape Theorem [39, 32]. However, the ideas from the proof of this result depend on the convex nature of the problem and are not extensible to non-convex cases, so we do not state it here.

An alternative approach to achieve an exact recovery condition for the ℓ_1 minimization problem is to interpret the kernel of $\mathbf{A}$ as a random subspace under the Gaussian assumption of the measurements and is RIP-free [47]. This method is similar to ideas described above in [39]. In fact, this RIP-free approach yields nearly all results that can be attained by RIP approaches. The analysis in [47] is the inspiration for our approach, so we shall summarize its main results.

The first idea is to note that a sufficient condition for $\mathbf{A}$ to satisfy the $(s, 1)$ -NSP is given by

$$\inf_{0 \neq \mathbf{h} \in \ker(\mathbf{A})} \frac{\|\mathbf{h}\|_1}{\|\mathbf{h}\|_2} > 2\sqrt{s}. \quad (2.6)$$

The condition (2.6) is concerned with the ratio between the ℓ_1 and ℓ_2 norms in a random subspace of dimension $n - m$, which can be analyzed using the tools from high-dimensional geometry. Indeed, a classical result in geometric functional analysis states that if the measurement $\mathbf{A}$ has iid Gaussian entries, then

$$\inf_{0 \neq \mathbf{h} \in \ker(\mathbf{A})} \frac{\|\mathbf{h}\|_1}{\|\mathbf{h}\|_2} > \frac{c\sqrt{m}}{\sqrt{1 + \log(n/s)}} \quad (2.7)$$

holds with overwhelming probability, where c is a dimension-free constant [17, 22]. Relations (2.6) and (2.7) together give a bound on m that is asymptotically equivalent to the classic result in [6].

A condition similar to (2.6) is given in [47] to guarantee that ℓ_1 minimization is stable. A specialization of the result reads as follows:

Theorem 2.6. Let $\tilde{\mathbf{x}}$ be the solution to the following minimization problem:

$$\min \|\mathbf{x}\|_1 \quad \text{subject to } \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2 \leq \varepsilon,$$

where $\mathbf{b} := \mathbf{A}\mathbf{x}_0 + \mathbf{e}$ with $\|\mathbf{e}\|_2 \leq \varepsilon$ and $\|\mathbf{x}_0\|_0 \leq s$. Let $\mathbf{u}$ and $\mathbf{w}$ be the orthogonal projections of $\tilde{\mathbf{x}} - \mathbf{x}_0$ to $\ker(\mathbf{A})$ and $\ker^\perp(\mathbf{A})$, respectively. If

$$s = \frac{v^2}{4} \frac{\|\mathbf{u}\|_1^2}{\|\mathbf{u}\|_2^2}$$

for some $v \in (0, 1)$, then for either $p = 1$ or $p = 2$,

$$\|\tilde{\mathbf{x}} - \mathbf{x}_0\|_p \leq 2\gamma_p \left(1 + \frac{1 + v\sqrt{2 - v^2}}{1 - v^2} \right) \|\mathbf{w}\|_2,$$

where $\gamma_1 = \sqrt{n}$ and $\gamma_2 = 1$.

It was also shown in [47] that $\|\mathbf{w}\|_2$ can be further bounded by $\varepsilon \|\mathbf{R}^{-T}\|_2$, where $\mathbf{R}$ is the triangular matrix in the QR decomposition of $\mathbf{A}^T$. It is worth noting that Theorem 2.6 is not implied by the RIP results in [6]. In fact, the constants involved in Theorem 2.6 are more revealing since they are directly related to the sparsity level s rather than the RIP parameters, which are not invariant under invertible transforms.

3 A local optimality criteria

In this section, we give a sufficient condition for an s -sparse signal $\mathbf{x}_0$ to be the local minimizer of (1.4) with $\mathbf{b} := \mathbf{A}\mathbf{x}_0$. Compared to the global optimality condition obtained later in Section 4, the local optimality condition in this section aids in understanding the behavior of ℓ_1/ℓ_2 optimization near $\mathbf{x}_0$. This is important in practice since many non-convex algorithms only have local convergence guarantees. Our local optimality result is signal-dependent but it offers asymptotically weaker and more interpretable conditions than those in [30]. Our characterization of local optimality depends on the (inverse) *dynamic range* $\rho = \rho(\mathbf{x})$ of a nontrivial vector $\mathbf{x}$, defined as

$$\rho := \frac{\min_{i \in \text{supp}(\mathbf{x})} |x_i|}{\max_{i \in \text{supp}(\mathbf{x})} |x_i|} = \frac{\min_{i \in \text{supp}(\mathbf{x})} |x_i|}{\|\mathbf{x}\|_\infty} \quad (3.1)$$

Smaller values of ρ indicate larger variation in the magnitude of the extremal nonzero entries in $\mathbf{x}_0$. It was observed in [43] that the performance of ℓ_1/ℓ_2 improves when the dynamic range increases. We will also need a quantity κ that is ratio of norm ratios:

$$\kappa = \kappa(\mathbf{x}_0) := \frac{\|\mathbf{x}_0\|_1 \|\mathbf{x}_0\|_\infty}{\|\mathbf{x}_0\|_2^2} = \frac{\frac{\|\mathbf{x}_0\|_1}{\|\mathbf{x}_0\|_2}}{\frac{\|\mathbf{x}_0\|_\infty}{\|\mathbf{x}_0\|_2}}. \quad (3.2)$$

The main result is the following:

Theorem 3.1 (Local optimality). *Let $\mathbf{x}_0$ be a nonzero s -sparse vector ($s > 1$), and $\mathbf{A}$ be a measurement matrix. Define*

$$c = c(\mathbf{A}) := \sup_{0 \neq \mathbf{h} \in \ker(\mathbf{A})} \frac{\|\mathbf{h}\|_2^2}{\|\mathbf{h}\|_1^2}. \quad (3.3)$$

Suppose that $\mathbf{x}_0, \mathbf{A}$ are such that

$$\rho(\kappa + 1) \leq \frac{1}{2c}, \quad (3.4)$$

where $\rho = \rho(\mathbf{x}_0)$ and $\kappa = \kappa(\mathbf{x}_0)$, and that $\mathbf{A}$ satisfies the NSP with parameters $(s, \frac{1}{2\kappa+1})$, i.e.,

$$\|\mathbf{h}_T\|_1 < \frac{1}{2\kappa+1} \|\mathbf{h}_{T^c}\|_1 \quad \text{for every } \mathbf{h} \in \ker(\mathbf{A}) \text{ and } T \subset [n]_s \quad (3.5)$$

Then $\mathbf{x}_0$ is the local minimizer of the constrained ℓ_1/ℓ_2 objective function with ℓ_1 convergence radius $\delta = \rho \|\mathbf{x}_0\|_\infty$, i.e., for any $\mathbf{x} \in \mathbb{R}^n$ satisfying $\mathbf{A}\mathbf{x} = \mathbf{A}\mathbf{x}_0$ and $\|\mathbf{x} - \mathbf{x}_0\|_1 \leq \delta$, then $\|\mathbf{x}_0\|_1/\|\mathbf{x}_0\|_2 < \|\mathbf{x}\|_1/\|\mathbf{x}\|_2$.

We prove this theorem later in this section, but first focus on some of its consequences. Theorem 3 is initially difficult to fully comprehend since the conditions (3.4) and (3.5) not only depend on the measurement matrix $\mathbf{A}$ but also on the sparse vector $\mathbf{x}_0$. However, a specialization is more transparent: A worst-case upper bound for κ results in a local optimality condition that is uniformly true for all s -sparse vectors.

Corollary 3.2. Assume $s > 6$. If $\mathbf{A}$ satisfies the NSP with parameters $(s, \frac{1}{\sqrt{s+2}})$, then $\mathbf{x}_0$ is a local minimizer of ℓ_1/ℓ_2 for all $\|\mathbf{x}_0\|_0 \leq s$.

Proof. Note that for any s -sparse $\mathbf{x}_0$,

$$\kappa \leq \frac{(\sqrt{s}+1)}{2}, \quad \rho \leq 1. \quad (3.6)$$

Indeed, since $\kappa(\mathbf{x}_0)$ is invariant under both scaling and permutation, we may assume $\mathbf{x}_0 = (x_1, \dots, x_s, 0, \dots, 0)^T$ with $x_1 \geq \dots \geq x_s \geq 0$ and $\|\mathbf{x}_0\|_2 = 1$. In this case, $\kappa(\mathbf{x}_0) = x_1^2 + x_1(\sum_{i=2}^s x_i) \leq x_1^2 + x_1\sqrt{(s-1)(1-x_1^2)}$, with equality achieved at $x_2 = \dots = x_s$. The maximum of the right-hand side is $(\sqrt{s}+1)/2$ and is attained when $x_1 = 1/2 + 1/(2\sqrt{s})$.

Since $\mathbf{A}$ satisfies the $(s, \frac{1}{\sqrt{s+2}})$ -NSP, then

$$\|\mathbf{h}_T\|_1 < \frac{1}{\sqrt{s}+2} \|\mathbf{h}_{T^c}\|_1 \quad \forall (\mathbf{h}, T) \in \ker(\mathbf{A}) \times [n]_s, \quad (3.7)$$

which combines with (3.6) to establish (3.5). Now we note that if

$$c \leq \frac{1}{\sqrt{s}+3}. \quad (3.8)$$

is achieved for all s -sparse vectors, then this along with (3.6) implies (3.4). We claim that (3.7) implies (3.8) for $s > 6$: Let $\mathbf{h} \in \ker(\mathbf{A})$ and note that (3.7) implies that

$$\|\mathbf{h}\|_1^2 \geq (\sqrt{s}+3)^2 \|\mathbf{h}_T\|_1^2, \quad \mathbf{h} \in \ker(\mathbf{A}).$$

For $r \in \mathbb{N}$, let T_r be the support of the $((r-1)s+1)$ -th to the rs -th components of $\mathbf{h}$ arranged in decreasing magnitude. This partition ensures that

$$\|\mathbf{h}_{T_r}\|_\infty \leq \min_{i \in T_{r-1}} |h_i| \leq \frac{1}{s} \sum_{i \in T_{r-1}} |h_i| = \frac{1}{s} \|\mathbf{h}_{T_{r-1}}\|_1. \quad (3.9)$$

Then applying a block-type argument, for $s > 6$,

$$\begin{aligned} \|\mathbf{h}\|_2^2 &= \|\mathbf{h}_{T_1}\|_2^2 + \sum_{r \geq 2} \|\mathbf{h}_{T_r}\|_2^2 \leq \frac{1}{(\sqrt{s}+3)^2} \|\mathbf{h}\|_1^2 + \sum_{r \geq 2} \|\mathbf{h}_{T_r}\|_2^2 \\ &\stackrel{(3.9)}{\leq} \frac{1}{(\sqrt{s}+3)^2} \|\mathbf{h}\|_1^2 + \sum_{r \geq 2} \sum_{i=1}^s \left(\frac{\|\mathbf{h}_{T_{r-1}}\|_1}{s} \right)^2 \\ &\leq \frac{1}{(\sqrt{s}+3)^2} \|\mathbf{h}\|_1^2 + \frac{1}{s} \sum_{r \geq 1} \|\mathbf{h}_{T_r}\|_1^2 \\ &= \frac{1}{(\sqrt{s}+3)^2} \|\mathbf{h}\|_1^2 + \frac{1}{s} \|\mathbf{h}\|_1^2 \leq \frac{1}{\sqrt{s}+3} \|\mathbf{h}\|_1^2. \end{aligned}$$

We have thus established both (3.4) and (3.5), so that the conclusion of Theorem 3.1 holds. $\square$

A similar technique does not, unfortunately, provide a uniform bound on the local convergence radius due to technical issues. Without asking for uniformity of the local convergence radius, Corollary 3.2 gives an asymptotic weaker condition for local optimality of sparse vectors compared to the result in [30], which requires a stronger NSP with parameters $(s, \frac{1}{s+1})$ for $s \geq 1$. In fact, in many situations of interest, κ is of mild order (which is specified in the following proposition), suggesting that (3.5) is not as restrictive as in Corollary 3.2.

Proposition 3.3. Suppose that the s nonzero components of $\mathbf{x}_0$ are iid with the same distribution as a scalar centered sub-gaussian random variable X . Then, for sufficiently large s , the following event holds with probability at least $1 - 1/s^2$:

$$\kappa(\mathbf{x}_0) \leq C\sqrt{\log s},$$

where C is a constant depending only on the sub-gaussian norm of $X/\sqrt{\mathbb{E}X^2}$.

Proof. Since κ is scale-invariant, we may assume that $\mathbb{E}X^2 = 1$, i.e., $\mathbb{E}|X| \leq 1$ and $\text{Var}X \leq 1$. Denote the sub-gaussian norm (see Definition 4.2) of X as σ_X . Applying Hoeffding's inequality to $\|\mathbf{x}_0\|_1$, Bernstein's inequality to $\|\mathbf{x}_0\|_2^2$ and the maximal inequality [31, Theorem 1.14] to $\|\mathbf{x}_0\|_\infty$ yields that for sufficiently large s ,

$$\begin{aligned}\mathbb{P}(\|\mathbf{x}_0\|_1 > 2s) &\leq \mathbb{P}(\|\mathbf{x}_0\|_1 - \mathbb{E}\|\mathbf{x}_0\|_1 > s\mathbb{E}|X|) \leq 2e^{-cs} \leq 1/3s^2 \\ \mathbb{P}(\|\mathbf{x}_0\|_2^2 < s/2) &\leq \mathbb{P}(\|\mathbf{x}_0\|_2^2 - \mathbb{E}\|\mathbf{x}_0\|_2^2 \leq -s\mathbb{E}X^2/2) \leq 2e^{-cs} \leq 1/3s^2 \\ \mathbb{P}(\|\mathbf{x}_0\|_\infty > 4\sigma_X\sqrt{\log s}) &\leq \mathbb{P}(\|\mathbf{x}_0 - \mathbb{E}\mathbf{x}_0\|_\infty > 3\sigma_X\sqrt{\log s}) \leq 1/3s^2,\end{aligned}$$

where c is a constant depending only on σ_X . It follows from a union bound that

$$\mathbb{P}(\kappa(\mathbf{x}_0) \leq 16\sigma_X\sqrt{\log s}) \geq 1 - 1/s^2.$$

The proof is complete. $\square$

We now provide the proof of Theorem 3.1:

Proof of Theorem 3.1. The main idea of the proof is that under condition (3.5), a large proportion of the perturbation of $\mathbf{x}_0$ by $\mathbf{h} \in \ker(\mathbf{A})$ will be well-spread outside the support of $\mathbf{x}_0$, which necessarily increases the value of the objective ℓ_1/ℓ_2 .

For any $\mathbf{x}, \mathbf{y} \neq 0$, define the ordering $\mathbf{x} \succ (\succ) \mathbf{y}$ as $\frac{\|\mathbf{x}\|_1}{\|\mathbf{x}\|_2} \geq (>) \frac{\|\mathbf{y}\|_1}{\|\mathbf{y}\|_2}$. For simplicity and without loss of generality, we assume the nonzero entries of $\mathbf{x}_0$ are arranged in order of decreasing magnitude in the first s components, i.e., that $\mathbf{x}_0 = (x_1, \dots, x_s, 0, \dots, 0)^T \in \mathbb{R}^n$ with $|x_1| \geq \dots \geq |x_s| > 0$. Define $\delta := \rho\|\mathbf{x}_0\|_\infty = |x_s| > 0$.

We claim that for any $\mathbf{h} = (h_1, \dots, h_n)^T \in \ker(\mathbf{A})$ with $\|\mathbf{h}\|_1 \leq \delta$, the perturbed vector $\mathbf{x}$ satisfies

$$\mathbf{x} := \mathbf{x}_0 + \mathbf{h} = (x_1 + h_1, \dots, x_s + h_s, h_{s+1}, \dots, h_n)^T \succ \mathbf{x}_0,$$

which would prove the desired result. To show the above relation, we will construct another vector $\mathbf{x}'$ from $\mathbf{x}$ and establish the ordering,

$$\mathbf{x} \succ \mathbf{x}' \succ \mathbf{x}_0. \quad (3.10)$$

To begin, introduce $\beta = \sum_{i=1}^s \text{sgn}(x_i)h_i$ and $\gamma = \sum_{i=1}^s |h_i|$, and augment entries 1 and s in $\mathbf{x}$ to obtain

$$\mathbf{x}' := \left(x_1 + \text{sgn}(x_1)\frac{\gamma + \beta}{2}, x_2, \dots, x_{s-1}, x_s - \text{sgn}(x_s)\frac{\gamma - \beta}{2}, h_{s+1}, \dots, h_n \right)^T.$$

Note that since $|\beta| \leq \gamma$, then

$$\begin{aligned}\|\mathbf{x}'\|_1 &= \left| x_1 + \text{sgn}(x_1)\frac{\gamma + \beta}{2} \right| + \left| x_s - \text{sgn}(x_s)\frac{\gamma - \beta}{2} \right| + \sum_{j=2}^{s-1} |x_j| + \sum_{j=s+1}^n |h_j| \\ &= \beta + \sum_{j=1}^s |x_j| + \sum_{j=s+1}^n |h_j| = \sum_{j=1}^s |x_j| + \text{sgn}(x_j)h_j + \sum_{j=s+1}^n |h_j| \\ &= \sum_{j=1}^s |x_j + h_j| + \sum_{j=s+1}^n |h_j| = \|\mathbf{x}\|_1,\end{aligned} \quad (3.11)$$

where the penultimate equality uses the fact that the assumption $\|\mathbf{h}\|_1 \leq \delta$ implies $|h_i| \leq |x_i|$ for $i = 1, \dots, s$. To show that $\|\mathbf{x}'\|_2 \geq \|\mathbf{x}\|_2$, we express the difference of their squares as

$$\|\mathbf{x}'\|_2^2 - \|\mathbf{x}\|_2^2 = \underbrace{(\gamma + \beta)|x_1| - (\gamma - \beta)|x_s| - 2 \sum_{i=1}^s \text{sgn}(x_i h_i) |x_i| |h_i|}_{(A)} + \underbrace{\frac{1}{2}(\gamma^2 + \beta^2) - \sum_{i=1}^s h_i^2}_{(B)} \quad (3.12a)$$

Term (A) satisfies

$$(A) = 2 \sum_{\text{sgn}(x_i h_i)=1} (|x_1| - |x_i|)|h_i| + 2 \sum_{\text{sgn}(x_i h_i)=-1} (|x_i| - |x_s|)|h_i| \geq 0, \quad (3.12b)$$

and term (B)

$$\begin{aligned} (B) &= \frac{1}{2} \left(2 \sum_{i=1}^s h_i^2 + 2 \sum_{1 \leq i < j \leq s} (|h_i||h_j| - \text{sgn}(x_i x_j) h_i h_j) \right) - \sum_{i=1}^s h_i^2 \\ &= \sum_{1 \leq i < j \leq s} (|h_i||h_j| - \text{sgn}(x_i x_j) h_i h_j) \geq 0. \end{aligned} \quad (3.12c)$$

Relations (3.11) and (3.12) establish the upper ordering in (3.10). To show the lower ordering, we first note that Taylor's Theorem applied to the function $y \mapsto \sqrt{y}$ along with that function's concavity implies that for any $\beta > 0$:

$$\sqrt{\|\mathbf{x}_0\|_2^2 + \beta} \leq \|\mathbf{x}_0\|_2 + \frac{1}{2\|\mathbf{x}_0\|_2} \beta. \quad (3.13)$$

Now we directly compare the ℓ_1/ℓ_2 norms of $\mathbf{x}_0$ and $\mathbf{x}'$:

$$\begin{aligned} \frac{\|\mathbf{x}'\|_1}{\|\mathbf{x}'\|_2} &= \frac{\|\mathbf{x}_0\|_1 + \beta + \sum_{i=s+1}^n |h_i|}{\sqrt{\|\mathbf{x}_0\|_2^2 + (\gamma + \beta)|x_1| - (\gamma - \beta)|x_s| + \frac{1}{2}(\gamma^2 + \beta^2) + \sum_{i=s+1}^n h_i^2}} \\ &\stackrel{|\beta| \leq \gamma}{\geq} \frac{\|\mathbf{x}_0\|_1 - \gamma + \sum_{i=s+1}^n |h_i|}{\sqrt{\|\mathbf{x}_0\|_2^2 + 2\gamma|x_1| + \gamma^2 + \|\mathbf{h}\|_2^2}} \\ &\stackrel{(3.13)}{\geq} \frac{\|\mathbf{x}_0\|_1 - \gamma + \sum_{i=s+1}^n |h_i|}{\|\mathbf{x}_0\|_2 + \frac{1}{2\|\mathbf{x}_0\|_2} (2\gamma|x_1| + \gamma^2 + \|\mathbf{h}\|_2^2)} \\ &\stackrel{(3.5), (3.1), (3.3)}{>} \frac{\|\mathbf{x}_0\|_1 + \frac{2\kappa}{2\kappa+2} \|\mathbf{h}\|_1}{\|\mathbf{x}_0\|_2 + \frac{1}{2\|\mathbf{x}_0\|_2} \left(\frac{2}{2\kappa+2} \|\mathbf{x}_0\|_\infty \|\mathbf{h}\|_1 + \frac{\rho}{(2\kappa+2)^2} \|\mathbf{x}\|_\infty \|\mathbf{h}\|_1 + c\rho \|\mathbf{x}\|_\infty \|\mathbf{h}\|_1 \right)} \\ &\geq \min \left(\frac{\|\mathbf{x}_0\|_1}{\|\mathbf{x}_0\|_2}, \frac{\frac{4\kappa}{2\kappa+2}}{\frac{2}{2\kappa+2} + \frac{\rho}{(2\kappa+2)^2} + c\rho} \frac{\|\mathbf{x}_0\|_2}{\|\mathbf{x}_0\|_\infty} \right) \\ &\stackrel{(3.2)}{\geq} \min \left(1, \frac{\frac{4}{2\kappa+2}}{\frac{2}{2\kappa+2} + \frac{\rho}{(2\kappa+2)^2} + c\rho} \right) \frac{\|\mathbf{x}_0\|_1}{\|\mathbf{x}_0\|_2} \\ &\stackrel{(3.4)}{=} \frac{\|\mathbf{x}_0\|_1}{\|\mathbf{x}_0\|_2}. \end{aligned}$$

We have thus established the lower relation in (3.10) and the proof is complete. $\square$

Remark 3.4. From Theorem 3.1, we notice that (3.4) is automatically satisfied if $\rho < 1/(\sqrt{s}+3)$. This suggests that the local optimality criteria is more likely to hold for vectors with a larger dynamic range, which is a possible explanation for why large dynamic range aids in the numerical performance of ℓ_1/ℓ_2 algorithms [43]. We will also numerically investigate this in Section 6.

Remark 3.5. It is possible to derive a sufficient and necessary condition for an s -sparse solution $\mathbf{x}_0$ to be a local minimum by taking directional derivatives, which is the approach taken in [30]. Without loss of generality we assume that $\|\mathbf{x}_0\|_2 = 1$. For $\mathbf{h} \in \ker(\mathbf{A})$, consider the function $L_{\mathbf{h}}(t) = \|\mathbf{x}_0 + t\mathbf{h}\|_1 / \|\mathbf{x}_0 + t\mathbf{h}\|_2, t \geq 0$. $\mathbf{x}_0$ is a local minimizer of the ℓ_1/ℓ_2 minimization subject to $\mathbf{A}\mathbf{x} = \mathbf{A}\mathbf{x}_0$ if and only if $L'_{\mathbf{h}}(0) \geq 0$ for all $\mathbf{h} \in \ker(\mathbf{A})$, and this is equivalent to the following condition:

$$\|\mathbf{h}_{S^c}\|_1 \geq \langle (\|\mathbf{x}_0\|_1(\mathbf{x}_0)_S - \text{sgn}((\mathbf{x}_0)_S), \mathbf{h}_S) \rangle \quad \forall 0 \neq \mathbf{h} \in \ker(\mathbf{A}),$$

where S is the support of $\mathbf{x}_0$. An unconditional upper bound on the right-hand side (independent of $\mathbf{x}_0$) is $(\sqrt{s} + 1)\|\mathbf{h}_S\|_1$, from which one can deduce that the NSP with parameters $(s, \frac{1}{\sqrt{s}+1})$ is sufficient to guarantee the uniform local optimality of s -sparse signals. This implies the uniform local optimality result in Corollary 3.2. However, the NSP condition is not tight since the equality cannot be attained, i.e., $\langle \text{sgn}((\mathbf{x}_0)_S), \mathbf{h}_S \rangle = \|\mathbf{h}_S\|_1$ implies $\langle \|\mathbf{x}_0\|_1(\mathbf{x}_0)_S, \mathbf{h}_S \rangle < 0$. As an alternative, Theorem 3.1 makes an attempt to understand the local optimality of a given sparse signal based on signal-dependent structure.

4 A sufficient condition for exact recovery

In this section, we propose a sufficient condition that guarantees the uniform exact recovery for sparse vectors using ℓ_1/ℓ_2 . As we will see, the condition to be obtained will hold with overwhelming probability for a large class of sub-gaussian random matrices. Since our approach for deriving this recoverability condition applies to other situations as well, we consider the following problem which is slightly more general than (1.4): For $0 < q \leq 1$ fixed, consider the optimization

$$\min \frac{\|\mathbf{x}\|_q}{\|\mathbf{x}\|_2} \quad \text{subject to } \mathbf{A}\mathbf{x} = \mathbf{b}. \quad (4.1)$$

We will refer to (4.1) as the ℓ_q/ℓ_2 minimization problem. The problem (4.1) is equivalent to (1.4) when $q = 1$. Clearly, (4.1) recovers all s -sparse vectors if and only if for any nonzero $\mathbf{h} \in \ker(\mathbf{A})$ and $\mathbf{x}_0$ with $\|\mathbf{x}_0\|_0 \leq s$,

$$\frac{\|\mathbf{x}_0\|_q^q}{\|\mathbf{x}_0\|_2^q} < \frac{\|\mathbf{x}_0 + \mathbf{h}\|_q^q}{\|\mathbf{x}_0 + \mathbf{h}\|_2^q}. \quad (4.2)$$

We will now directly work with (4.2) to find a sufficient condition establishing this property. As we will see, the condition to be obtained from our analysis is strictly sufficient and *stronger* than the NSP assumption used in the state-of-art ℓ_q recovery (in particular with $q = 1$). This is mainly due to the technical difficulty in analyzing the ratio form in a uniform way. The main exact recovery result is as follows.

Theorem 4.1 (Uniform recoverability). *If, for some $s \in \mathbb{N}$, the matrix $\mathbf{A}$ satisfies*

$$\inf_{\mathbf{h} \in \ker(\mathbf{A}) \setminus \{\mathbf{0}\}} \frac{\|\mathbf{h}\|_q}{\|\mathbf{h}\|_2} > 3^{1/q} s^{1/q-1/2}. \quad (4.3)$$

then (4.2) holds, establishing exact recovery for the optimization (4.1).

Condition (4.3) should be compared to (2.6) in Section 2. When $q = 1$, the right-hand side of (4.3) is $3\sqrt{s}$, which is slightly larger (worse) than $2\sqrt{s}$ in (2.6). When $\mathbf{A}$ is Gaussian, (2.7) ensures that (4.3) is satisfied with overwhelming probability provided that m behaves linearly in s (up to a logarithmic factor). The next theorem generalizes (2.7) to the class of isotropic sub-gaussian matrices, but with a slightly worse logarithmic factor. For convenience, we first recall the definition of sub-gaussian and isotropic vectors.

Definition 4.2 (Sub-gaussian random vectors). A random vector $X \in \mathbb{R}^n$ is said to be sub-gaussian if its one-dimensional marginal $\mathbf{a}^T X$ is sub-gaussian for all $\mathbf{a} \in \mathbb{R}^n$. In other words, X is sub-gaussian if $\|\mathbf{a}^T X\|_{\Psi_2} := \inf\{t > 0, \mathbb{E}[e^{|\mathbf{a}^T X|^2/t^2}] \leq 2\} < \infty$ for all $\mathbf{a} \in \mathbb{R}^n$, where the sub-gaussian norm $\|\cdot\|_{\Psi_2}$ of X is defined as

$$\|X\|_{\Psi_2} = \sup_{\|\mathbf{a}\|_2=1} \|\mathbf{a}^T X\|_{\Psi_2}.$$

If X is standard normal, all of its marginals are standard normal in 1D, therefore $\|X\|_{\Psi_2} = \|\mathcal{N}(0, 1)\|_{\Psi_2}$.

Definition 4.3 (Isotropic random vectors). A random vector $X \in \mathbb{R}^n$ is said to be isotropic if $\mathbb{E}[XX^T] = \mathbf{I}_n$, where $\mathbf{I}_n$ is the identity matrix.

Theorem 4.4 (Uniform recoverability for sub-gaussian matrices). *Let $\mathbf{A} \in \mathbb{R}^{m \times n}$ be a random matrix whose rows are independent, isotropic and sub-gaussian random vectors in $\mathbb{R}^n$. Suppose that s satisfies,*

$$\frac{m}{s} > DF^4 u \log n, \quad (4.4)$$

where F is the maximum sub-gaussian norm of rows of $\mathbf{A}$, $u \geq 1$, and D is an absolute constant. Then, for $q = 1$, (4.3) holds with probability at least $1 - 2e^{-u}$.

The class of sub-gaussian random matrices considered in Theorem 4.4 is general enough for practical purposes: Gaussian, symmetric Bernoulli, and many other random matrices whose entries are sampled from bounded distributions with appropriate centralization and normalization fall into this realm.

We now provide the proof of Theorem 4.1.

Proof of Theorem 4.1. To find a sufficient condition only, we may require the left-hand side of (4.2) to be smaller than a quantity which is unconditionally smaller than the right-hand side of (4.2). First recall the q -triangle inequality, which states that for $0 < q \leq 1$,

$$\|\mathbf{x}_1 + \mathbf{x}_2\|_q^q \leq \|\mathbf{x}_1\|_q^q + \|\mathbf{x}_2\|_q^q, \quad \forall \mathbf{x}_1, \mathbf{x}_2 \in \mathbb{R}^n.$$

Fix an $\mathbf{x}_0$ such that $\|\mathbf{x}_0\|_0 \leq s$, and let $S \subset [n]$ denote the support of $\mathbf{x}_0$. The right-hand side of (4.2) satisfies

$$\begin{aligned} \frac{\|\mathbf{x}_0 + \mathbf{h}\|_q^q}{\|\mathbf{x}_0 + \mathbf{h}\|_2^q} &\geq \frac{\|(\mathbf{x}_0 + \mathbf{h})_S\|_q^q + \|\mathbf{h}_{S^c}\|_q^q}{(\|\mathbf{x}_0\|_2 + \|\mathbf{h}\|_2)^q} \\ &\geq \frac{\|\mathbf{x}_0\|_q^q + \|\mathbf{h}_{S^c}\|_q^q - \|\mathbf{h}_S\|_q^q}{\|\mathbf{x}_0\|_2^q + \|\mathbf{h}\|_2^q}, \end{aligned} \quad (4.5)$$

where for the first inequality we used $\|\mathbf{x}_0 + \mathbf{h}\|_2 \leq \|\mathbf{x}_0\|_2 + \|\mathbf{h}\|_2$, and for the second inequality, we used the q -triangle inequality for the numerator, $\|(\mathbf{x}_0 + \mathbf{h})_S\|_q^q \geq \|\mathbf{x}_0\|_q^q - \|\mathbf{h}_S\|_q^q$, and concavity of $y \mapsto y^q$ for the denominator, $(\|\mathbf{x}_0\|_2 + \|\mathbf{h}\|_2)^q \leq \|\mathbf{x}_0\|_2^q + \|\mathbf{h}\|_2^q$. Continuing, we have

$$\begin{aligned} \frac{\|\mathbf{x}_0\|_q^q + \|\mathbf{h}_{S^c}\|_q^q - \|\mathbf{h}_S\|_q^q}{\|\mathbf{x}_0\|_2^q + \|\mathbf{h}\|_2^q} &\geq \min \left\{ \frac{\|\mathbf{x}_0\|_q^q}{\|\mathbf{x}_0\|_2^q}, \frac{\|\mathbf{h}_{S^c}\|_q^q - \|\mathbf{h}_S\|_q^q}{\|\mathbf{h}\|_2^q} \right\} \\ &= \min \left\{ \frac{\|\mathbf{x}_0\|_q^q}{\|\mathbf{x}_0\|_2^q}, \frac{\|\mathbf{h}\|_q^q - 2\|\mathbf{h}_S\|_q^q}{\|\mathbf{h}\|_2^q} \right\} \end{aligned}$$

for all nonzero $\mathbf{h} \in \ker(\mathbf{A})$. Thus, we have established that

$$\frac{\|\mathbf{x}_0 + \mathbf{h}\|_q^q}{\|\mathbf{x}_0 + \mathbf{h}\|_2^q} \geq \min \left\{ \frac{\|\mathbf{x}_0\|_q^q}{\|\mathbf{x}_0\|_2^q}, \frac{\|\mathbf{h}\|_q^q - 2\|\mathbf{h}_S\|_q^q}{\|\mathbf{h}\|_2^q} \right\}, \quad (4.6a)$$

where equality holds if and only if $\frac{\|\mathbf{x}_0\|_q^q}{\|\mathbf{x}_0\|_2^q} = \frac{\|\mathbf{h}\|_q^q - 2\|\mathbf{h}_S\|_q^q}{\|\mathbf{h}\|_2^q}$. Now by the generalized Hölder's inequality, we have

$$s^{1-q/2} \geq \frac{\|\mathbf{x}_0\|_q^q}{\|\mathbf{x}_0\|_2^q} \quad (4.6b)$$

which is a sharp estimate since we consider a general s -sparse $\mathbf{x}_0$. Therefore the conditions (4.6) in tandem with

$$\frac{\|\mathbf{h}\|_q^q - 2\|\mathbf{h}_S\|_q^q}{\|\mathbf{h}\|_2^q} > s^{1-q/2}. \quad (4.7)$$

are sufficient to conclude Theorem 4.1. In order to prove (4.7), we note that under assumption (4.3) then

$$\begin{aligned} \frac{\|\mathbf{h}\|_q^q}{\|\mathbf{h}\|_2^q} &\stackrel{(4.3)}{>} 3s^{1-q/2} \stackrel{(4.6b)}{\geq} s^{1-q/2} + 2 \frac{\|\mathbf{h}_S\|_q^q}{\|\mathbf{h}_S\|_2^q} \\ &\geq s^{1-q/2} + 2 \frac{\|\mathbf{h}_S\|_q^q}{\|\mathbf{h}\|_2^q}, \end{aligned}$$

where the second inequality uses the generalized Hölder inequality (4.6b) for the s -sparse vector $\mathbf{h}_S$. This inequality is equivalent to (4.7), proving Theorem 4.1. $\square$

Remark 4.5. In the derivation above, the condition (4.7) is very different from (4.2). Since (4.7) is scale-invariant, it holds for all $\mathbf{h} \in \ker(\mathbf{A})$ if and only if it holds for $\mathbf{h} \in \ker(\mathbf{A}) \cap K$, where K is any $\mathbf{0}$ -starshaped set in $\mathbb{R}^n$. E.g., K can be the unit ball in ℓ_q . However, for (4.2), it is generally not true that $\|\mathbf{x}_0\|_q/\|\mathbf{x}_0\|_2 < \|\mathbf{x}_0 + \mathbf{h}\|_q/\|\mathbf{x}_0 + \mathbf{h}\|_2$ implies the same inequality for $\mathbf{h}$ replaced by $t\mathbf{h}$, $t \in \mathbb{R} \setminus \{0\}$. This suggests that (4.7) is strictly stronger than (4.2).

Next we give the proof for Theorem 4.4, based on a matrix deviation inequality inspired by [40].

Proof. Under the assumption of Theorem 4.4, it follows from [40, Exercise 9.13] that for any bounded $T \subset \mathbb{R}^n$ and $u > 0$, with probability at least $1 - 2e^{-u^2}$,

$$\sup_{\mathbf{x} \in T} \left| \|\mathbf{A}\mathbf{x}\|_2 - \sqrt{m}\|\mathbf{x}\|_2 \right| \leq cF^2(\gamma(T) + u \cdot \text{rad}(T)), \quad (4.8)$$

where

$$\gamma(T) := \mathbb{E} \sup_{\mathbf{x} \in K} |\langle \mathbf{g}, \mathbf{x} \rangle|$$

is a variant of the Gaussian width $w(T)$ from Definition 2.3, $\text{rad}(T)$ is the radius of T , c is some absolute constant, and $F = \max_i \|\mathbf{A}_i\|_{\psi_2}$ with $\mathbf{A}_i$ the i -th row of $\mathbf{A}$.

Take $T = \ker(\mathbf{A}) \cap B_1^n$. In this case, T is symmetric with respect to the origin, so we have

$$\begin{aligned} \gamma(T) &= w(T) \leq w(B_1^n) \stackrel{(2.3)}{\leq} \sqrt{8 \log n} < 2\sqrt{\pi \log n}, \\ \text{rad}(T) &= \frac{1}{2} \text{diam}(T) \leq \frac{\sqrt{2\pi}}{2} w(T) \stackrel{(2.3)}{\leq} 2\sqrt{\pi \log n}. \end{aligned}$$

Using these in (4.8), we have the following inequality with probability at least $1 - 2e^{-u^2}$:

$$\sup_{\mathbf{x} \in \ker(\mathbf{A}) \cap B_1^n} \|\mathbf{x}\|_2 \leq 2cF^2\sqrt{\pi}(1+u) \sqrt{\frac{\log n}{m}} \leq 4cF^2\sqrt{\pi}u \sqrt{\frac{\log n}{m}}, \quad (4.9)$$

where the last inequality uses $u \geq 1$. Thus, with probability at least $1 - 2e^{-u^2}$,

$$\inf_{\mathbf{h} \in \ker(\mathbf{A}) \setminus \{\mathbf{0}\}} \frac{\|\mathbf{h}\|_1}{\|\mathbf{h}\|_2} \stackrel{(4.9)}{\geq} \frac{1}{4cF^2\sqrt{\pi}u} \sqrt{\frac{m}{\log n}} \stackrel{(4.4)}{>} 3\sqrt{s},$$

where we have taken $D = 144\pi c^2$ in our use of (4.4). Renaming u^2 by u leads to the final inequality (4.3). $\square$

Remark 4.6. Theorem 4.4 only shows that (4.3) is satisfied with high probability for isotropic subgaussian matrices when $q = 1$. It is interesting to know if the same holds true for all $q < 1$ and if the corresponding constant will get better. The answer to this question is not completely known to us, but there is some evidence that it might be true. Indeed, for fixed $q \leq 1$, if $\mathbf{A}$ satisfies

$$\inf_{\mathbf{x} \in \ker(\mathbf{A}) \setminus \{\mathbf{0}\}} \frac{\|\mathbf{x}\|_q}{\|\mathbf{x}\|_2} \geq c_q^{1/q} \left(\frac{m}{\log(n/m) + 1} \right)^{1/q-1/2} \quad (4.10)$$

for some $c_q > 0$, then

$$m \geq \left(\frac{c_q}{3}\right)^{1-q/2} s \log n$$

is sufficient for (4.3). The right-hand side of (4.10) is closely related to the Gelfand's widths of ℓ_p balls, see [16]. Assuming that $\mathbf{A}$ is a Gaussian random matrix and $n \geq m^2$, a result in [12] states that there exists $c_q > 0$ such that (4.10) is satisfied with probability approaching 1 as $n \rightarrow \infty$. However, it is not clear whether the best attainable c_q in (4.10) ensures $(c_q/3)^{1-q/2}$ as a decreasing function in q .

5 Robustness analysis

In this section, we discuss the robustness of ℓ_1/ℓ_2 minimization when noise is present. As in other compressed sensing results, we assume that $\mathbf{b}$ is contaminated by some noise $\mathbf{e}$: $\mathbf{b} = \mathbf{A}\mathbf{x}_0 + \mathbf{e} \in \mathbb{R}^m$, where $\mathbf{x}_0$ is a sparse vector. If the size (say the ℓ_2 norm) of $\mathbf{e}$ is bounded by an *a priori* known quantity ε , then the ℓ_1/ℓ_2 denoising problem can be stated as

$$\min \frac{\|\mathbf{x}\|_1}{\|\mathbf{x}\|_2} \quad \text{subject to } \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2 \leq \varepsilon. \quad (5.1)$$

Let $\mathbf{x}^*$ be a minimizer of (5.1). Then $\|\mathbf{A}\mathbf{x}^* - \mathbf{A}\mathbf{x}_0\|_2 \leq 2\varepsilon$, and necessarily,

$$\frac{\|\mathbf{x}^*\|_1}{\|\mathbf{x}^*\|_2} \leq \frac{\|\mathbf{x}_0\|_1}{\|\mathbf{x}_0\|_2} \leq \sqrt{s}. \quad (5.2)$$

This inequality will play an important role in our following discussion. Since ℓ_1/ℓ_2 is scale-invariant, it would be difficult to distinguish $\mathbf{x}^*$ from $\mathbf{x}_0$ when $\mathbf{x}^*$ is nearly parallel to $\mathbf{x}_0$ with similar magnitude. This behavior is quantified in our main robustness result:

Theorem 5.1 (Robustness). *Let $\mathbf{x}_0 \in \mathbb{R}^n$ be an s -sparse vector and $\mathbf{x}^*$ be a minimizer of (5.1). Let $\mathbf{x}^* - \mathbf{x}_0 = \mathbf{u} + \mathbf{w}$ with $\langle \mathbf{u}, \mathbf{w} \rangle = 0$ and assume that*

$$\beta := 4\sqrt{2s} \frac{\|\mathbf{u}\|_2}{\|\mathbf{u}\|_1} < 1. \quad (5.3)$$

Let $\alpha \in (\beta, 1)$ be any number. The following holds:

- *If both of the conditions*

$$\langle \mathbf{x}_0, \mathbf{x}^* \rangle \geq (1 - \alpha^2/2) \|\mathbf{x}_0\|_2 \|\mathbf{x}^*\|_2, \quad (5.4a)$$

$$\|\mathbf{x}_0\|_2 \leq \|\mathbf{x}^*\|_2 \leq (1 + \alpha) \|\mathbf{x}_0\|_2, \quad (5.4b)$$

hold, then

$$\|\mathbf{x}^* - \mathbf{x}_0\|_2 \leq 2\sqrt{\alpha} \|\mathbf{x}_0\|_2. \quad (5.5)$$

- *If at least one of (5.4) is violated, then*

$$\|\mathbf{x}^* - \mathbf{x}_0\|_p \leq \frac{2\alpha - \beta}{\alpha - \beta} \|\mathbf{w}\|_p, \quad (5.6)$$

for either $p = 1$ or $p = 2$.

We provide some intuitive interpretation of the results in Theorem 5.1. The conditions (5.4) codify the regime when suboptimal behavior of ℓ_1/ℓ_2 optimization is expected due to a confounding geometric positions of $\mathbf{x}^*$ and $\mathbf{x}_0$. If (5.4) is violated, then $\mathbf{x}^*$ either points to a different direction than $\mathbf{x}_0$ or has a different magnitude. In either case, it can be shown that $\|\mathbf{x}^* - \mathbf{x}_0\|_1 / \|\mathbf{x}^* - \mathbf{x}_0\|_2$ is reasonably small. This combined with the orthogonal decomposition $\mathbf{x}^* - \mathbf{x}_0 = \mathbf{u} + \mathbf{w}$ and the fact that $\mathbf{u}$ has large ℓ_1/ℓ_2 ratio implies that $\|\mathbf{u}\| \ll \|\mathbf{w}\|$ under some appropriate norm. This gives an informal justification for (5.6). Note that (5.6) provides an upper bound on $\|\mathbf{x}^* - \mathbf{x}_0\|_2$ relative to $\|\mathbf{w}\|_1$ regardless of the value of p since $\|\cdot\|_2 \leq \|\cdot\|_1$. We demonstrate the utility of Theorem 5.1 by showing how it can be used to show the robustness of (5.1).

Corollary 5.2. Let $\mathbf{A} \in \mathbb{R}^{m \times n}$ be an isotropic sub-gaussian random matrix, and let $\mathbf{x}^*$ be a solution to (5.1). Let $\alpha \in (0, 1)$ be such that for some $u > \log 5$,

$$\frac{n}{4} + 1 \geq m \geq K \frac{u}{\alpha^2} s \log n \quad (5.7)$$

where K is a constant depending only the sub-gaussian distribution of $\mathbf{A}$. Then with probability exceeding $1 - 5e^{-cu}$, for all s -sparse vectors $\mathbf{x}_0 \in \mathbb{R}^n$, at least one of the following is true: (i) (5.5) is true, or (ii)

$$\|\mathbf{x}^* - \mathbf{x}_0\|_p \leq C_p \varepsilon, \quad (5.8)$$

is true for either $p = 1$ or $p = 2$, where $C_1 = \sqrt{n}C_2$, and c, C_2 are constants depending only on the sub-gaussian distribution of $\mathbf{A}$.

Proof. Decompose $\mathbf{x}^* - \mathbf{x}_0 = \mathbf{u} + \mathbf{w}$ with $\mathbf{u} \in \ker(\mathbf{A})$ and $\mathbf{w} \in \ker(\mathbf{A})^\perp$. By assumption (5.7), we can apply Theorem 4.4 to conclude that with probability at least $1 - 2e^{-u}$, we have

$$\frac{\|\mathbf{u}\|_1}{\|\mathbf{u}\|_2} \geq \frac{8\sqrt{2}s}{\alpha} \implies \frac{\beta}{\alpha} \leq \frac{1}{2}, \quad (5.9)$$

with β as in (5.3). Note in particular that this implies $\beta < 1$. We can now apply Theorem 5.1, so that either (5.5) holds (as desired) or (5.6) holds. We now investigate (5.6):

$$\|\mathbf{x} - \mathbf{x}_0\|_p \leq \frac{2\alpha - \beta}{\alpha - \beta} \|\mathbf{w}\|_p \stackrel{(5.9)}{\leq} 3\|\mathbf{w}\|_p, \quad (5.10)$$

for either $p = 1$ or $p = 2$. In order to compute bounds for $\|\mathbf{w}\|_p$, we appeal to a well-known result on the lower bound of the singular value of sub-gaussian random matrices: Theorem 1.1 in [33] shows that the smallest nonzero singular value $\sigma_{\min}(\mathbf{A})$ of $\mathbf{A}$ is bounded below by some positive constant with high probability. More precisely,

$$\mathbb{P}\left(\sigma_n(\mathbf{A}) \geq 0.5C \left(1 - \sqrt{\frac{m-1}{n}}\right)\right) \geq 1 - e^{-cn} - 2^{-n-m+1} \geq 1 - 3e^{-c_1n}, \quad (5.11)$$

where $c, c_1, C > 0$ are absolute constants only depending on the sub-gaussian distribution of $\mathbf{A}$. Since $u < n$ (otherwise (5.7) implies more measurements than unknowns), then the events (5.11) and (5.9) occur simultaneously with probability at least $1 - 5e^{-cu}$ for some constant c . Under this simultaneous event, with $\mathbf{A}^\dagger$ the Moore-Penrose pseudoinverse of $\mathbf{A}$, then

$$\begin{aligned} \|\mathbf{w}\|_2 &= \|\text{Proj}_{\ker(\mathbf{A})}(\mathbf{x}_0 - \mathbf{x}^*)\|_2 = \|\mathbf{A}^\dagger \mathbf{A}(\mathbf{x}_0 - \mathbf{x}^*)\|_2 \leq \|\mathbf{A}^\dagger\|_2 \|\mathbf{A}(\mathbf{x}_0 - \mathbf{x}^*)\|_2 \\ &\leq 2\varepsilon \sigma_{\min}(\mathbf{A}) \leq 4C^{-1} \left(1 - \sqrt{\frac{m-1}{n}}\right)^{-1} \varepsilon \leq 8C^{-1} \varepsilon. \end{aligned}$$

holds with the probability on the right-hand side of (5.11). Using this in (5.10) with $p = 2$ yields (5.8). To show the $p = 1$ result, we first use the inequality $\|\mathbf{w}\|_1 \leq \sqrt{n}\|\mathbf{w}\|_2$ in (5.10), which yields (5.8) with $p = 1$. Note that this $\sqrt{n}$ factor is sharp when $\mathbf{A}$ is a Gaussian random matrix, cf. Theorem 2.1. $\square$

Note that both bounds (5.5) and (5.8) are comparable if one can choose $\alpha \lesssim \varepsilon^2$, but this unfortunately makes (5.7) quite restrictive. Similar issue also rises in the analysis of the ℓ_1 minimization in Theorem 2.4. Note also that (5.8) is consistent with the result in Theorem 2.6.

Now we provide the proof of Theorem 5.1:

Proof of Theorem 5.1. If both conditions (5.4) hold, then

$$\begin{aligned} \|\mathbf{x}^* - \mathbf{x}_0\|_2^2 &= \|\mathbf{x}^*\|_2^2 + \|\mathbf{x}_0\|_2^2 - 2\langle \mathbf{x}^*, \mathbf{x}_0 \rangle \\ &\leq (1 + (1 + \alpha)^2) \|\mathbf{x}_0\|_2^2 - 2 \left(1 - \frac{\alpha^2}{2}\right) \|\mathbf{x}_0\|_2^2 \leq 4\alpha \|\mathbf{x}_0\|_2^2, \end{aligned}$$

where the last inequality uses $\alpha^2 < \alpha$ for $0 < \alpha < 1$. Thus, we now restrict our attention to when (5.4) do not hold simultaneously; the goal in this case is to compute upper and lower bounds for $\|\mathbf{x}^* - \mathbf{x}_0\|_1 / \|\mathbf{x}^* - \mathbf{x}_0\|_2$, with the intuition for the upper bound being motivated by (5.2). In the sequel, we denote the violation of (5.4a) or (5.4b) as (5.4a)^G and (5.4b)^G, respectively. When (5.4) is violated, our desired upper bound will take the form

$$\frac{\|\mathbf{x}^* - \mathbf{x}_0\|_1}{\|\mathbf{x}^* - \mathbf{x}_0\|_2} \leq \frac{4\sqrt{s}}{\alpha}. \quad (5.12)$$

To show this, suppose first that (5.4a) is violated, then

$$\begin{aligned} \frac{\|\mathbf{x}^* - \mathbf{x}_0\|_1}{\|\mathbf{x}^* - \mathbf{x}_0\|_2} &\leq \frac{\|\mathbf{x}^*\|_1 + \|\mathbf{x}_0\|_1}{\sqrt{\|\mathbf{x}^*\|_2^2 + \|\mathbf{x}_0\|_2^2 - 2\langle \mathbf{x}^*, \mathbf{x}_0 \rangle}} \stackrel{(5.4a)^G}{\leq} \frac{\|\mathbf{x}^*\|_1 + \|\mathbf{x}_0\|_1}{\sqrt{\frac{\alpha^2}{2}(\|\mathbf{x}^*\|_2^2 + \|\mathbf{x}_0\|_2^2)}} \\ &\leq \frac{2}{\alpha} \cdot \frac{\|\mathbf{x}^*\|_1 + \|\mathbf{x}_0\|_1}{\|\mathbf{x}^*\|_2 + \|\mathbf{x}_0\|_2} \leq \frac{2}{\alpha} \cdot \frac{\|\mathbf{x}_0\|_1}{\|\mathbf{x}_0\|_2} \stackrel{(5.2)}{\leq} \frac{2}{\alpha} \sqrt{s} \leq \frac{4\sqrt{s}}{\alpha}, \end{aligned}$$

which is the desired inequality (5.12). A violation of the lower condition in (5.4b) in addition to (5.2) implies $\|\mathbf{x}^*\|_1 \leq \|\mathbf{x}_0\|_1$. We therefore split the case when (5.4b) is violated into two dichotomous sub-cases.

1. (5.4b)^G, case 1: $\|\mathbf{x}^*\|_1 \leq \|\mathbf{x}_0\|_1$. Let S be the support of $\mathbf{x}_0$. The following inequality is true:

$$\begin{aligned} \|\mathbf{x}_0\|_1 &\geq \|\mathbf{x}^*\|_1 = \|\mathbf{x}_0 + (\mathbf{x}^* - \mathbf{x}_0)\|_1 = \|(\mathbf{x}_0 + (\mathbf{x}^* - \mathbf{x}_0))_S\|_1 + \|\mathbf{x}_{S^c}^*\|_1 \\ &\geq \|\mathbf{x}_0\|_1 - \|(\mathbf{x}^* - \mathbf{x}_0)_S\|_1 + \|\mathbf{x}_{S^c}^*\|_1, \end{aligned}$$

so that we must have $\|\mathbf{x}_{S^c}^*\|_1 \leq \|(\mathbf{x}^* - \mathbf{x}_0)_S\|_1$. Therefore, the Cauchy-Schwarz inequality implies

$$\begin{aligned} \frac{\|\mathbf{x}^* - \mathbf{x}_0\|_1}{\|\mathbf{x}^* - \mathbf{x}_0\|_2} &\leq \sqrt{2} \cdot \frac{\|(\mathbf{x}^* - \mathbf{x}_0)_S\|_1 + \|\mathbf{x}_{S^c}^*\|_1}{\|(\mathbf{x}^* - \mathbf{x}_0)_S\|_2 + \|\mathbf{x}_{S^c}^*\|_2} \\ &\leq \sqrt{8} \cdot \frac{\|(\mathbf{x}^* - \mathbf{x}_0)_S\|_1}{\|(\mathbf{x}^* - \mathbf{x}_0)_S\|_2} \\ &\leq \sqrt{8s} \leq \frac{4\sqrt{s}}{\alpha}, \end{aligned}$$

which is the desired inequality (5.12).

2. (5.4b)^G, case 2: $\|\mathbf{x}^*\|_1 > \|\mathbf{x}_0\|_1$ and $\|\mathbf{x}^*\|_2 > (1 + \alpha)\|\mathbf{x}_0\|_2$. The condition $\|\mathbf{x}^*\|_1 > \|\mathbf{x}_0\|_1$ and (5.2) together imply $\|\mathbf{x}^*\|_2 > \|\mathbf{x}_0\|_2$. With this, we have

$$\begin{aligned} \frac{\|\mathbf{x}^* - \mathbf{x}_0\|_1}{\|\mathbf{x}^* - \mathbf{x}_0\|_2} &\leq \frac{\|\mathbf{x}_0\|_1 + \|\mathbf{x}^*\|_1}{\|\mathbf{x}^*\|_2 - \|\mathbf{x}_0\|_2} \\ &\leq \frac{\alpha + 1}{\alpha} \cdot \frac{\|\mathbf{x}_0\|_1 + \|\mathbf{x}^*\|_1}{\|\mathbf{x}^*\|_2} \\ &\leq \frac{\alpha + 1}{\alpha} \left(\frac{\|\mathbf{x}_0\|_1}{\|\mathbf{x}_0\|_2} + \frac{\|\mathbf{x}^*\|_1}{\|\mathbf{x}^*\|_2} \right) \\ &\leq \left(2 + \frac{2}{\alpha} \right) \sqrt{s} \leq 4 \frac{\sqrt{s}}{\alpha}, \end{aligned}$$

where the last inequality uses the fact that $\alpha < 1$.

When (5.4) is violated, we have established (5.12). We now begin the main part of the proof: decompose $\mathbf{x} = \mathbf{u} + \mathbf{w}$ as in the assumption. If $\|\mathbf{u}\|_p \leq \|\mathbf{w}\|_p$ for either $p = 1, 2$, then

$$\|\mathbf{x}_0 - \mathbf{x}^*\|_p \leq \|\mathbf{u}\|_p + \|\mathbf{w}\|_p \leq 2\|\mathbf{w}\|_p \leq \frac{2 - \beta}{1 - \beta} \|\mathbf{w}\|_p,$$

for any $\beta \in (0, 1)$, proving (5.6). Therefore, we can assume $\|\mathbf{u}\|_p > \|\mathbf{w}\|_p$ for both $p = 1, 2$. In this case,

$$\frac{\|\mathbf{x}^* - \mathbf{x}_0\|_1}{\|\mathbf{x}^* - \mathbf{x}_0\|_2} \geq \frac{\|\mathbf{u}\|_1 - \|\mathbf{w}\|_1}{\sqrt{\|\mathbf{u}\|_2^2 + \|\mathbf{w}\|_2^2}} \geq \frac{1 - \frac{\|\mathbf{w}\|_1}{\|\mathbf{u}\|_1}}{\sqrt{1 + \frac{\|\mathbf{w}\|_2^2}{\|\mathbf{u}\|_2^2}}} \cdot \frac{\|\mathbf{u}\|_1}{\|\mathbf{u}\|_2} \geq h(v) \frac{\|\mathbf{u}\|_1}{\|\mathbf{u}\|_2},$$

where

$$h(v) := \frac{1 - v}{\sqrt{1 + v^2}}, \quad v := \max \left\{ \frac{\|\mathbf{w}\|_1}{\|\mathbf{u}\|_1}, \frac{\|\mathbf{w}\|_2}{\|\mathbf{u}\|_2} \right\}. \quad (5.13)$$

If $h(v) \frac{\|\mathbf{u}\|_1}{\|\mathbf{u}\|_2} > \frac{4\sqrt{s}}{\alpha}$, then this contradicts (5.12), so that (5.4) must hold, showing (5.5). Therefore it only remains to consider when $h(v) \frac{\|\mathbf{u}\|_1}{\|\mathbf{u}\|_2} \leq \frac{4\sqrt{s}}{\alpha}$. Since $h(v) \geq \frac{1}{\sqrt{2}}(1 - v)$ for $v > 0$, then

$$\frac{\beta}{\alpha\sqrt{2}} = \frac{4\sqrt{s}}{\alpha} \frac{\|\mathbf{u}\|_2}{\|\mathbf{u}\|_1} \geq h(v) \geq \frac{1}{\sqrt{2}}(1 - v),$$

showing that $v \geq 1 - \frac{\beta}{\alpha}$. Thus, if $p = 1$ or 2 corresponds to whichever norm maximizes (5.13), then

$$\|\mathbf{u}\|_p = v^{-1} \|\mathbf{w}\|_p \leq (1 - \beta/\alpha)^{-1} \|\mathbf{w}\|_p$$

If we add $\|\mathbf{w}\|_p$ to both sides and apply the triangle inequality to the left-hand side, we obtain the desired result (5.6). $\square$

Remark 5.3. The discussion above splits into two cases based on the conditions $\langle \mathbf{x}_0, \mathbf{x}^* \rangle \geq (1 - \alpha^2/2)\|\mathbf{x}_0\|_2\|\mathbf{x}^*\|_2$ and $\|\mathbf{x}_0\|_2 \leq \|\mathbf{x}^*\|_2 \leq (1 + \alpha)\|\mathbf{x}_0\|_2$. This would be unnecessary if one can show that $\|\mathbf{x}^* - \mathbf{x}_0\|_1/\|\mathbf{x}^* - \mathbf{x}_0\|_2$ is bounded by some universal constant (independent of n and $\mathbf{x}^*$) times $\sqrt{s}$. Even though there is strong intuition that this is correct, a proof is currently elusive.

6 Numerical experiments

In this section, we present several numerical simulations to complement our previous theoretical investigation. All the results in this section are repeatable on a standard laptop installed with R [29]. For simplicity, we will restrict to the noiseless case. The noisy case can be carried out similarly by tuning the parameter of the penalty term arising from the constraint. Since ℓ_1/ℓ_2 minimization problems are non-convex, our algorithms solve the problem approximately. In particular, we utilize the algorithms from [43], which is essentially the Inverse Power Method [19] with an extra augmented quadratic term in $\mathbf{x}$ -update. For completeness, we briefly explain how the algorithm works: It was observed in [43] that subject to $\mathbf{A}\mathbf{x} = \mathbf{b}$, minimizing $\|\mathbf{x}\|_1/\|\mathbf{x}\|_2$ is equivalent to minimizing $\|\mathbf{x}\|_1 - \alpha\|\mathbf{x}\|_2$ for some $\alpha \in [1, \sqrt{n}]$, where α is some case-dependent parameter. Since the true value of α is unknown, one can start with an initial guess and update it using a bisection search. At iteration k , a full bisection search requires solving a minimization problem of the form $\min_{\mathbf{A}\mathbf{x}=\mathbf{b}} \|\mathbf{x}\|_1 - \alpha^{(k)}\|\mathbf{x}\|_2$ and the minimizer will be used to update $\alpha^{(k)}$. To accelerate, an adaptive algorithm based on the difference of convex functions algorithm (see [36]) was proposed in [43] by replacing $\alpha^{(k)}\|\mathbf{x}\|_2$ by its linearization at the previous iterate $\mathbf{x}^{(k-1)}$ with an additional regularization term scaled by a tunable parameter β . Given an initialization $\mathbf{x}^{(0)}$ and $\alpha^{(0)}$, the alternating algorithm can be summarized as follows:

$$\begin{cases} \mathbf{x}^{(k+1)} &= \arg \min_{\mathbf{A}\mathbf{x}=\mathbf{b}} \left\{ \|\mathbf{x}\|_1 - \frac{\alpha^{(k)}}{\|\mathbf{x}^{(k)}\|_2} \langle \mathbf{x}, \mathbf{x}^{(k)} \rangle + \frac{\beta}{2} \|\mathbf{x} - \mathbf{x}^{(k)}\|_2^2 \right\} \\ \alpha^{(k+1)} &= \frac{\|\mathbf{x}^{(k+1)}\|_1}{\|\mathbf{x}^{(k+1)}\|_2}. \end{cases}$$

The $\mathbf{x}$ -subproblem can be efficiently solved using the Alternating Direction Method of Multipliers (ADMM), see [4]. It was shown in [43] that small regularization parameters β tend to yield better

local decay rate. In our simulations we use $\beta = 0.5$. A choice for the penalty parameter ρ in the ADMM algorithm (not shown above) is slightly more tricky. In our experiment we choose $\rho = 20$, but other kinds of simulations may require tuning for a different value of ρ . We emphasize that the lack of certainty about the choice of these parameters is a drawback of the algorithms in general, and is not introduced by our implementation or choice of application. Since ADMM is a relaxation scheme, it can only approximately solve the original problem, resulting in a solution vector upon termination many of whose components have small magnitude. To increase the stability of the algorithm, a box constraint based on prior information will be incorporated, and the details will be specified later.

6.1 Initialization and support selection

A common issue in solving non-convex optimization problems is that algorithms may become trapped at local minimizers. This phenomenon is particularly worrying in our case. Indeed, due to the scale-invariant structure, the objective function ℓ_1/ℓ_2 may have infinitely many local minimizers in the feasible set. As a result, global convergence of the above algorithm depends on a good initialization. A natural choice would be the ℓ_1 minimizer or the first a few steps of the iterative reweighted least squares (IRLS) solving ℓ_0 , see [15], [10] and [24]. The intuition of these choices is that the ℓ_1 (or ℓ_q if IRLS is used) minimizer is not too far from $\mathbf{x}_0$, and $\mathbf{x}_0$ is one of the minimizers of ℓ_1/ℓ_2 . Therefore, success of the above algorithm under such initialization heavily relies on the ‘approximate’ success of the ℓ_1 minimization. This observation will be numerically verified later. To overcome the strong dependence on the ℓ_1 (ℓ_q) minimization, we will propose a novel initialization approach based on a support selection process to make the algorithm less reliant on the ℓ_1 minimizer and leads to improved results.

We propose to initialize $\mathbf{x}^{(0)}$ in a way that utilizes the information of the support of $\mathbf{x}_0$. Unfortunately, the support of $\mathbf{x}_0$ is generally unknown. As a substitute, one can use the support of the recovered solution from other algorithms. Here we will interpret the support of the ℓ_1 minimizer as near-oracle identification of the support of $\mathbf{x}_0$. Indeed the theoretical uniform recoverability of ℓ_1 minimization makes it superior to most greedy algorithms, and algorithmic implementations of ℓ_1 minimization have better convergence guarantees than many other non-convex algorithms. For a fixed sparsity s , we first compute the best s -term approximation of the ℓ_1 minimizer, which we denote by $\mathbf{x}_s$. Instead of using $\mathbf{x}_s$ directly for initialization, we will consider each element of the support separately. For every $i \in \text{supp}(\mathbf{x}_s)$, we consider the initialization $\mathbf{x}^{(0)}$ as defined by a vector whose i -th component is $\langle \mathbf{b}, \mathbf{a}_i \rangle / \|\mathbf{a}_i\|_2^2$ and the other components are 0, where $\mathbf{a}_i$ is the i -th column vector of $\mathbf{A}$. The idea behind this is to counteract the influence of incorrectly detected components in the support on the correctly detected support. One may also view this as a way to mitigate algorithm failure (see [30]) via multi-initialization. In total, one needs to solve s subproblems of ℓ_1/ℓ_2 minimization and choose the solution that gives the smallest ℓ_1/ℓ_2 value. Thus, the computational complexity will be s times more than that of solving a single ℓ_1/ℓ_2 minimization problem. This increased cost can be mitigated via parallel computing since the subproblems are embarrassingly parallel. We will call the ℓ_1/ℓ_2 algorithm with this proposed initialization “ ℓ_1/ℓ_2 +SS”, where SS stands for *support selection*. In fact, this multi-initialization approach is also applicable to other iterative algorithms.

6.2 ℓ_1/ℓ_2 simulation particulars

In the simulation below, we choose the measurement matrix $\mathbf{A}$ to be a 50×250 Gaussian random matrix with iid standard normal entries. s is the sparsity level of generated vectors $\mathbf{x}_0$. For each fixed value of s , we generate $\mathbf{x}_0$ by randomly choosing s of its components to be nonzero. The nonzero components are independently drawn from distributions with different dynamic range: Uniform($[-10, 10]$) and Uniform($[-10, 5] \cup [5, 10]$).

An additional box constraint $\|\mathbf{x}\|_\infty \leq 10$ based on this prior information on magnitude of entries is computationally imposed during iterations to solve the ℓ_1/ℓ_2 problem in both cases. Note that this extra constraint does not change the problem if $\mathbf{x}_0$ is the global minimizer. If $\mathbf{x}_0$ is not the global minimizer, the box constraint will disallow solution vectors with erroneous large magnitude, making the algorithm more stable in practice.

To detect exact recovery, we will use a slightly different criteria than the commonly used relative error threshold in other literature. We say that a computed solution $\mathbf{x}$ recovers $\mathbf{x}_0$ if the support of the best 50-term approximation of $\mathbf{x}$ under the ℓ_1 norm contains the support of $\mathbf{x}_0$. Indeed, in this case $\mathbf{x}_0$ can be easily reconstructed by solving $\mathbf{Ax} = \mathbf{b}$ with $\mathbf{A}$'s columns restricted to the support of the 50 largest components of $\mathbf{x}$. The reason for this criteria is due to computational reasons. Based on our choice of regularization and relaxation parameters, the computational error of ℓ_1/ℓ_2 cannot be made as small as many other algorithms with known convergence guarantees.

6.3 List of algorithms

Now we compare the ℓ_1/ℓ_2 +SS algorithm (l1/l2+SS) described above with the box constraint against the following popular non-convex (and ℓ_1) methods in sparse recovery:

- ℓ_1 minimization (l1): The box constraint is included for consistency in comparison. We will use the linear programming package `lpSolve` ([1]) in R to solve it.
- Reweighted ℓ_1 minimization (RWl1+l1): The box constraint is included for consistency in comparison. We will use the algorithm in Section 2.2 of [7] to solve it, with the regularization parameter $\varepsilon = 0.1$. The initialization is set as the ℓ_1 minimizer.
- $\ell_{1/2} + \ell_1$ minimization (l1/2+l1): We will use the IRLS algorithm in [15] to solve it. (We do not use the improved versions in [10] or [24] as they require knowledge on the NSP/RIP of $\mathbf{A}$, which is hard to compute in practice. For technical reasons, we are not able to incorporate the box constraint in this case.) The initialization is set as the ℓ_1 minimizer.
- $\ell_{1/2}$ +SS minimization (l1/2+SS): This is the same as the above $\ell_{1/2}$ algorithm but initialized with the additional support selection process introduced above.
- $\ell_1 - \ell_2$ minimization (l1-l2+l1): The box constraint is included for consistency in comparison. We will use the algorithm in [45] to solve it. The Lasso penalty parameter and the ADMM penalty parameter are chosen to be 0.01 and 100, respectively. The stopping rules are the same as the one proposed in [45]. The initialization is set as the ℓ_1 minimizer.
- $\ell_1/\ell_2 + \ell_1$ minimization (l1/l2+l1): The box constraint is included for consistency in comparison. We use the adaptive algorithm in [43] to solve it with the ℓ_1 initialization.
- Orthogonal Marching Pursuit (OMP): We use the OMP algorithm in [38]. The stopping criterion is either the length of the residual falls below 10^{-8} or the size of the detected support exceeds the total number of measurements, which in our case equals 50.
- Compressive Sampling Marching Pursuit (CoSaMP): We use the CoSaMP algorithm in [27]. The stopping criterion is either the length of the residual falls below 10^{-8} or the total iteration step exceeds 100. The use of the maximal iteration step in the halting rule is due to that convergence of the CoSaMP is guaranteed if the measurement matrix satisfies certain RIP conditions [34]. When sparsity gets larger, the required RIP condition is no longer valid thus the algorithm may not converge.

The sparsity s ranges from 6 to 24 by increments of 2, and for each s we perform 100 independent experiments with the average recovery rates and relative error recorded. The complete average comparison simulation is run 100 times with the quantiles plotted to demonstrate the uncertainty in the simulation. The quantile levels are chosen as 0.2-0.5-0.8 for each method. The results are given below.

To compare the scalability of the algorithms, we compute the average processing time of each algorithm applied to signals of varying length. The range of length of the signal n is chosen between 2^6 and 2^{12} , increasing by a multiple constant 2 at a time. The number of measurements

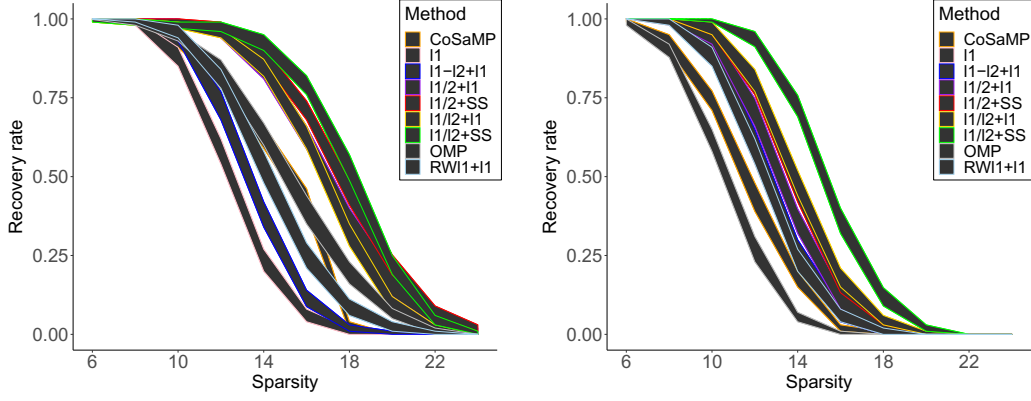


Figure 1: 0.2-0.5-0.8 quantile band for the average recovery rate: Left: Uniform($[-10, 10]$) coefficients. Right: Uniform($[-10, -5] \cup [5, 10]$) coefficients.

m is set as $m = n/4$, and the sparsity level s is set as $s = m/4$. The components in the support of the signal are uniformly generated from $[-10, 10]$. For each algorithm, its average processing time is computed as the average time of running 10 independent samples. Since the support selection procedure is algorithmically equivalent to applying a single-initialization algorithm s times, its average processing time is taken as s multiplied by the time for the same algorithm without support selection. The results are given in Figure 2. Many of the algorithms exhibit similar asymptotic computational complexity, although CoSaMP, l1-l2+l1, and l1/2+l1 have slightly better complexity. The computational cost for the support selection procedures is higher than most of the rest, but the asymptotic complexity is similar.

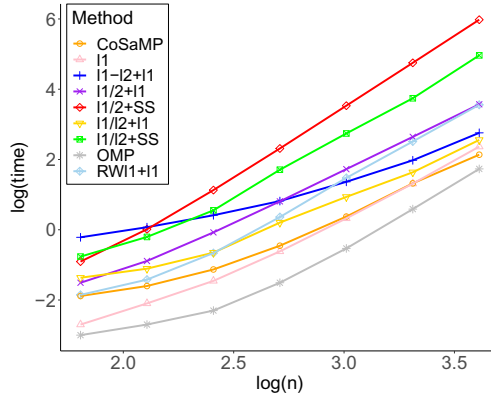


Figure 2: Average processing time of the algorithms applied to signals of different length. Both the x and y axis are plotted using the logarithm with based 10.

6.4 Simulation results

It can be seen from Figure 1 that for both types of coefficients, ℓ_1/ℓ_2 with the box constraint and SS-initialization has the best performance among all non-convex optimization methods under comparison. $\ell_{1/2}$ also performs fairly well but is slightly inferior to ℓ_1/ℓ_2 . This is no surprise since $\ell_{1/2}$ utilizes the box constraint which is absent in the ℓ_1/ℓ_2 algorithm. On the other hand, by taking the SS-initialization, the recovery rate of $\ell_{1/2}$ has significantly improved compared to a similar step taken for ℓ_q . This implies that ℓ_1/ℓ_2 is more sensitive to the initial value and the multi-initialization step enhances the success rate of the algorithm. By comparing the left- and right-hand panels in Figure 1, it is easy to observe that all the methods under comparison perform better when the dynamic range of the coefficients is large. This phenomenon can be

well explained for reweighted ℓ_1 and ℓ_q , in which a reweighting step is used to reduce the bias between ℓ_q ($0 < q \leq 1$) and ℓ_0 . However, for ℓ_1/ℓ_2 , this is not well understood. We provide some theoretical evidence in Theorem 3 for this behavior in terms of the local optimality condition; nevertheless, a complete understanding of this is still absent.

It can be seen from Figure 2 that based on our choice of algorithms, ℓ_1/ℓ_2 with single initialization demonstrates a reasonable computational time asymptotically. It is more expensive than the greedy algorithms and ℓ_1 , of which the solution is used to give a good initialization for ℓ_1/ℓ_2 . It is almost at the same level as the $\ell_1 - \ell_2$ since both used the ADMM relaxation in the computation. Meanwhile, it is cheaper than the other non-convex algorithms such as ℓ_q and reweighted ℓ_1 which either require matrix inversion or solving a linear programming problem in each iteration (more advanced numerical methods can help accelerate computation in practice, but we do not investigate it here). The support selection procedure increases the processing time of the algorithms by a multiple factor of the sparsity level. When sparsity is large, this effect is not negligible but can be mitigated via parallel computing.

It is worth pointing out that although ℓ_1/ℓ_2 algorithms yield better recovery results when the magnitude of the entries of $\mathbf{x}_0$ are known *a priori* to be bounded from above, their recovery rate is closely related to the accuracy of the ℓ_1 (ℓ_q) minimizer. If the solution $\mathbf{x}$ obtained from minimizing ℓ_1 (ℓ_q) is incoherent with $\mathbf{x}_0$, then it is unlikely that ℓ_1/ℓ_2 will give substantially better result. Our initialization approach proposed earlier is not able to completely remove such a dependence, and only mitigates the impact. However, it is likely that the support of the ℓ_1 minimizer contains at least one component that lies in the true support of $\mathbf{x}_0$. If one of these components happens to be close to the ℓ_1/ℓ_2 convergence regime of $\mathbf{x}_0$, then the support selection process will promote convergence to $\mathbf{x}_0$ by removing the influence of other elements in the detected support. Figure 3 below verifies this point.

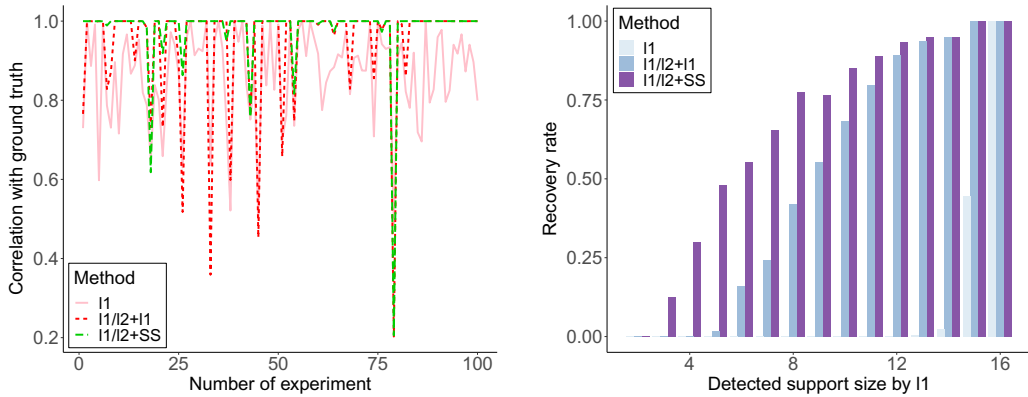


Figure 3: Numerical experiments on the support selection initialization. Left: Correlation between the minimizers of ℓ_1 , $\ell_1/\ell_2 + \ell_1$, $\ell_1/\ell_2 + \text{SS}$ and the ground truth $\mathbf{x}_0$ in 100 experiments in the case of coefficients drawn from the distribution $\text{Uniform}([-10, 10])$ with sparsity level $s = 16$. Right: Recovery rate of ℓ_1 , $\ell_1/\ell_2 + \ell_1$, and $\ell_1/\ell_2 + \text{SS}$ as a function of ℓ_1 -detected support size over 2000 experiments (total), with $\text{Uniform}([-10, 10])$ coefficients and sparsity level $s = 16$.

In Figure 3, the left panel illustrates the correlation between the minimizers of ℓ_1 , $\ell_1/\ell_2 + \ell_1$, $\ell_1/\ell_2 + \text{SS}$ and the true signal $\mathbf{x}_0$ in 100 experiments when $s = 16$ and the coefficients are chosen from $\text{Uniform}[-10, -10]$. It can be seen that the general trend of the three curves is similar. When the correlation between the ℓ_1 minimizer and $\mathbf{x}_0$ is low, say below 0.5, it is also low for both the ℓ_1/ℓ_2 minimizers. This implies that success of the ℓ_1/ℓ_2 algorithms heavily relies on the ℓ_1 minimizer being reasonably close to $\mathbf{x}_0$. On the other hand, the right panel compares the recovery rate of the three methods for different sizes of the ℓ_1 -detected support over 2000 experiments. The ℓ_1 -detected support refers to the number of indices in the support of $\mathbf{x}_0$ that are correctly detected by the best s -sparse truncation of the ℓ_1 minimizer. When the detected support is close to the true support, there is a large chance that both ℓ_1/ℓ_2 algorithms can successfully push it towards $\mathbf{x}_0$. As the detected support size diminishes, ℓ_1/ℓ_2 with the

support-selection based initialization appears to do a better job. This provides some numerical evidence that utilizing support information of the ℓ_1 minimizer through support selection helps reduce algorithm failure.

In Figure 4, we give a concrete example in which both ℓ_1 and ℓ_1/ℓ_2 initialized with ℓ_1 failed to recover $\mathbf{x}_0$ but with the additional support selection process, it succeeded. The left panel visualizes the structures of both the true and recovered signals. The support of the true signal is

$$\text{supp}(\mathbf{x}_0) = \{212, 194, 66, 73, 132, 248, 234, 70, 12, 249, 226, 102, 69, 85, 201, 106\},$$

arranged according to decreasing magnitude of the entries. In particular, $\mathbf{x}_0[212] = -9.832$ and $\mathbf{x}_0[106] = 0.158$. The ℓ_1 minimization completely fails to recover the solution in the sense that $\text{supp}(\mathbf{x}_0) \not\subseteq \text{supp}(\mathbf{x}_{\ell_1})$, where $\mathbf{x}_{\ell_1}$ is the ℓ_1 minimizer. In this case, the ℓ_1/ℓ_2 with ℓ_1 initialization moves to a local minimizer near $\mathbf{x}_{\ell_1}$ which is different from $\mathbf{x}_0$, whereas the same algorithm with support selection initialization successfully detects $\mathbf{x}_0$. The right panel gives a more careful comparison between the true components and the recovered components using $\ell_1/\ell_2 + \text{SS}$ on the support.

To better understand the success of the support selection procedure, we compare the ℓ_1/ℓ_2 objective value of the s minimizers obtained from initializing on each component in the support of the best s -approximation of $\mathbf{x}_{\ell_1}$. This is given in Table 1. In our example, the support of the best s -approximation of $\mathbf{x}_{\ell_1}$ is

$$\text{supp}(\mathbf{x}_{\ell_1}|_s) = \{234, 137, 212, 66, 145, 85, 87, 102, 194, 132, 99, 110, 205, 40, 246, 128\}.$$

The detected support by the ℓ_1 is

$$\text{supp}(\mathbf{x}_0) \cap \text{supp}(\mathbf{x}_{\ell_1}|_s) = \{212, 194, 66, 132, 234, 102, 85\}.$$

The ℓ_1/ℓ_2 objective values of $\mathbf{x}_0$ and the solutions found by the ℓ_1 and $\ell_1/\ell_2 + \ell_1$ are 3.456, 5.173 and 4.510, respectively. It is clear from Table 1 that support selection procedure initialized at indices 212, 194, 66 and 132 succeeded in recovering $\mathbf{x}_0$ (up to some computational error). These indices are in the support of $\mathbf{x}_0$. Note that initialization at other indices which are also in $\text{supp}(\mathbf{x}_0)$ such as 234, 102 and 85 results algorithm failure. A possible explanation for this is that the magnitude of $\mathbf{x}_0$ on these indices is relatively small compared to that on the indices leading to success.

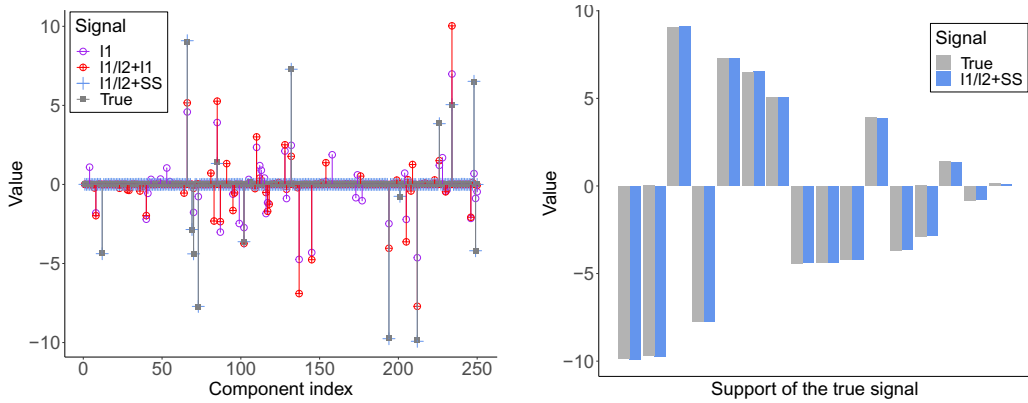


Figure 4: A case where both ℓ_1 and $\ell_1/\ell_2 + \ell_1$ failed to recover $\mathbf{x}_0$ but $\ell_1/\ell_2 + \text{SS}$ succeeded in the case of coefficients $\text{Uniform}([-10, 10])$ and sparsity level $s = 16$: Left: General distribution of magnitude of the true components and recovered components using ℓ_1 , $\ell_1/\ell_2 + \ell_1$ and $\ell_1/\ell_2 + \text{SS}$. Right: Careful comparison between the true components and the recovered components using $\ell_1/\ell_2 + \text{SS}$ on the support.

Initialized index in $\text{supp}(\mathbf{x}_{\ell_1 s})$	ℓ_1/ℓ_2
234	4.235
137	4.450
212	3.483
66	3.483
145	4.401
85	4.277
87	4.282
102	4.248
194	3.483
132	3.483
99	4.520
110	4.361
205	4.493
40	4.505
246	4.520
128	4.480

Table 1: Comparison of the ℓ_1/ℓ_2 objective value from the 16 single-support initializations in the Figure 4. Initializations started from support indices 212, 66, 194 and 132 lead to solutions with the best ℓ_1/ℓ_2 value 3.483, which approximately matches the ℓ_1/ℓ_2 value (3.456) of $\mathbf{x}_0$.

7 Conclusion and future work

We have theoretically and numerically investigated the ℓ_1/ℓ_2 minimization problem in the context of recovery of sparse signals from a small number of measurements. We have provided a novel local optimality criterion in Theorem 3.1, which gives some theoretical justification to the empirical observation that ℓ_1/ℓ_2 performs better when the nonzero entries of the sparse solution have a large dynamic range. We also provide a uniform recoverability condition in Theorem 4.1 for the ℓ_1/ℓ_2 minimization problem. Our final theoretical contribution is a robustness result in Theorem 5.1 that can be used to provide stability for noisy ℓ_1/ℓ_2 minimization problems, see Corollary 5.2. We have also proposed a new type of initialization for this nonconvex optimization problem called *support selection* that empirically improves the recovery rate for ℓ_1/ℓ_2 minimization. Investigations that give a better theoretical understanding of why large dynamic range improves this type of minimization, along with additional analysis to better quantify stability in noisy cases, will be the subject of future research.

Although our analysis in this article arrives in similar recoverability and stability conditions analogous to the ones given by ℓ_1 , it does not give anything better. This may be due to the fact that the inequalities originally sharp for ℓ_1 become less optimal when additional division steps are taken in the estimates. Also, norm ratios are more of objectives promoting the compressibility of a signal rather than the sparsity defined by ℓ_0 , which is highly discontinuous. Whereas in many practical problems, a sparse signal comes from the approximation of a compressible signal. This suggests that ℓ_1/ℓ_2 itself could be an alternative objective in terms of defining the goal of compressed sensing. More theoretical work in this direction is also worth exploring in the future.

Conflict of interest

There is none.

Acknowledgements

We would like to thank the anonymous referees for their very helpful comments which significantly improve the presentation of the paper. The first author (yxu@utah.math.edu) thanks for useful discussions with Tom Alberts, You-Cheng Chou and Dong Wang. The first and second

authors (yxu@math.utah.edu, akil@sci.utah.edu) acknowledge partial support from NSF DMS-1848508. The third author (tranha@ornl.gov) acknowledges support from Scientific Discovery through Advanced Computing (SciDAC) program through the FASTMath Institute under Contract No. DE-AC02-05CH11231. The last author (cwebst13@utk.edu) acknowledges the U.S. Department of Energy, Office of Science, Early Career Research Program under award number ERKJ314; U.S. Department of Energy, Office of Advanced Scientific Computing Research under award numbers ERKJ331 and ERKJ345; and the National Science Foundation, Division of Mathematical Sciences, Computational Mathematics program under contract number DMS1620280.

References

- [1] M. Berkelaar et al. *lpSolve: Interface to lpSolve v. 5.5 to Solve Linear/Integer Programs*. R package version 5.6.13.3. 2019.
- [2] T. Blumensath and M. E. Davies. “Iterative hard thresholding for compressed sensing”. *Applied and computational harmonic analysis* 27.3 (2009), pp. 265–274.
- [3] R. I. Boş, M. N. Dao, and G. Li. “Extrapolated Proximal Subgradient Algorithms for Nonconvex and Nonsmooth Fractional Programs”. *arXiv preprint arXiv:2003.04124* (2020).
- [4] S. Boyd. *Distributed Optimization and Statistical Learning via the Alternating Direction Method of Multipliers*. now Publishers Inc, 2010. DOI: [10.1561/9781601984616](https://doi.org/10.1561/9781601984616).
- [5] E. Candès and B. Recht. “Exact matrix completion via convex optimization”. *Communications of the ACM* 55.6 (2012), p. 111. DOI: [10.1145/2184319.2184343](https://doi.org/10.1145/2184319.2184343).
- [6] E. J. Candès and T. Tao. “Near-Optimal Signal Recovery From Random Projections: Universal Encoding Strategies?” *IEEE Transactions on Information Theory* 52.12 (2006), pp. 5406–5425. DOI: [10.1109/tit.2006.885507](https://doi.org/10.1109/tit.2006.885507).
- [7] E. J. Candès, M. B. Wakin, and S. P. Boyd. *Enhancing Sparsity by Reweighted ℓ_1 Minimization*. Tech. rep. 2008. DOI: [10.21236/ada528514](https://doi.org/10.21236/ada528514).
- [8] R. Chartrand. “Exact Reconstruction of Sparse Signals via Nonconvex Minimization”. *IEEE Signal Processing Letters* 14.10 (2007), pp. 707–710. DOI: [10.1109/lsp.2007.898300](https://doi.org/10.1109/lsp.2007.898300).
- [9] A. Cohen, W. Dahmen, and R. DeVore. “Compressed sensing and best k -term approximation”. *Journal of the American Mathematical Society* 22.1 (2008), pp. 211–231. DOI: [10.1090/s0894-0347-08-00610-3](https://doi.org/10.1090/s0894-0347-08-00610-3).
- [10] I. Daubechies, R. DeVore, M. Fornasier, and C. S. Gunturk. *Iteratively Re-weighted Least Squares Minimization for Sparse Recovery*. Tech. rep. 2008. DOI: [10.21236/ada528510](https://doi.org/10.21236/ada528510).
- [11] D. L. Donoho and M. Elad. “Optimally sparse representation in general (nonorthogonal) dictionaries via ℓ_1 minimization”. *Proceedings of the National Academy of Sciences* 100.5 (2003), pp. 2197–2202. DOI: [10.1073/pnas.0437847100](https://doi.org/10.1073/pnas.0437847100).
- [12] D. Donoho. “Compressed sensing”. *IEEE Transactions on Information Theory* 52.4 (2006), pp. 1289–1306. DOI: [10.1109/tit.2006.871582](https://doi.org/10.1109/tit.2006.871582).
- [13] D. Donoho and X. Huo. “Uncertainty principles and ideal atomic decomposition”. *IEEE Transactions on Information Theory* 47.7 (2001), pp. 2845–2862. DOI: [10.1109/18.959265](https://doi.org/10.1109/18.959265).
- [14] E. Esser, Y. Lou, and J. Xin. “A Method for Finding Structured Sparse Solutions to Nonnegative Least Squares Problems with Applications”. *SIAM Journal on Imaging Sciences* 6.4 (2013), pp. 2010–2046. DOI: [10.1137/13090540x](https://doi.org/10.1137/13090540x).
- [15] S. Foucart and M.-J. Lai. “Sparsest solutions of underdetermined linear systems via ℓ^q -minimization for $0 < q < 1$ ”. *Applied and Computational Harmonic Analysis* 26.3 (2009), pp. 395–407.
- [16] S. Foucart, A. Pajor, H. Rauhut, and T. Ullrich. “The Gelfand widths of ℓ^p -balls for $0 < p \leq 1$ ”. *Journal of Complexity* 26.6 (2010), pp. 629–640.
- [17] E. D. Gluskin. “Norms of random matrices and widths of finite-dimensional sets”. *Mathematics of the USSR-Sbornik* 48.1 (1984), pp. 173–182. DOI: [10.1070/sm1984v048n01abeh002667](https://doi.org/10.1070/sm1984v048n01abeh002667).
- [18] R. Gribonval and M. Nielsen. “Highly sparse representations from dictionaries are unique and independent of the sparseness measure”. *Applied and Computational Harmonic Analysis* 22.3 (2007), pp. 335–355. DOI: [10.1016/j.acha.2006.09.003](https://doi.org/10.1016/j.acha.2006.09.003).

- [19] M. Hein and T. Bühler. “An inverse power method for nonlinear eigenproblems with applications in 1-spectral clustering and sparse PCA”. In: *Advances in Neural Information Processing Systems*. 2010, pp. 847–855.
- [20] P. Hoyer. “Non-negative sparse coding”. In: *Proceedings of the 12th IEEE Workshop on Neural Networks for Signal Processing*. IEEE. DOI: [10.1109/nnspp.2002.1030067](#).
- [21] N. Hurley and S. Rickard. “Comparing measures of sparsity”. In: *2008 IEEE Workshop on Machine Learning for Signal Processing*. IEEE, 2008. DOI: [10.1109/mlsp.2008.4685455](#).
- [22] B. Kasin. “The widths of certain finite-dimensional sets and classes of smooth functions”. *Izv. Akad. Nauk SSSR Ser. Mat* 41.2 (1977), pp. 334–351.
- [23] D. Krishnan, T. Tay, and R. Fergus. “Blind deconvolution using a normalized sparsity measure”. In: *CVPR 2011*. IEEE, 2011. DOI: [10.1109/cvpr.2011.5995521](#).
- [24] M.-J. Lai, Y. Xu, and W. Yin. “Improved Iteratively Reweighted Least Squares for Unconstrained Smoothed ℓ_q Minimization”. *SIAM Journal on Numerical Analysis* 51.2 (2013), pp. 927–957. DOI: [10.1137/110840364](#).
- [25] J. Lv and Y. Fan. “A unified approach to model selection and sparse recovery using regularized least squares”. *The Annals of Statistics* 37.6A (2009), pp. 3498–3528. DOI: [10.1214/09-aos683](#).
- [26] B. K. Natarajan. “Sparse Approximate Solutions to Linear Systems”. *SIAM Journal on Computing* 24.2 (1995), pp. 227–234. DOI: [10.1137/s0097539792240406](#).
- [27] D. Needell and J. A. Tropp. “CoSaMP”. *Communications of the ACM* 53.12 (2010), p. 93. DOI: [10.1145/1859204.1859229](#).
- [28] A. Petrosyan, H. Tran, and C. Webster. “Reconstruction of jointly sparse vectors via manifold optimization”. *Applied Numerical Mathematics* 144 (2019), pp. 140–150. DOI: [10.1016/j.apnum.2019.05.022](#).
- [29] R Core Team. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing. Vienna, Austria, 2020.
- [30] Y. Rahimi, C. Wang, H. Dong, and Y. Lou. “A Scale-Invariant Approach for Sparse Signal Recovery”. *SIAM Journal on Scientific Computing* 41.6 (2019), A3649–A3672. DOI: [10.1137/18m123147x](#).
- [31] P. Rigollet and J.-C. Hütter. “High dimensional statistics”. *Lecture notes for course 18S997* (2015).
- [32] M. Rudelson and R. Vershynin. “On sparse reconstruction from Fourier and Gaussian measurements”. *Communications on Pure and Applied Mathematics* 61.8 (2008), pp. 1025–1045. DOI: [10.1002/cpa.20227](#).
- [33] M. Rudelson and R. Vershynin. “Smallest singular value of a random rectangular matrix”. *Communications on Pure and Applied Mathematics* 62.12 (2009), pp. 1707–1739. DOI: [10.1002/cpa.20294](#).
- [34] S. Satpathi and M. Chakraborty. “On the number of iterations for convergence of CoSaMP and Subspace Pursuit algorithms”. *Applied and Computational Harmonic Analysis* 43.3 (2017), pp. 568–576. DOI: [10.1016/j.acha.2016.10.001](#).
- [35] M. TAO and Y. LOU. “MINIMIZATION OF L1 OVER L2 FOR SPARSE SIGNAL RECOVERY WITH CONVERGENCE GUARANTEE” ().
- [36] P. D. Tao and L. T. H. An. “A D.C. Optimization Algorithm for Solving the Trust-Region Subproblem”. *SIAM Journal on Optimization* 8.2 (1998), pp. 476–505. DOI: [10.1137/s1052623494274313](#).
- [37] H. Tran and C. Webster. “A class of null space conditions for sparse recovery via nonconvex, non-separable minimizations”. *Results in Applied Mathematics* 3 (2019), p. 100011. DOI: [10.1016/j.rinam.2019.100011](#).
- [38] J. A. Tropp and A. C. Gilbert. “Signal Recovery From Random Measurements Via Orthogonal Matching Pursuit”. *IEEE Transactions on Information Theory* 53.12 (2007), pp. 4655–4666. DOI: [10.1109/tit.2007.909108](#).
- [39] R. Vershynin. “Estimation in High Dimensions: A Geometric Perspective”. In: *Sampling Theory, a Renaissance*. Springer International Publishing, 2015, pp. 3–66. DOI: [10.1007/978-3-319-19749-4_1](#).
- [40] R. Vershynin. *High-Dimensional Probability*. Cambridge University Press, 2018. DOI: [10.1017/9781108231596](#).

- [41] C. Wang, M. Tao, C.-N. Chuah, J. Nagy, and Y. Lou. “Minimizing L1 over L2 norms on the gradient”. *arXiv preprint arXiv:2101.00809* (2021).
- [42] C. Wang, M. Tao, J. Nagy, and Y. Lou. “Limited-angle CT reconstruction via the L1/L2 minimization”. *arXiv preprint arXiv:2006.00601* (2020).
- [43] C. Wang, M. Yan, Y. Rahimi, and Y. Lou. “Accelerated Schemes for the ℓ_1/ℓ_2 Minimization”. *IEEE Transactions on Signal Processing* 68 (2020), pp. 2660–2669. DOI: [10.1109/tsp.2020.2985298](https://doi.org/10.1109/tsp.2020.2985298).
- [44] P. Yin, E. Esser, and J. Xin. “Ratio and difference of ℓ_1 and ℓ_2 norms and sparse representation with coherent dictionaries”. *Communications in Information and Systems* 14.2 (2014), pp. 87–109. DOI: [10.4310/cis.2014.v14.n2.a2](https://doi.org/10.4310/cis.2014.v14.n2.a2).
- [45] P. Yin, Y. Lou, Q. He, and J. Xin. “Minimization of ℓ_1 – ℓ_2 for Compressed Sensing”. *SIAM Journal on Scientific Computing* 37.1 (2015), A536–A563. DOI: [10.1137/140952363](https://doi.org/10.1137/140952363).
- [46] L. Zeng, P. Yu, and T. K. Pong. “Analysis and algorithms for some compressed sensing models based on L1/L2 minimization”. *arXiv preprint arXiv:2007.12821* (2020).
- [47] Y. Zhang. “Theory of Compressive Sensing via ℓ_1 -Minimization: a Non-RIP Analysis and Extensions”. *Journal of the Operations Research Society of China* 1.1 (2013), pp. 79–105. DOI: [10.1007/s40305-013-0010-2](https://doi.org/10.1007/s40305-013-0010-2).