

Iteratively Re-weighted Least Squares Minimization for Sparse Recovery

Ingrid Daubechies*, Ronald DeVore†, Massimo Fornasier‡, C. Sinan Güntürk§

June 14, 2008

Abstract

Under certain conditions (known as the *Restricted Isometry Property* or RIP) on the $m \times N$ -matrix Φ (where $m < N$), vectors $x \in \mathbb{R}^N$ that are sparse (i.e. have most of their entries equal to zero) can be recovered exactly from $y := \Phi x$ even though $\Phi^{-1}(y)$ is typically an $(N - m)$ -dimensional hyperplane; in addition x is then equal to the element in $\Phi^{-1}(y)$ of minimal ℓ_1 -norm. This minimal element can be identified via linear programming algorithms.

We study an alternative method of determining x , as the limit of an *Iteratively Re-weighted Least Squares* (IRLS) algorithm. The main step of this IRLS finds, for a given weight vector w , the element in $\Phi^{-1}(y)$ with smallest $\ell_2(w)$ -norm. If $x^{(n)}$ is the solution at iteration step n , then the new weight $w^{(n)}$ is defined by $w_i^{(n)} := \left[|x_i^{(n)}|^2 + \epsilon_n^2 \right]^{-1/2}$, $i = 1, \dots, N$, for a decreasing sequence of adaptively defined ϵ_n ; this updated weight is then used to obtain $x^{(n+1)}$ and the process is repeated. We prove that when Φ satisfies the RIP conditions, the sequence $x^{(n)}$ converges for all y , regardless of whether $\Phi^{-1}(y)$ contains a sparse vector. If there is a sparse vector in $\Phi^{-1}(y)$, then the limit is this sparse vector, and when $x^{(n)}$ is sufficiently close to the limit, the remaining steps of the algorithm converge exponentially fast (*linear convergence* in the terminology of numerical optimization). The same algorithm with the “heavier” weight $w_i^{(n)} = \left[|x_i^{(n)}|^2 + \epsilon_n^2 \right]^{-1+\tau/2}$, $i = 1, \dots, N$, where $0 < \tau < 1$, can recover sparse solutions as well; more importantly, we show its local convergence is superlinear and approaches a *quadratic* rate for τ approaching to zero.

1 Introduction

Let Φ be an $m \times N$ matrix with $m < N$ and let $y \in \mathbb{R}^m$. (In the compressed sensing application that motivated this study, Φ typically has full rank, i.e. $\text{Ran}(\Phi) = \mathbb{R}^m$. We shall implicitly assume,

*Princeton University, Department of Mathematics and Program in Applied and Computational Mathematics, ingrid@math.princeton.edu.

†University of South Carolina, Department of Mathematics, devore@math.sc.edu.

‡Johann Radon Institute for Computational and Applied Mathematics, Austrian Academy of Sciences, massimo.fornasier@oeaw.ac.at.

§New York University, Courant Institute of Mathematical Sciences, gunturk@courant.nyu.edu.

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 14 JUN 2008		2. REPORT TYPE		3. DATES COVERED 00-00-2008 to 00-00-2008	
4. TITLE AND SUBTITLE Iteratively Re-weighted Least Squares Minimization for Sparse Recovery				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Princeton University, Department of Mathematics and Program in Applied and Computational Mathematics, Princeton, NJ, 08544				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT Under certain conditions (known as the Restricted Isometry Property or RIP) on the $m \times N$ matrix (where $m < N$), vectors $x \in \mathbb{R}^N$ that are sparse (i.e. have most of their entries equal to zero) can be recovered exactly from $y := Ax$ even though y is typically an $(N - m)$-dimensional hyperplane; in addition x is then equal to the element in y of minimal ℓ_1-norm. This minimal element can be identified via linear programming algorithms. We study an alternative method of determining x, as the limit of an Iteratively Re-weighted Least Squares (IRLS) algorithm. The main step of this IRLS finds, for a given weight vector w, the element in y with smallest $\ell_2(w)$-norm. If $x(n)$ is the solution at iteration step n then the new weight $w(n)$ is defined by $w(n)_i := \frac{1}{ x(n)_i + \epsilon}$, $i = 1, \dots, N$, for a decreasing sequence of adaptively defined ϵ; this updated weight is then used to obtain $x(n+1)$ and the process is repeated. We prove that when A satisfies the RIP conditions, the sequence $x(n)$ converges for all y, regardless of whether y contains a sparse vector. If there is a sparse vector in y, then the limit is this sparse vector, and when $x(n)$ is sufficiently close to the limit, the remaining steps of the algorithm converge exponentially fast (linear convergence in the terminology of numerical optimization). The same algorithm with the "heavier" weight $w(n)_i := \frac{1}{ x(n)_i + \epsilon} + \frac{1}{2}$, $i = 1, \dots, N$, where $0 < \epsilon < 1$, can recover sparse solutions as well more importantly, we show its local convergence is superlinear and approaches a quadratic rate for ϵ approaching to zero.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 35	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

throughout the paper, that this is the case. Our results still hold for the case where $\text{Ran}(\Phi) \subsetneq \mathbb{R}^m$, with the proviso that y must then lie in $\text{Ran}(\Phi)$.)

The linear system of equations

$$\Phi x = y \tag{1.1}$$

is underdetermined, and has infinitely many solutions. If $\mathcal{N} := \mathcal{N}(\Phi)$ is the null space of Φ and x_0 is any solution to (1.1) then the set $\mathcal{F}(y) := \Phi^{-1}(y)$ of all solutions to (1.1) is given by $\mathcal{F}(y) = x_0 + \mathcal{N}$.

In the absence of any other information, no solution to (1.1) is to be preferred over any other. However, many scientific applications work under the assumption that the desired solution $x \in \mathcal{F}(y)$ is either sparse or well approximated by (a) sparse vector(s). Here and later, we say a vector *has sparsity* k (or is *k-sparse*) if it has *at most* k nonzero coordinates. Suppose then that we know that the desired solution of (1.1) is k -sparse, where $k < m$ is known. How could we find such an x ? One possibility is to consider any set T of k column indices and find the least squares solution $x^T := \operatorname{argmin}_{z \in \mathcal{F}(y)} \|\Phi_T z - y\|_{\ell_2^m}$, where Φ_T is obtained from Φ by setting to zero all entries that are not in columns from T . Finding x^T is numerically simple (see (1.9)). After finding each x^T , we choose the particular set T^* that minimizes the residual $\|\Phi_T z - y\|_{\ell_2^m}$. This would find a k -sparse solution (if it exists), $x^* = x^{T^*}$. However, this naive method is numerically prohibitive when N and k are large, since it requires solving $\binom{N}{k}$ least squares problems.

An attractive alternative to the naive minimization is its convex relaxation that consists in selecting the element in $\mathcal{F}(y)$ which has minimal ℓ_1 -norm:

$$x := \operatorname{argmin}_{z \in \mathcal{F}(y)} \|z\|_{\ell_1^N}. \tag{1.2}$$

Here and later we use the ℓ_p -norms

$$\|x\|_{\ell_p} := \|x\|_{\ell_p^N} := \begin{cases} \left(\sum_{j=1}^N |x_j|^p \right)^{1/p}, & 0 < p < \infty, \\ \max_{j=1, \dots, N} |x_j|, & p = \infty. \end{cases} \tag{1.3}$$

Under certain assumptions on Φ and y that we shall describe in §2, it is known that (1.2) has a unique solution (which we shall denote by x^*), and that, when there is a k -sparse solution to (1.1), (1.2) will find this solution [3, 7, 20, 21]. Because the problem (1.2) can be formulated as a linear program, it is numerically tractable.

Solving underdetermined systems by ℓ_1 -minimization has a long history. It is at the heart of many numerical algorithms for approximation, compression, and statistical estimation. The use of the ℓ_1 -norm as a sparsity-promoting functional can be found first in reflection seismology and in deconvolution of seismic traces [16, 37, 38]. Rigorous results for ℓ_1 -minimization began to appear in the late-1980's, with Donoho and Stark [23] and Donoho and Logan [22]. Applications for ℓ_1 -minimization in statistical estimation began in the mid-1990's with the introduction of the LASSO and related formulations [39] (iterative soft-thresholding), also known as Basis Pursuit [15], proposed in compression applications for extracting the sparsest signal representation from highly overcomplete frames. Around the same time other signal processing groups started using ℓ_1 -minimization for the analysis of sparse signals; see, e.g. [32]. The applications and understanding

of ℓ_1 -minimization saw a dramatic increase in the last 5 years [20, 24, 21, 25, 7, 4, 3, 6], with the development of fairly general mathematical frameworks in which ℓ_1 -minimization, known heuristically to be sparsity-promoting, can be *proved* to recover sparse solutions *exactly*. We shall not trace all the relevant results and applications; a detailed history is beyond the scope of this introduction. We refer the reader to the survey papers [5, 1]. The reader can also find a comprehensive collection of the ongoing recent developments at the web-site <http://www.dsp.ece.rice.edu/cs/>. In fact, ℓ_1 -minimization has been so surprisingly effective in several applications, that Candès, Wakin, and Boyd call it the “modern least squares” in [8]. We thus clearly need efficient algorithms for the minimization problem (1.2).

Several alternatives to (1.2), see, e.g., [26, 31], have been proposed as possibly more efficient numerically, or simpler to implement by non-experts, than standard algorithms for linear programming (such as interior point or barrier methods). In this paper we clarify fine convergence properties of one such alternative method, called *Iteratively Re-weighted Least Squares minimization* (IRLS). It begins with the following observation (see §2 for details). If (1.2) has a solution x^* that has no vanishing coordinates, then the (unique!) solution x^w of the weighted least squares problem

$$x^w := \operatorname{argmin}_{z \in \mathcal{F}(y)} \|z\|_{\ell_2^N(w)}, \quad w := (w_1, \dots, w_N), \quad \text{where } w_j := |x_j^*|^{-1}, \quad (1.4)$$

coincides with x^* . (The following argument provides a short proof by contradiction of this statement. Assume that x^* is not the $\ell_2^N(w)$ -minimizer. Then there exists $\eta \in \mathcal{N}$ such that $\|x^* + \eta\|_{\ell_2^N(w)}^2 < \|x^*\|_{\ell_2^N(w)}^2$ or equivalently $\frac{1}{2}\|\eta\|_{\ell_2^N(w)}^2 < -\sum_{j=1}^N w_j \eta_j x_j^* = \sum_{j=1}^N \eta_j \operatorname{sign}(x_j^*)$. However, because x^* is an ℓ_1 -minimizer, we have $\|x^*\|_{\ell_1} \leq \|x^* + h\eta\|_{\ell_1}$ for all $h \neq 0$; taking h sufficiently small, this implies $\sum_{j=1}^N \eta_j \operatorname{sign}(x_j^*) = 0$, a contradiction.)

Since we do not know x^* , this observation cannot be used directly. However, it leads to the following paradigm for finding x^* . We choose a starting weight w^0 and solve (1.4) for this weight. We then use this solution to define a new weight w^1 and repeat this process. An IRLS algorithm of this type appears for the first time in the approximation practice in the Ph.D. thesis of Lawson in 1961 [30], in the form of an algorithm for solving uniform approximation problems, in particular by Chebyshev polynomials, by means of limits of weighted ℓ_p -norm solutions. This iterative algorithm is now well-known in classical approximation theory as Lawson’s algorithm. In [17] it is proved that this algorithm has in principle a linear convergence rate. In the 1970s extensions of Lawson’s algorithm for ℓ_p -minimization, and in particular ℓ_1 -minimization, were proposed. In signal analysis, IRLS was proposed as a technique to build algorithms for sparse signal reconstruction in [28]. Perhaps the most comprehensive mathematical analysis of the performance of IRLS for ℓ_p -minimization was given in the work of Osborne [33].

Osborne proves that a suitable IRLS method is convergent for $1 < p < 3$. For $p = 1$, if w^n denotes the weight at the n th iteration and x^n the minimal weighted least squares solution for this weight, then the algorithm considered by Osborne defines the new weight w^{n+1} coordinatewise as $w_j^{n+1} := |x_j^n|^{-1}$. His main conclusion in this case is that if the ℓ_1 minimization problem (1.2) has a unique solution, then the algorithm converges to this solution, in principle with linear convergence

rate, i.e. exponentially fast, with a constant “contraction factor”.

However, the analysis of Osborne does not take into consideration what happens if one of the coordinates vanishes at some iteration n , i.e. $x_j^n = 0$. Taking this to impose that the corresponding weight component w_j^{n+1} must “equal” ∞ leads to $x_j^{n+1} = 0$ at the next iteration as well; this then persists in all later iterations. If $x_j^* = 0$, all is well, but if there is an index j for which $x_j^* \neq 0$, yet $x_j^n = 0$ at some iteration step n , then this “infinite weight” prescription leads to problems. In practice, this is avoided by changing the definition of the weight at coordinates j where $x_j^n = 0$ (see [31] and [10, 27] where a variant for total variation minimization is studied); such modified algorithms need no longer converge to x^* , however). Because Osborne’s convergence proof is local, it implies that if the iterations begin with a vector sufficiently close to the solution, and if the solution is unique and has only nonzero entries, then none of the $x_j^n = 0$ vanish, and the weight-change is not required; Osborne’s analysis does indeed show the linear convergence rate of the algorithm under these assumptions. Unfortunately, as we will see in Remark 2.2, the uniqueness of the solution necessarily implies that it has vanishing components. In other words, the set of vectors to which Osborne’s analysis applies is vacuous.

The purpose of the present paper is to put forward an IRLS algorithm that gives a re-weighting without infinite components in the weight, and to provide an analysis of this algorithm, with various results about its convergence and rate of convergence. It turns out that care must be taken in just how the new weight w^{n+1} is derived from the solution x^n of the current weighted least squares problem. To manage this difficulty, we shall consider a very specific recipe for generating the weights. Other recipes are certainly possible.

Given a real number $\epsilon > 0$ and a weight vector $w \in \mathbb{R}^N$, with $w_j > 0$, $j = 1, \dots, N$, we define

$$\mathcal{J}(z, w, \epsilon) := \frac{1}{2} \left[\sum_{j=1}^N z_j^2 w_j + \sum_{j=1}^N (\epsilon^2 w_j + w_j^{-1}) \right], \quad z \in \mathbb{R}^N. \quad (1.5)$$

Given w and ϵ , the element $z \in \mathbb{R}^N$ that minimizes $\mathcal{J}$ is unique because $\mathcal{J}$ is strictly convex.

Our algorithm will use an alternating method for choosing minimizers and weights based on the functional $\mathcal{J}$. To describe this, we define for $z \in \mathbb{R}^N$ the non-increasing rearrangement $r(z)$ of the absolute values of the entries of z . Thus $r(z)_i$ is the i -th largest element of the set $\{|z_j|, j = 1, \dots, N\}$, and a vector v is k -sparse iff $r(v)_{k+1} = 0$.

Algorithm 1 We initialize by taking $w^0 := (1, \dots, 1)$. We also set $\epsilon_0 := 1$. We then recursively define for $n = 0, 1, \dots$,

$$x^{n+1} := \operatorname{argmin}_{z \in \mathcal{F}(y)} \mathcal{J}(z, w^n, \epsilon_n) = \operatorname{argmin}_{z \in \mathcal{F}(y)} \|z\|_{\ell_2(w^n)} \quad (1.6)$$

and

$$\epsilon_{n+1} := \min(\epsilon_n, \frac{r(x^{n+1})_{K+1}}{N}), \quad (1.7)$$

where K is a fixed integer that will be described more fully later. We also define

$$w^{n+1} := \operatorname{argmin}_{w>0} \mathcal{J}(x^{n+1}, w, \epsilon_{n+1}). \quad (1.8)$$

We stop the algorithm if $\epsilon_n = 0$; in this case we define $x^j := x^n$ for $j > n$. However, in general, the algorithm will generate an infinite sequence $(x^n)_{n \in \mathbb{N}}$ of distinct vectors. $\square$

Each step of the algorithm requires the solution of a least squares problem. In matrix form

$$x^{n+1} = D_n \Phi^t (\Phi D_n \Phi^t)^{-1} y, \quad (1.9)$$

where D_n is the $N \times N$ diagonal matrix whose j -th diagonal entry is w_j^n and A^t denotes the transpose of the matrix A . Once x^{n+1} is found, the weight w^{n+1} is given by

$$w_j^{n+1} = [(x_j^{n+1})^2 + \epsilon_{n+1}^2]^{-1/2}, \quad j = 1, \dots, N. \quad (1.10)$$

We shall prove several results about the convergence and rate of convergence of this algorithm. This will be done under the following assumption on Φ .

The Restricted Isometry Property (RIP): We say that the matrix Φ satisfies the Restricted Isometry Property of order L with constant $\delta \in (0, 1)$ if for each vector z with sparsity L we have

$$(1 - \delta) \|z\|_{\ell_2^N} \leq \|\Phi z\|_{\ell_2^m} \leq (1 + \delta) \|z\|_{\ell_2^N}. \quad (1.11)$$

The RIP was introduced by Candès and Tao [7, 4] in their study of compressed sensing and ℓ_1 -minimization. It has several analytical and geometrical interpretations that will be discussed in §3. To mention just one of these results (see [18]), it is known that if Φ has the RIP of order $L := J + J'$, with $\delta < \frac{\sqrt{J'} - \sqrt{J}}{\sqrt{J'} + \sqrt{J}}$ (here $J' > J$) and if (1.1) has a J -sparse solution $z \in \mathcal{F}(y)$, then this solution is the unique ℓ_1 minimizer in $\mathcal{F}(y)$. (This can still be sharpened: in [9], Candès showed that if $\mathcal{F}(y)$ contains a J -sparse vector, and if Φ has RIP of order $2J$ with $\delta < \sqrt{2} - 1$, then that J -sparse vector is unique and is the unique ℓ_1 minimizer in $\mathcal{F}(y)$.)

The main result of this paper (Theorem 5.3) is that whenever Φ satisfies the RIP of order $K + K'$ (for some $K' > K$) and δ sufficiently close to zero, then Algorithm 1 converges to a solution $\bar{x}$ of (1.1) for each $y \in \mathbb{R}^m$. Moreover, if there is a solution z to (1.1) that has sparsity $k \leq K - \kappa$, then $\bar{x} = z$. Here $\kappa > 1$ depends on the RIP constant δ and can be made arbitrarily close to 1 when δ is made small. The result cited in our previous paragraph implies that in this case $\bar{x} = x^*$, where x^* is the ℓ_1 -minimal solution to (1.1).

A second part of our analysis concerns rates of convergence. We shall show that if (1.1) has a k -sparse solution with, e.g., $k \leq K - 4$ and if Φ satisfies the RIP of order $3K$ with δ sufficiently close to zero, then Algorithm 1 converges exponentially fast to $\bar{x} = x^*$. Namely, once x^{n_0} is sufficiently close to its limit $\bar{x}$, we have

$$\|\bar{x} - x^{n+1}\|_{\ell_1^N} \leq \mu \|\bar{x} - x^n\|_{\ell_1^N}, \quad n \geq n_0, \quad (1.12)$$

where $\mu < 1$ is a fixed constant (depending on δ). From this result it follows that we have exponential convergence to $\bar{x}$ whenever $\bar{x}$ is k -sparse; however we have no real information on how long it will take before the iterates enter the region where we can control μ . (Note that this is similar to convergence results for the interior point algorithms that can be used for direct ℓ_1 -minimization.)

The potential of IRLS algorithms, tailored to mimic ℓ_1 -minimization and so recover sparse solutions, has recently been investigated numerically by Chartrand and several co-authors [11, 12, 14]. Our work provides proofs of several findings listed in these works.

One of the virtues of our approach is that, with minor technical modifications, it allows a similar detailed analysis of IRLS algorithms with weights that promote the *non-convex* optimization of ℓ_τ -norms for $0 < \tau < 1$. We can show not only that these algorithms can again recover sparse solutions, but also that their local rate of convergence is superlinear and tends to quadratic when τ tends to zero. Thus we also justify theoretically the recent numerical results by Chartrand et al. concerning such non-convex ℓ_τ -norm optimization [11, 12, 13, 36].

An outline of our paper is the following. In the next section we make some remarks about ℓ_1 - and weighted ℓ_2 -minimization, upon which we shall call in our proof. In the following section, we recall the Restricted Isometry Property and the Null Space Property including some of its consequences that are important to our analysis. In section 4, we gather some preliminary results we shall need to prove our main convergence result, Theorem 5.3, which is formulated and proved in section 5. We then turn to the issue on rate of convergence in section 6. In section 7 we generalize the convergence results obtained for ℓ_1 -minimization to the case of ℓ_τ -spaces for $0 < \tau < 1$; in particular, we show, with Theorem 7.9, the local superlinear convergence of the IRLS algorithm in this setting. We conclude the paper with a short section dedicated to a few numerical examples that dovetail nicely with the theoretical results.

2 Characterization of ℓ_1 - and weighted ℓ_2 -minimizers

We fix $y \in \mathbb{R}^m$ and consider the underdetermined system $\Phi x = y$. Given a norm $\|\cdot\|$, the problem of minimizing $\|z\|$ over $z \in \mathcal{F}(y)$ can be viewed as a problem of approximation. Namely, for any $x_0 \in \mathcal{F}(y)$, we can characterize the minimizers in $\mathcal{F}(y)$ as exactly those elements $z \in \mathcal{F}(y)$ that can be written as $z = x_0 + \eta$, with η a best approximation to $-x_0$ from $\mathcal{N}$. In this way one can characterize minimizers z from classical results on best approximation in normed spaces. We consider two examples of this in the present section, corresponding to the ℓ_1 -norm and the weighted $\ell_2(w)$ -norm.

Throughout this paper, we shall denote by x any element from $\mathcal{F}(y)$ that has smallest ℓ_1 -norm, as in (1.2). When x is unique, we shall emphasize this by denoting it by x^* . In general, x and x^* need not be sparse, although we will often consider cases where they are. We begin with the following well-known lemma (see for example Pinkus [34]) which characterizes the minimal ℓ_1 -norm elements from $\mathcal{F}(y)$.

Lemma 2.1 *An element $x \in \mathcal{F}(y)$ has minimal ℓ_1 -norm among all elements $z \in \mathcal{F}(y)$ if and only if*

$$\left| \sum_{x_i \neq 0} \text{sign}(x_i) \eta_i \right| \leq \sum_{x_i = 0} |\eta_i|, \quad \eta \in \mathcal{N}. \quad (2.1)$$

Moreover, x is unique if and only if we have strict inequality in (2.1) for all $\eta \in \mathcal{N}$ which are not identically zero.

Proof: We give the simple proof for completeness of this paper. If $x \in \mathcal{F}(y)$ has minimum ℓ_1 -norm, then we have, for any $\eta \in \mathcal{N}$ and any $t \in \mathbb{R}$,

$$\sum_{i=1}^N |x_i + t\eta_i| \geq \sum_{i=1}^N |x_i|. \quad (2.2)$$

Fix $\eta \in \mathcal{N}$. If t is sufficiently small then $x_i + t\eta_i$ and x_i will have the same sign $s_i := \text{sign}(x_i)$ whenever $x_i \neq 0$. Hence, (2.2) can be written as

$$t \sum_{x_i \neq 0} s_i \eta_i + \sum_{x_i = 0} |t\eta_i| \geq 0.$$

Choosing t of an appropriate sign, we see that (2.1) is a necessary condition.

For the opposite direction, we note that if (2.1) holds then for each $\eta \in \mathcal{N}$, we have

$$\begin{aligned} \sum_{i=1}^N |x_i| &= \sum_{x_i \neq 0} s_i x_i = \sum_{x_i \neq 0} s_i (x_i + \eta_i) - \sum_{x_i \neq 0} s_i \eta_i \\ &\leq \sum_{x_i \neq 0} s_i (x_i + \eta_i) + \sum_{x_i = 0} |\eta_i| \leq \sum_{i=1}^N |x_i + \eta_i|, \end{aligned} \quad (2.3)$$

where the first inequality uses (2.1).

If x is unique then we have strict inequality in (2.2) and hence subsequently in (2.1). If we have strict inequality in (2.1) then the subsequent strict inequality in (2.3) implies uniqueness. ■

Remark 2.2 Applying Lemma 2.1 to the special case of ℓ_1 -minimizers with no vanishing entries, we see that a vector $x \in \mathcal{F}(y)$, with $x_i \neq 0$ for all $i = 1, \dots, N$, is a minimal ℓ_1 -norm solution if and only if

$$\sum_{i=1}^N s_i \eta_i = 0, \quad \text{for all } \eta \in \mathcal{N}. \quad (2.4)$$

This implies that a minimal ℓ_1 -norm solution to $\Phi x = y$ for which all entries are non-vanishing is necessarily non-unique, by the following argument. Suppose that $x_i \neq 0$ for all $i = 1, \dots, N$ and that $x \in \mathcal{F}(y)$ is a minimal ℓ_1 -norm solution. Pick now any $\eta \in \mathcal{N}$, $\eta \neq 0$, and pick $t > 0$ so that $t < \min_{\eta_i \neq 0} |x_i|/|\eta_i|$; it then follows that $s_i = \text{sign}(x_i + t\eta_i)$ for all $i = 1, \dots, N$. But then we have $\sum_{i=1}^N |x_i + t\eta_i| = \sum_{i=1}^N s_i (x_i + t\eta_i) = \sum_{i=1}^N |x_i|$ by (2.4), so that $x + t\eta$ is also a minimal solution, different from x . Hence, unique ℓ_1 -minimizers are necessarily k -sparse for some $k < N$. □

We next consider minimization in a weighted $\ell_2(w)$ -norm. We suppose that the weight w is *strictly positive* which we define to mean that $w_j > 0$ for all $j \in \{1, \dots, N\}$. In this case, $\ell_2(w)$ is

a Hilbert space with the inner product

$$\langle u, v \rangle_w := \sum_{j=1}^N w_j u_j v_j. \quad (2.5)$$

We define

$$x^w := \operatorname{argmin}_{z \in \mathcal{F}(y)} \|z\|_{\ell_2^N(w)}. \quad (2.6)$$

Because the $\|\cdot\|_{\ell_2^N(w)}$ -norm is strictly convex, the minimizer x^w is necessarily unique; it is completely characterized by the orthogonality conditions

$$\langle x^w, \eta \rangle_w = 0, \quad \forall \eta \in \mathcal{N}. \quad (2.7)$$

Namely, x^w necessarily satisfies (2.7); on the other hand, any element $z \in \mathcal{F}(y)$ that satisfies $\langle z, \eta \rangle_w = 0$ for all $\eta \in \mathcal{N}$ is automatically equal to x^w .

At this point, we would like to tabulate some of the notation we have used in this paper to denote various kinds of minimizers and other solutions alike (such as limits of algorithms).

z	an (arbitrary) element of $\mathcal{F}(y)$
x	any solution of $\min_{z \in \mathcal{F}(y)} \ z\ _{\ell_1}$
x^*	unique solution of $\min_{z \in \mathcal{F}(y)} \ z\ _{\ell_1}$ (notation used only when the minimizer is unique)
x^w	unique solution of $\min_{z \in \mathcal{F}(y)} \ z\ _{\ell_2(w)}$, $w_j > 0$ for all j
$\bar{x}$	limit of Algorithm 1
x^ϵ	unique solution of $\min_{z \in \mathcal{F}(y)} f_\epsilon(z)$; see (5.4)

Table 1: Notation for solutions and minimizers.

3 The Restricted Isometry and the Null Space Properties

To analyze the convergence of our algorithm, we shall impose the Restricted Isometry Property (RIP) already mentioned in the introduction, or a slightly weaker version, the *Null Space Property*, which will be defined below. Recall that Φ satisfies RIP of order L for $\delta \in (0, 1)$ (see (1.11)) iff

$$(1 - \delta)\|z\|_{\ell_2^N} \leq \|\Phi z\|_{\ell_2^m} \leq (1 + \delta)\|z\|_{\ell_2^N}, \quad \text{for all } L\text{-sparse } z. \quad (3.1)$$

It is known that many families of matrices satisfy the RIP. While there are deterministic families that are known to satisfy RIP, the largest range of L , (asymptotically, as $N \rightarrow \infty$, with e.g. m/N kept constant) is obtained (to date) by using random families. For example, random families in which the entries of the matrix Φ are independent realizations of a (fixed) Gaussian or Bernoulli

random variable are known to have the RIP with high probability for each $L \leq c_0(\delta) n / \log n$ (see [7, 4, 2, 35] for a discussion of these results).

We shall say that Φ has the *Null Space Property* (NSP) of order L for $\gamma > 0$ if ¹

$$\|\eta_T\|_{\ell_1} \leq \gamma \|\eta_{T^c}\|_{\ell_1}, \quad (3.2)$$

for all sets T of cardinality not exceeding L and all $\eta \in \mathcal{N}$. Here and later, we denote by η_S the vector obtained from η by setting to zero all coordinates η_i for $i \notin S \subset \{1, 2, \dots, N\}$; T^c denotes the complement of the set T . It is shown in Lemma 4.1 of [18] that if Φ has the RIP of order $L := J + J'$ for a given $\delta \in (0, 1)$, where $J, J' \geq 1$ are integers, then Φ has the NSP of order K for $\gamma := \frac{1+\delta}{1-\delta} \sqrt{\frac{J}{J'}}$. Note that if J' is sufficiently large then $\gamma < 1$.

Another result in [18] (see also Lemma 4.3 below) states that in order to guarantee that a k -sparse vector x^* is the unique ℓ_1 -minimizer in $\mathcal{F}(y)$, it is sufficient that Φ has the NSP of order $L \geq k$ and $\gamma < 1$. (In fact, the argument in [4], proving that for Φ with the RIP, ℓ_1 -minimization identifies sparse vectors in $\mathcal{F}(y)$, can be split into two steps: one that implicitly derives the NSP from the RIP, and the remainder of the proof, which uses only the NSP.)

Note that if the NSP holds for some order L_0 and constant γ_0 (not necessarily < 1), then, by choosing $a > 0$ sufficiently small, one can ensure that Φ has the NSP of order $L = aL_0$ with constant $\gamma < 1$ (see [18] for details). So the effect of requiring that $\gamma < 1$ is tantamount to reducing the range of L slightly.

When proving results on the convergence of our algorithm later in this paper, we shall state them under the assumptions that Φ has the NSP for some $\gamma < 1$ and an appropriate value of L . Using the observations above, they can easily be rephrased in terms of RIP bounds for Φ .

4 Preliminary results

We first make some comments about the decreasing rearrangement $r(z)$ and the j -term approximation errors for vectors in $\mathbb{R}^N$. Let us denote by Σ_k the set of all $x \in \mathbb{R}^N$ such that $\#(\text{supp}(x)) \leq k$. For any $z \in \mathbb{R}^N$ and any $j = 1, 2, \dots, N$, we denote by

$$\sigma_j(z)_{\ell_1} := \inf_{w \in \Sigma_j} \|z - w\|_{\ell_1^N} \quad (4.1)$$

the ℓ_1 -error in approximating a general vector $z \in \mathbb{R}^N$ by a j -sparse vector. Note that these approximation errors can be written as a sum of entries of $r(u)$: $\sigma_j(z)_{\ell_1} = \sum_{\nu > j} r(z)_\nu$. We have the following lemma:

¹This definition of the Null Space Property is a slight variant of that given in [18] but is more convenient for the results in the present paper.

Lemma 4.1 *The map $z \mapsto r(z)$ is Lipschitz continuous on $(\mathbb{R}^N, \|\cdot\|_{\ell_\infty})$: for any $z, z' \in \mathbb{R}^N$, we have*

$$\|r(z) - r(z')\|_{\ell_\infty} \leq \|z - z'\|_{\ell_\infty}. \quad (4.2)$$

Moreover, for any j , we have

$$|\sigma_j(z)_{\ell_1} - \sigma_j(z')_{\ell_1}| \leq \|z - z'\|_{\ell_1}, \quad (4.3)$$

and for any $J > j$, we have

$$(J - j)r(z)_J \leq \|z - z'\|_{\ell_1} + \sigma_j(z')_{\ell_1}. \quad (4.4)$$

Proof: For any pair of points z and z' , and any $j \in \{1, \dots, N\}$, let Λ be a set of $j - 1$ indices corresponding to the $j - 1$ largest entries in z' . Then

$$r(z)_j \leq \max_{i \in \Lambda^c} |z_i| \leq \max_{i \in \Lambda^c} |z'_i| + \|z - z'\|_{\ell_\infty} = r(z')_j + \|z - z'\|_{\ell_\infty}. \quad (4.5)$$

We can also reverse the roles of z and z' . Therefore, we obtain (4.2). To prove (4.3), we approximate z by a j -term best approximation $u \in \Sigma_j$ of z' in ℓ_1 . Then

$$\sigma_j(z)_{\ell_1} \leq \|z - u\|_{\ell_1} \leq \|z - z'\|_{\ell_1} + \sigma_j(z')_{\ell_1},$$

and the result follows from symmetry.

To prove (4.4), it suffices to note that $(J - j)r(z)_J \leq \sigma_j(z)_{\ell_1}$. ■

Our next result is an approximate reverse triangle inequality for points in $\mathcal{F}(y)$. Its importance to us lies in its implication that whenever two points $z, z' \in \mathcal{F}(y)$ have close ℓ_1 -norms and one of them is close to a k -sparse vector, then they necessarily are close to each other. (Note that it also implies that the other vector must then also be close to that k -sparse vector.) This is a geometric property of the null space.

Lemma 4.2 *Assume that (3.2) holds for some L and $\gamma < 1$. Then, for any $z, z' \in \mathcal{F}(y)$, we have*

$$\|z' - z\|_{\ell_1} \leq \frac{1 + \gamma}{1 - \gamma} (\|z'\|_{\ell_1} - \|z\|_{\ell_1} + 2\sigma_L(z)_{\ell_1}). \quad (4.6)$$

Proof: Let T be a set of indices of the L largest entries in z . Then

$$\begin{aligned} \|(z' - z)_{T^c}\|_{\ell_1} &\leq \|z'_{T^c}\|_{\ell_1} + \|z_{T^c}\|_{\ell_1} \\ &= \|z'\|_{\ell_1} - \|z'_T\|_{\ell_1} + \sigma_L(z)_{\ell_1} \\ &= \|z\|_{\ell_1} + \|z'\|_{\ell_1} - \|z\|_{\ell_1} - \|z'_T\|_{\ell_1} + \sigma_L(z)_{\ell_1} \\ &= \|z_T\|_{\ell_1} - \|z'_T\|_{\ell_1} + \|z'\|_{\ell_1} - \|z\|_{\ell_1} + 2\sigma_L(z)_{\ell_1} \\ &\leq \|(z' - z)_T\|_{\ell_1} + \|z'\|_{\ell_1} - \|z\|_{\ell_1} + 2\sigma_L(z)_{\ell_1}. \end{aligned} \quad (4.7)$$

Using (3.2), this gives

$$\|(z' - z)_T\|_{\ell_1} \leq \gamma \|(z' - z)_{T^c}\|_{\ell_1} \leq \gamma (\|(z' - z)_T\|_{\ell_1} + \|z'\|_{\ell_1} - \|z\|_{\ell_1} + 2\sigma_L(z)_{\ell_1}). \quad (4.8)$$

In other words,

$$\|(z' - z)_T\|_{\ell_1} \leq \frac{\gamma}{1 - \gamma} (\|z'\|_{\ell_1} - \|z\|_{\ell_1} + 2\sigma_L(z)_{\ell_1}). \quad (4.9)$$

Using this, together with (4.7), we obtain

$$\|z' - z\|_{\ell_1} = \|(z' - z)_{T^c}\|_{\ell_1} + \|(z' - z)_T\|_{\ell_1} \leq \frac{1 + \gamma}{1 - \gamma} (\|z'\|_{\ell_1} - \|z\|_{\ell_1} + 2\sigma_L(z)_{\ell_1}), \quad (4.10)$$

as desired. $\blacksquare$

This result then allows the following simple proof of some of the results of [18]:

Lemma 4.3 *Assume that (3.2) holds for some L and $\gamma < 1$. Suppose that $\mathcal{F}(y)$ contains an L -sparse vector. Then this vector is the unique ℓ_1 -minimizer in $\mathcal{F}(y)$; denoting it by x^* , we have moreover, for all $v \in \mathcal{F}(y)$,*

$$\|v - x^*\|_{\ell_1} \leq 2 \frac{1 + \gamma}{1 - \gamma} \sigma_L(v)_{\ell_1}. \quad (4.11)$$

Proof: For the time being, we denote the L -sparse vector in $\mathcal{F}(y)$ by x_s . Applying (4.6) with $z' = v$ and $z = x_s$, we find

$$\|v - x_s\|_{\ell_1} \leq \frac{1 + \gamma}{1 - \gamma} [\|v\|_{\ell_1} - \|x_s\|_{\ell_1}];$$

since $v \in \mathcal{F}(y)$ is arbitrary, this implies that $\|v\|_{\ell_1} - \|x_s\|_{\ell_1} \geq 0$ for all $v \in \mathcal{F}(y)$, so that x_s is an ℓ_1 -norm minimizer in $\mathcal{F}(y)$.

If x' were another ℓ_1 -minimizer in $\mathcal{F}(y)$, then it would follow that $\|x'\|_{\ell_1} = \|x_s\|_{\ell_1}$, and the inequality we just derived would imply $\|x' - x_s\|_{\ell_1} = 0$, or $x' = x_s$. It follows that x_s is the unique ℓ_1 -minimizer in $\mathcal{F}(y)$, which we denote by x^* , as proposed earlier.

Finally, we apply (4.6) with $z' = x^*$ and $z = v$, and we obtain

$$\|v - x^*\| \leq \frac{1 + \gamma}{1 - \gamma} (\|x^*\|_{\ell_1} - \|v\|_{\ell_1} + 2\sigma_L(v)_{\ell_1}) \leq 2 \frac{1 + \gamma}{1 - \gamma} \sigma_L(v)_{\ell_1},$$

where we have used the ℓ_1 -minimization property of x^* . $\blacksquare$

Our next set of remarks centers around the functional $\mathcal{J}$ defined by (1.5). Note that for each $n = 1, 2, \dots$, we have

$$\mathcal{J}(x^{n+1}, w^{n+1}, \epsilon_{n+1}) = \sum_{j=1}^N [(x_j^{n+1})^2 + \epsilon_{n+1}^2]^{1/2}. \quad (4.12)$$

We also have the following monotonicity property which holds for all $n \geq 0$:

$$\mathcal{J}(x^{n+1}, w^{n+1}, \epsilon_{n+1}) \leq \mathcal{J}(x^{n+1}, w^n, \epsilon_{n+1}) \leq \mathcal{J}(x^{n+1}, w^n, \epsilon_n) \leq \mathcal{J}(x^n, w^n, \epsilon_n). \quad (4.13)$$

Here the first inequality follows from the minimization property that defines w^{n+1} , the second inequality from $\epsilon_{n+1} \leq \epsilon_n$, and the last inequality from the minimization property that defines

x^{n+1} . For each n , x^{n+1} is completely determined by w^n ; for $n = 0$, in particular, x^1 is determined solely by w^0 , and independent of the choice of $x^0 \in \mathcal{F}(y)$. (With the initial weight vector defined by $w^0 = (1, \dots, 1)$, x^1 is the classical minimum ℓ_2 -norm element of $\mathcal{F}(y)$.) The inequality (4.13) for $n = 0$ thus holds for arbitrary $x^0 \in \mathcal{F}(y)$.

Lemma 4.4 *For each $n \geq 1$ we have*

$$\|x^n\|_{\ell_1} \leq \mathcal{J}(x^1, w^0, \epsilon_0) =: A \quad (4.14)$$

and

$$w_j^n \geq A^{-1}, \quad j = 1, \dots, N. \quad (4.15)$$

Proof: The bound (4.14) follows from (4.13) and

$$\|x^n\|_{\ell_1} \leq \sum_{j=1}^N [(x_j^n)^2 + \epsilon_n^2]^{1/2} = \mathcal{J}(x^n, w^n, \epsilon_n).$$

The bound (4.15) follows from $(w_j^n)^{-1} = [(x_j^n)^2 + \epsilon_n^2]^{1/2} \leq \mathcal{J}(x^n, w^n, \epsilon_n) \leq A$, where the last inequality uses (4.13). $\blacksquare$

5 Convergence of the algorithm

In this section, we prove that the algorithm converges. Our starting point is the following lemma that establishes $(x^n - x^{n+1}) \rightarrow 0$ for $n \rightarrow \infty$.

Lemma 5.1 *Given any $y \in \mathbb{R}^m$, the x^n satisfy*

$$\sum_{n=1}^{\infty} \|x^{n+1} - x^n\|_{\ell_2}^2 \leq 2A^2. \quad (5.1)$$

where A is the constant of Lemma 4.4. In particular, we have

$$\lim_{n \rightarrow \infty} (x^n - x^{n+1}) = 0. \quad (5.2)$$

Proof: For each $n = 1, 2, \dots$, we have

$$\begin{aligned} 2[\mathcal{J}(x^n, w^n, \epsilon_n) - \mathcal{J}(x^{n+1}, w^{n+1}, \epsilon_{n+1})] &\geq 2[\mathcal{J}(x^n, w^n, \epsilon_n) - \mathcal{J}(x^{n+1}, w^n, \epsilon_n)] \\ &= \langle x^n, x^n \rangle_{w^n} - \langle x^{n+1}, x^{n+1} \rangle_{w^n} \\ &= \langle x^n + x^{n+1}, x^n - x^{n+1} \rangle_{w^n} \\ &= \langle x^n - x^{n+1}, x^n - x^{n+1} \rangle_{w^n} \\ &= \sum_{j=1}^N w_j^n (x_j^n - x_j^{n+1})^2 \end{aligned}$$

$$\geq A^{-1} \|x^n - x^{n+1}\|_{\ell_2}^2, \quad (5.3)$$

where the third equality uses the fact that $\langle x^{n+1}, x^n - x^{n+1} \rangle_{w^n} = 0$ (observe that $x^{n+1} - x^n \in \mathcal{N}$ and invoke (2.7)), and the inequality uses the bound (4.15) on the weights. If we now sum these inequalities over $n \geq 1$, we arrive at (5.1). $\blacksquare$

From the monotonicity of ϵ_n , we know that $\epsilon := \lim_{n \rightarrow \infty} \epsilon_n$ exists and is non-negative. The following functional will play an important role in our proof of convergence:

$$f_\epsilon(z) := \sum_{j=1}^N (z_j^2 + \epsilon^2)^{1/2}. \quad (5.4)$$

Notice that if we knew that x^n converged to x then, in view of (4.12), $f_\epsilon(x)$ would be the limit of $\mathcal{J}(x^n, w^n, \epsilon_n)$. When $\epsilon > 0$ the functional f_ϵ is strictly convex and therefore has a unique minimizer

$$x^\epsilon := \operatorname{argmin}_{z \in \mathcal{F}(y)} f_\epsilon(z). \quad (5.5)$$

This minimizer is characterized by the following lemma:

Lemma 5.2 *Let $\epsilon > 0$ and $z \in \mathcal{F}(y)$. Then $z = x^\epsilon$ if and only if $\langle z, \eta \rangle_{\tilde{w}(z, \epsilon)} = 0$ for all $\eta \in \mathcal{N}$, where $\tilde{w}(z, \epsilon)_i = [z_i^2 + \epsilon^2]^{-1/2}$.*

Proof: For the “only if” part, let $z = x^\epsilon$ and $\eta \in \mathcal{N}$ be arbitrary. Consider the analytic function

$$G_\epsilon(t) := f_\epsilon(z + t\eta) - f_\epsilon(z).$$

We have $G_\epsilon(0) = 0$, and by the minimization property $G_\epsilon(t) \geq 0$ for all $t \in \mathbb{R}$. Hence, $G'_\epsilon(0) = 0$. A simple calculation reveals that

$$G'_\epsilon(0) = \sum_{j=1}^N \frac{\eta_j z_j}{[z_j^2 + \epsilon^2]^{1/2}} = \langle z, \eta \rangle_{\tilde{w}(z, \epsilon)},$$

which gives the desired result.

For the “if” part, assume that $z \in \mathcal{F}(y)$ and $\langle z, \eta \rangle_{\tilde{w}(z, \epsilon)} = 0$ for all $\eta \in \mathcal{N}$, where $\tilde{w}(z, \epsilon)$ is defined as above. We shall show that z is a minimizer of f_ϵ on $\mathcal{F}(y)$. Indeed, consider the convex univariate function $[u^2 + \epsilon^2]^{1/2}$. For any point u_0 we have from convexity that

$$[u^2 + \epsilon^2]^{1/2} \geq [u_0^2 + \epsilon^2]^{1/2} + [u_0^2 + \epsilon^2]^{-1/2} u_0 (u - u_0), \quad (5.6)$$

because the right side is the linear function which is tangent to this function at u_0 . It follows that for any point $v \in \mathcal{F}(y)$ we have

$$f_\epsilon(v) \geq f_\epsilon(z) + \sum_{j=1}^N [z_j^2 + \epsilon^2]^{-1/2} z_j (v_j - z_j) = f_\epsilon(z) + \langle z, v - z \rangle_{\tilde{w}(z, \epsilon)} = f_\epsilon(z), \quad (5.7)$$

where we have used the orthogonality condition (5.13) and the fact that $v - z$ is in $\mathcal{N}$. Since v is arbitrary, it follows that $z = x^\epsilon$, as claimed. $\blacksquare$

We now give the convergence of the algorithm.

Theorem 5.3 *Let K (the same index as used in the update rule (1.7)) be chosen so that Φ satisfies the Null Space Property (3.2) of order K , with $\gamma < 1$. Then, for each $y \in \mathbb{R}^m$, the output of Algorithm 1 converges to a vector $\bar{x}$, with $r(\bar{x})_{K+1} = N \lim_{n \rightarrow \infty} \epsilon_n$ and the following hold:*

(i) *If $\epsilon = \lim_{n \rightarrow \infty} \epsilon_n = 0$, then $\bar{x}$ is K -sparse; in this case there is therefore a unique ℓ_1 -minimizer x^* , and $\bar{x} = x^*$; moreover, we have, for $k \leq K$, and any $z \in \mathcal{F}(y)$,*

$$\|z - \bar{x}\|_{\ell_1} \leq c \sigma_k(z)_{\ell_1}, \quad \text{with } c := \frac{2(1+\gamma)}{1-\gamma} \quad (5.8)$$

(ii) *If $\epsilon = \lim_{n \rightarrow \infty} \epsilon_n > 0$, then $\bar{x} = x^\epsilon$;*

(iii) *In this last case, if γ satisfies the stricter bound $\gamma < 1 - \frac{2}{K+2}$ (or, equivalently, if $\frac{2\gamma}{1-\gamma} < K$), then we have, for all $z \in \mathcal{F}(y)$ and any $k < K - \frac{2\gamma}{1-\gamma}$, that*

$$\|z - \bar{x}\|_{\ell_1} \leq \tilde{c} \sigma_k(z)_{\ell_1}, \quad \text{with } \tilde{c} := \frac{2(1+\gamma)}{1-\gamma} \left[\frac{K - k + \frac{3}{2}}{K - k - \frac{2\gamma}{1-\gamma}} \right] \quad (5.9)$$

As a consequence, this case is excluded if $\mathcal{F}(y)$ contains a vector of sparsity $k < K - \frac{2\gamma}{1-\gamma}$.

The constant $\tilde{c}$ can be quite reasonable; for instance, if $\gamma \leq 1/2$ and $k \leq K - 3$, then we have $\tilde{c} \leq 9 \frac{1+\gamma}{1-\gamma} \leq 27$.

Proof: Note that since $\epsilon_{n+1} \leq \epsilon_n$, the ϵ_n always converge. We start by considering the case $\epsilon := \lim_{n \rightarrow \infty} \epsilon_n = 0$.

Case $\epsilon = 0$: In this case, we want to prove that x^n converges, and that its limit is an ℓ_1 -minimizer. Suppose that $\epsilon_{n_0} = 0$ for some n_0 . Then by the definition of the algorithm, we know that the iteration is stopped at $n = n_0$, and $x^n = x^{n_0}$, $n \geq n_0$. Therefore $\bar{x} = x^{n_0}$. From the definition of ϵ_n , it then also follows that $r(x^{n_0})_{K+1} = 0$ and so $\bar{x} = x^{n_0}$ is K -sparse. As noted in §3 and Lemma 4.3, if a K -sparse solution exists when Φ satisfies the NSP of order K with $\gamma < 1$, then it is the unique ℓ_1 -minimizer. Therefore, $\bar{x}$ equals x^* , this unique minimizer.

Suppose now that $\epsilon_n > 0$ for all n . Since $\epsilon_n \rightarrow 0$, there is an increasing sequence of indices (n_i) such that $\epsilon_{n_i} < \epsilon_{n_{i-1}}$ for all i . By the definition (1.7) of $(\epsilon_n)_{n \in \mathbb{N}}$, we must have $r(x^{n_i})_{K+1} < N \epsilon_{n_{i-1}}$ for all i . Noting that $(x^n)_{n \in \mathbb{N}}$ is a bounded sequence, there exists a subsequence $(p_j)_{j \in \mathbb{N}}$ of $(n_i)_{i \in \mathbb{N}}$ such that $(x^{p_j})_{j \in \mathbb{N}}$ converges to a point $\tilde{x} \in \mathcal{F}(y)$. By Lemma 4.1, we know that $r(x^{p_j})_{K+1}$ converges to $r(\tilde{x})_{K+1}$. Hence we get

$$r(\tilde{x})_{K+1} = \lim_{j \rightarrow \infty} r(x^{p_j})_{K+1} \leq \lim_{j \rightarrow \infty} N \epsilon_{p_j-1} = 0, \quad (5.10)$$

which means that the support-width of $\tilde{x}$ is at most K , i.e. $\tilde{x}$ is K -sparse. By the same token used above, we again have that $\tilde{x} = x^*$, the unique ℓ_1 -minimizer. We must still show that $x^n \rightarrow x^*$. Since $x^{p_j} \rightarrow x^*$ and $\epsilon_{p_j} \rightarrow 0$, (4.12) implies $\mathcal{J}(x^{p_j}, w^{p_j}, \epsilon_{p_j}) \rightarrow \|x^*\|_{\ell_1}$. By the monotonicity property stated in (4.13), we get $\mathcal{J}(x^n, w^n, \epsilon_n) \rightarrow \|x^*\|_{\ell_1}$. Since (4.12) implies

$$\mathcal{J}(x^n, w^n, \epsilon_n) - N\epsilon_n \leq \|x^n\|_{\ell_1} \leq \mathcal{J}(x^n, w^n, \epsilon_n), \quad (5.11)$$

we obtain $\|x^n\|_{\ell_1} \rightarrow \|x^*\|_{\ell_1}$. Finally, we invoke Lemma 4.2 with $z' = x^n$, $z = x^*$, and $k = K$ to get

$$\limsup_{n \rightarrow \infty} \|x^n - x^*\|_{\ell_1} \leq \frac{1 + \gamma}{1 - \gamma} \left(\lim_{n \rightarrow \infty} \|x^n\|_{\ell_1} - \|x^*\|_{\ell_1} \right) = 0, \quad (5.12)$$

which completes the proof that $x^n \rightarrow x^*$ in this case.

Finally, (5.8) follows from (4.11) of Lemma 4.3 (with $L = K$), and the observation that $\sigma_n(z) \geq \sigma_{n'}(z)$ if $n \leq n'$.

Case $\epsilon > 0$: We shall first show that $x^n \rightarrow x^\epsilon$, $n \rightarrow \infty$, with x^ϵ as defined by (5.5). By Lemma 4.4, we know that $(x^n)_{n=1}^\infty$ is a bounded sequence in $\mathbb{R}^N$ and hence this sequence has accumulation points. Let (x^{n_i}) be any convergent subsequence of (x^n) and let $\tilde{x} \in \mathcal{F}(y)$ be its limit. We want to show that $\tilde{x} = x^\epsilon$.

Since $w_j^n = [(x_j^n)^2 + \epsilon_n^2]^{-1/2} \leq \epsilon^{-1}$, it follows that $\lim_{i \rightarrow \infty} w_j^{n_i} = [(\tilde{x}_j)^2 + \epsilon^2]^{-1/2} = \tilde{w}(\tilde{x}, \epsilon)_j =: \tilde{w}_j$, $j = 1, \dots, N$. On the other hand, by invoking Lemma 5.1, we now find that $x^{n_i+1} \rightarrow \tilde{x}$, $i \rightarrow \infty$. It then follows from the orthogonality relations (2.7) that for every $\eta \in \mathcal{N}$, we have

$$\langle \tilde{x}, \eta \rangle_{\tilde{w}} = \lim_{i \rightarrow \infty} \langle x^{n_i+1}, \eta \rangle_{w^{n_i}} = 0. \quad (5.13)$$

Now the “if” part of Lemma 5.2 implies that $\tilde{x} = x^\epsilon$. Hence x^ϵ is the unique accumulation point of $(x^n)_{n \in \mathbb{N}}$ and therefore its limit. This establishes (ii).

To prove the error estimate (5.9) stated in (iii), we first note that for any $z \in \mathcal{F}(y)$, we have

$$\|x^\epsilon\|_{\ell_1} \leq f_\epsilon(x^\epsilon) \leq f_\epsilon(z) \leq \|z\|_{\ell_1} + N\epsilon, \quad (5.14)$$

where the second inequality uses the minimizing property of x^ϵ . Hence it follows that $\|x^\epsilon\|_{\ell_1} - \|z\|_{\ell_1} \leq N\epsilon$. We now invoke Lemma 4.2 to obtain

$$\|x^\epsilon - z\|_{\ell_1} \leq \frac{1 + \gamma}{1 - \gamma} [N\epsilon + 2\sigma_k(z)_{\ell_1}]. \quad (5.15)$$

From Lemma 4.1 and (1.7), we obtain

$$N\epsilon = \lim_{n \rightarrow \infty} N\epsilon_n \leq \lim_{n \rightarrow \infty} r(x^n)_{K+1} = r(x^\epsilon)_{K+1}. \quad (5.16)$$

It follows from (4.4) that

$$\begin{aligned} (K + 1 - k)N\epsilon &\leq (K + 1 - k)r(x^\epsilon)_{K+1} \\ &\leq \|x^\epsilon - z\|_{\ell_1} + \sigma_k(z)_{\ell_1} \end{aligned}$$

$$\leq \frac{1+\gamma}{1-\gamma} [N\epsilon + 2\sigma_k(z)_{\ell_1}] + \sigma_k(z)_{\ell_1}, \quad (5.17)$$

where the last inequality uses (5.15). Since by assumption on K , we have $K - k > \frac{2\gamma}{1-\gamma}$, i.e. $K + 1 - k > \frac{1+\gamma}{1-\gamma}$, we obtain

$$N\epsilon + 2\sigma_k(z)_{\ell_1} \leq \frac{2(K - k) + 3}{(K - k) - \frac{2\gamma}{1-\gamma}} \sigma_k(z)_{\ell_1}.$$

Using this back in (5.15), we arrive at (5.9).

Finally, notice that if $\mathcal{F}(y)$ contains a k -sparse vector (with $k < K - \frac{2\gamma}{1-\gamma}$), then we know already (see §3) that this must be the unique ℓ_1 -minimizer x^* ; it then follows from our arguments above that we must have $\epsilon = 0$. Indeed, if we had $\epsilon > 0$, then (5.17) would hold for $z = x^*$; since x^* is k -sparse, $\sigma_k(x^*)_{\ell_1} = 0$, implying $\epsilon = 0$, a contradiction with the assumption $\epsilon > 0$. This finishes the proof. ■

Remark 5.4 Let us briefly compare our analysis of the IRLS algorithm with ℓ_1 minimization. The latter recovers a k -sparse solution (when one exists) if Φ has the NSP of order K and $k \leq K$. The analysis given in our proof of Theorem 5.3 guarantees that our IRLS algorithm recovers k -sparse x for a slightly smaller range of values k than ℓ_1 -minimization, namely for $k < K - \frac{2\gamma}{1-\gamma}$. Notice that this “gap” vanishes for vanishingly small γ . Although we have no examples to demonstrate, our arguments cannot exclude the case where $\mathcal{F}(y)$ contains a k -sparse vector x^* with $K - \frac{2\gamma}{1-\gamma} \leq k \leq K$ (e.g., if $\gamma \geq 1/3$ and $k = K - 1$), and our IRLS algorithm converges to $\bar{x}$, yet $\bar{x} \neq x^*$. However, note that unless γ is close to 1, the range of k -values in this “gap” is fairly small; for instance, for $\gamma < \frac{1}{3}$, this non-recovery of a k -sparse x^* can happen only if $k = K$. □

Remark 5.5 The constant c in (5.8) is clearly smaller than the constant $\tilde{c}$ in (5.9); it follows that when $k < K - \frac{2\gamma}{1-\gamma}$, the estimate (5.9) holds for all cases, regardless of whether $\epsilon = 0$ or not. □

6 Rate of Convergence

Under the conditions of Theorem 5.3 the algorithm converges to a limit $\bar{x}$; if there is a k -sparse vector in $\mathcal{F}(y)$ with $k < K - \frac{2\gamma}{1-\gamma}$, then this limit coincides with that k -sparse vector, which is then also automatically the unique ℓ_1 -minimizer x^* . In this section our goal is to establish a bound for the rate of convergence in both the sparse and non-sparse cases. In the latter case, the goal is to establish the rate at which x^n approaches to a ball of radius $C_1\sigma_k(x^*)_{\ell_1}$ centered at x^* . We shall work under the same assumptions as in Theorem 5.3.

6.1 Case of k -sparse vectors

Let us begin by assuming that $\mathcal{F}(y)$ contains the k -sparse vector x^* . The algorithm produces the sequence x^n , which converges to x^* , as established above. Let us denote the (unknown) support of the k -sparse vector x^* by T .

We introduce an auxiliary sequence of error vectors $\eta^n \in \mathcal{N}$ via $\eta^n := x^n - x^*$ and

$$E_n := \|\eta^n\|_{\ell_1} = \|x^* - x^n\|_{\ell_1^N}.$$

We know that $E_n \rightarrow 0$. The following theorem gives a bound on the rate of convergence of E_n to zero.

Theorem 6.1 *Assume Φ satisfies NSP of order K with constant γ such that $0 < \gamma < 1 - \frac{2}{K+2}$. Suppose that $k < K - \frac{2\gamma}{1-\gamma}$, $0 < \rho < 1$, and $0 < \gamma < 1 - \frac{2}{K+2}$ are such that*

$$\mu := \frac{\gamma(1+\gamma)}{1-\rho} \left(1 + \frac{1}{K+1-k} \right) < 1.$$

Assume that $\mathcal{F}(y)$ contains a k -sparse vector x^ and let $T = \text{supp}(x^*)$. Let n_0 be such that*

$$E_{n_0} \leq R^* := \rho \min_{i \in T} |x_i^*|. \quad (6.1)$$

Then for all $n \geq n_0$, we have

$$E_{n+1} \leq \mu E_n.$$

Consequently x^n converges to x^ exponentially.*

Remark 6.2 Notice that if γ is sufficiently small, e.g. $\gamma(1+\gamma) < \frac{2}{3}$, then for any $k < K$, there is a $\rho > 0$ for which $\mu < 1$, so we have exponential convergence to x^* whenever x^* is k -sparse. $\square$

Proof: We start with the relation (2.7) with $w = w^n$, $x^w = x^{n+1} = x^* + \eta^{n+1}$, and $\eta = x^{n+1} - x^* = \eta^{n+1}$, which gives

$$\sum_{i=1}^N (x_i^* + \eta_i^{n+1}) \eta_i^{n+1} w_i^n = 0.$$

Rearranging the terms and using the fact that x^* is supported on T , we get

$$\sum_{i=1}^N |\eta_i^{n+1}|^2 w_i^n = - \sum_{i \in T} x_i^* \eta_i^{n+1} w_i^n = - \sum_{i \in T} \frac{x_i^*}{[(x_i^n)^2 + \epsilon_n^2]^{1/2}} \eta_i^{n+1}. \quad (6.2)$$

We will prove the theorem by induction. Let us assume that we have shown $E_n \leq R^*$ already. We then have, for all $i \in T$,

$$|\eta_i^n| \leq \|\eta^n\|_{\ell_1^N} = E_n \leq \rho |x_i^*|,$$

so that

$$\frac{|x_i^*|}{[(x_i^n)^2 + \epsilon_n^2]^{1/2}} \leq \frac{|x_i^*|}{|x_i^n|} = \frac{|x_i^*|}{|x_i^* + \eta_i^n|} \leq \frac{1}{1-\rho}, \quad (6.3)$$

and hence (6.2) combined with (6.3) and NSP gives

$$\sum_{i=1}^N |\eta_i^{n+1}|^2 w_i^n \leq \frac{1}{1-\rho} \|\eta_T^{n+1}\|_{\ell_1} \leq \frac{\gamma}{1-\rho} \|\eta_{T^c}^{n+1}\|_{\ell_1}$$

At the same time, the Cauchy-Schwarz inequality combined with the above estimate yields

$$\begin{aligned} \|\eta_{T^c}^{n+1}\|_{\ell_1}^2 &\leq \left(\sum_{i \in T^c} |\eta_i^{n+1}|^2 w_i^n \right) \left(\sum_{i \in T^c} [(x_i^n)^2 + \epsilon_n^2]^{1/2} \right) \\ &\leq \left(\sum_{i=1}^N |\eta_i^{n+1}|^2 w_i^n \right) \left(\sum_{i \in T^c} [(\eta_i^n)^2 + \epsilon_n^2]^{1/2} \right) \\ &\leq \frac{\gamma}{1-\rho} \|\eta_{T^c}^{n+1}\|_{\ell_1} (\|\eta^n\|_{\ell_1} + N\epsilon_n). \end{aligned} \quad (6.4)$$

If $\eta_{T^c}^{n+1} = 0$, then $x_{T^c}^{n+1} = 0$. In this case x^{n+1} is k -sparse and the algorithm has stopped by definition; since $x^{n+1} - x^*$ is in the null space $\mathcal{N}$, which contains no k -sparse elements other than 0, we have already obtained the solution $x^{n+1} = x^*$. If $\eta_{T^c}^{n+1} \neq 0$, then after canceling the factor $\|\eta_{T^c}^{n+1}\|_{\ell_1}$ in (6.4), we get

$$\|\eta_{T^c}^{n+1}\|_{\ell_1} \leq \frac{\gamma}{1-\rho} (\|\eta^n\|_{\ell_1} + N\epsilon_n),$$

and thus

$$\|\eta^{n+1}\|_{\ell_1} = \|\eta_T^{n+1}\|_{\ell_1} + \|\eta_{T^c}^{n+1}\|_{\ell_1} \leq (1+\gamma) \|\eta_{T^c}^{n+1}\|_{\ell_1} \leq \frac{\gamma(1+\gamma)}{1-\rho} (\|\eta^n\|_{\ell_1} + N\epsilon_n). \quad (6.5)$$

Now, we also have by (1.7) and (4.4)

$$N\epsilon_n \leq r(x^n)_{K+1} \leq \frac{1}{K+1-k} (\|x^n - x^*\|_{\ell_1} + \sigma_k(x^*)_{\ell_1}) = \frac{\|\eta^n\|_{\ell_1}}{K+1-k}, \quad (6.6)$$

since by assumption $\sigma_k(x^*) = 0$. This, together with (6.5), yields the desired bound,

$$E_{n+1} = \|\eta^{n+1}\|_{\ell_1} \leq \frac{\gamma(1+\gamma)}{1-\rho} \left(1 + \frac{1}{K+1-k} \right) \|\eta^n\|_{\ell_1} = \mu E_n.$$

In particular, since $\mu < 1$, we have $E_{n+1} \leq R^*$, which completes the induction step. It follows that $E_{n+1} \leq \mu E_n$ for all $n \geq n_0$. $\blacksquare$

Remark 6.3 Note that the precise update rule (1.7) for ϵ_n does not really intervene in this analysis. If $E_{n_0} \leq R^*$, then the estimate

$$E_{n+1} \leq \mu_0(E_n + N\epsilon_n) \quad \text{with } \mu_0 := \gamma(1+\gamma)/(1-\rho), \quad (6.7)$$

guarantees that all further E_n will be bounded by R^* as well, provided $N\epsilon_n \leq (\mu_0^{-1} - 1)R^*$. It is only in guaranteeing that (6.1) must be satisfied for some n_0 that the update rule plays a role:

indeed, by Theorem 5.3, $E_n \rightarrow 0$ for $n \rightarrow \infty$ if ϵ_n is updated following (1.7), so that (6.1) has to be satisfied eventually.

Other update rules may work as well. If $(\epsilon_n)_{n \in \mathbb{N}}$ is defined so that it is a monotonically decreasing sequence with limit ϵ , then the relation (6.7) immediately implies that

$$\limsup_{n \rightarrow \infty} E_n \leq \frac{\mu_0 N \epsilon}{1 - \mu_0}.$$

In particular, if $\epsilon = 0$, then $E_n \rightarrow 0$. The rate at which $E_n \rightarrow 0$ in this case will depend on μ_0 as well as on the rate with which $\epsilon_n \rightarrow 0$. We shall not quantify this relation, except to note that if $\epsilon_n = O(\beta^n)$ for some $\beta < 1$, then $E_n = O(n\tilde{\mu}^n)$ where $\tilde{\mu} = \max(\mu_0, \beta)$. $\square$

6.2 Case of noisy k -sparse vectors

We show here that the exponential rate of convergence to a k -sparse limit vector can be extended to the case where the “ideal” (i.e. k -sparse) target vector has been corrupted by noise and is therefore only “approximately k -sparse”. More precisely, we no longer assume that $\mathcal{F}(y)$ contains a k -sparse vector; consequently the limit $\bar{x}$ of the x^n need not be an ℓ_1 -minimizer (see Theorem 5.3). If x is any ℓ_1 -minimizer in $\mathcal{F}(y)$, Theorem 5.3 guarantees $\|\bar{x} - x\|_{\ell_1} \leq C\sigma_k(x)_{\ell_1}$; since this is the best level of accuracy guaranteed in the limit, we are in this case interested only in how fast x^n will converge to a ball centered at x with radius given by some (prearranged) multiple of $\sigma_k(x)_{\ell_1}$. (Note that if $\mathcal{F}(y)$ contains several ℓ_1 -minimizers, they all lie within a distance $C'\sigma_k(x)_{\ell_1}$ of each other, so that it does not matter which x we pick.) We shall express the notion that z is “approximately k -sparse with gap ratio C ”, or a “noisy version of a k -sparse vector, with gap ratio C ” by the condition

$$r(z)_k \geq C\sigma_k(z)_{\ell_1}$$

where k is such that Φ has the NSP for some pair K, γ such that $0 \leq k < K - \frac{2\gamma}{1-\gamma}$ (e.g. we could have $K = k + 1$ if $\gamma < 1/2$). If the gap ratio C is much greater than the constant C_1 in (5.9), then exponential convergence can be exhibited for a meaningful number of iterations. Note that this class includes perturbations of any k -sparse vector for which the perturbation is sufficiently small in ℓ^1 -norm (when compared to the unperturbed k -sparse vector).

Our argument for the noisy case will closely resemble the case for the exact k -sparse vectors. However there are some crucial differences that justify our decision to separate these two cases.

We will be interested in only the case $\epsilon > 0$ where we recall that ϵ is the limit of the ϵ_n occurring in the algorithm. This assumption implies $\sigma_k(x)_{\ell_1} > 0$, and can only happen if x is not K -sparse. (As noted earlier, the exact k -sparse case always corresponds to $\epsilon = 0$ if $k < K - \frac{2\gamma}{1-\gamma}$. For k in the region $K - \frac{2\gamma}{1-\gamma} \leq k \leq K$, both $\epsilon = 0$ and $\epsilon > 0$ are theoretical possibilities.)

First, we redefine $\eta^n = x^n - x^\epsilon$, where x^ϵ is the minimizer of f_ϵ on $\mathcal{F}(y)$ and $\epsilon > 0$. We know from Theorem 5.3 that $\eta^n \rightarrow 0$. We again set $E_n = \|\eta^n\|_{\ell_1}$.

Theorem 6.4 Given $0 < \rho < 1$, and integers k, K with $k < K$, assume that Φ satisfies the NSP of order K with constant γ such that all the conditions of Theorem 5.3 are satisfied and, in addition,

$$\mu := \frac{\gamma(1+\gamma)}{1-\rho} \left(1 + \frac{1}{K+1-k} \right) < 1.$$

Suppose $z \in \mathcal{F}(y)$ is “approximately k -sparse with gap ratio C ”, i.e.

$$r(z)_k \geq C \sigma_k(z)_{\ell_1} \quad (6.8)$$

with $C \geq C_1$, where C_1 is as in Theorem 5.3. Let T stand for the set of indices of the k largest entries of x^ϵ , and n_0 be such that

$$E_{n_0} \leq R^* := \rho \min_{i \in T} |x_i^\epsilon| = \rho r(x^\epsilon)_k. \quad (6.9)$$

Then for all $n \geq n_0$, we have

$$E_{n+1} \leq \mu E_n + B \sigma_k(z)_{\ell_1}, \quad (6.10)$$

where $B > 0$ is a constant. Similarly, if we define $\tilde{E}_n = \|x^n - z\|_{\ell_1}$, then

$$\tilde{E}_{n+1} \leq \mu \tilde{E}_n + \tilde{B} \sigma_k(z)_{\ell_1}, \quad (6.11)$$

for $n \geq n_0$, where $\tilde{B} > 0$ is a constant. This implies that x^n converges at an exponential (linear) rate to the ball of radius $\tilde{B}(1-\mu)^{-1} \sigma_k(z)_{\ell_1}$ centered at z .

Remark 6.5 Note that Theorem 5.3 trivially implies the inequalities (6.10) and (6.11) in the limit $n \rightarrow \infty$ since $E_n \rightarrow 0$, $\sigma_k(z)_{\ell_1} > 0$, and $\|\bar{x} - z\|_{\ell_1} \leq C_1 \sigma_k(z)_{\ell_1}$. However, Theorem 6.4 quantifies the event when it is guaranteed that the two measures of error, E_n and $\tilde{E}_n$, must shrink (at least) by a factor $\mu < 1$ at each iteration. As noted above, this corresponds to the range $\sigma_k(z)_{\ell_1} \lesssim E_n, \tilde{E}_n \lesssim r(x^\epsilon)_k$, and would be realized if, say, z is the sum of a k -sparse vector and a fully supported “noise” vector which is sufficiently small in ℓ_1 norm. In this sense, the theorem shows that the rate estimate of Theorem 5.3 extends to a neighborhood of k -sparse vectors.

Proof: First, note that the existence of n_0 is guaranteed by the fact that $E_n \rightarrow 0$ and $R^* > 0$. For the latter, note that Lemma 4.1 and Theorem 5.3 imply

$$r(x^\epsilon)_k \geq r(z)_k - \|z - x^\epsilon\|_{\ell_1} \geq (C - C_1) \sigma_k(z)_{\ell_1},$$

so that $R^* \geq \rho(C - C_1) \sigma_k(z)_{\ell_1} > 0$.

We follow the proof of Theorem 6.1 and consider the orthogonality relation (6.2). Since x^ϵ is not sparse in general, we rewrite (6.2) as

$$\sum_{i=1}^N |\eta_i^{n+1}|^2 w_i^n = - \sum_{i=1}^N x_i^\epsilon \eta_i^{n+1} w_i^n = - \sum_{i \in T \cup T^c} \frac{x_i^\epsilon}{[(x_i^n)^2 + \epsilon_n^2]^{1/2}} \eta_i^{n+1}. \quad (6.12)$$

We deal with the contribution on T in the same way as before:

$$\left| \sum_{i \in T} \frac{x_i^\epsilon}{[(x_i^n)^2 + \epsilon_n^2]^{1/2}} \eta_i^{n+1} \right| \leq \frac{1}{1-\rho} \|\eta_T^{n+1}\|_{\ell_1} \leq \frac{\gamma}{1-\rho} \|\eta_{T^c}^{n+1}\|_{\ell_1}$$

For the contribution on T^c , note that

$$\beta_n := \max_{i \in T^c} \frac{|\eta_i^{n+1}|}{[(x_i^n)^2 + \epsilon_n^2]^{1/2}} \leq \epsilon^{-1} \|\eta^{n+1}\|_{\ell_\infty}.$$

Since $\eta^n \rightarrow 0$ we have $\beta_n \rightarrow 0$. It follows that

$$\left| \sum_{i \in T^c} \frac{x_i^\epsilon}{[(x_i^n)^2 + \epsilon_n^2]^{1/2}} \eta_i^{n+1} \right| \leq \beta_n \sigma_k(x^\epsilon)_{\ell_1} \leq \beta_n (\sigma_k(z)_{\ell_1} + \|x^\epsilon - z\|_{\ell_1}) \leq C_2 \beta_n \sigma_k(z)_{\ell_1}, \quad (6.13)$$

where the second inequality is due to Lemma 4.1, the last one to Theorem 5.3, and $C_2 = C_1 + 1$. Combining these two bounds, we get

$$\sum_{i=1}^N |\eta_i^{n+1}|^2 w_i^n \leq \frac{\gamma}{1-\rho} \|\eta_{T^c}^{n+1}\|_{\ell_1} + C_2 \beta_n \sigma_k(z)_{\ell_1}$$

We combine this again with a Cauchy-Schwarz estimate, to obtain

$$\begin{aligned} \|\eta_{T^c}^{n+1}\|_{\ell_1}^2 &\leq \left(\sum_{i \in T^c} |\eta_i^{n+1}|^2 w_i^n \right) \left(\sum_{i \in T^c} [(x_i^n)^2 + \epsilon_n^2]^{1/2} \right) \\ &\leq \left(\sum_{i=1}^N |\eta_i^{n+1}|^2 w_i^n \right) \left(\sum_{i \in T^c} [|\eta_i^n| + |x_i^\epsilon| + \epsilon_n] \right) \\ &\leq \left(\frac{\gamma}{1-\rho} \|\eta_{T^c}^{n+1}\|_{\ell_1} + C_2 \beta_n \sigma_k(z)_{\ell_1} \right) (\|\eta_{T^c}^n\|_{\ell_1} + \sigma_k(x^\epsilon)_{\ell_1} + N\epsilon_n) \\ &\leq \left(\frac{\gamma}{1-\rho} \|\eta_{T^c}^{n+1}\|_{\ell_1} + C_2 \beta_n \sigma_k(z)_{\ell_1} \right) (\|\eta_{T^c}^n\|_{\ell_1} + C_2 \sigma_k(z)_{\ell_1} + N\epsilon_n), \end{aligned} \quad (6.14)$$

It is easy to check that if $u^2 \leq Au + B$, where A and B are positive, then $u \leq A + B/A$. Applying this to $u = \|\eta_{T^c}^{n+1}\|_{\ell_1}$ in the above estimate, we get

$$\|\eta_{T^c}^{n+1}\|_{\ell_1} \leq \frac{\gamma}{1-\rho} [\|\eta_{T^c}^n\|_{\ell_1} + C_2 \sigma_k(z)_{\ell_1} + N\epsilon_n] + C_3 \beta_n \sigma_k(z)_{\ell_1}, \quad (6.15)$$

where $C_3 = C_2(1-\rho)/\gamma$. Similar to (6.6), we also have, by combining (4.4) with (part of) the chain of inequalities (6.13),

$$N\epsilon_n \leq r(x^n)_{K+1} \leq \frac{1}{K+1-k} (\|x^n - x^\epsilon\|_{\ell_1} + \sigma_k(x^\epsilon)_{\ell_1}) \leq \frac{1}{K+1-k} (\|\eta^n\|_{\ell_1} + C_2 \sigma_k(z)_{\ell_1}), \quad (6.16)$$

and consequently (6.15) becomes

$$\begin{aligned} \|\eta^{n+1}\|_{\ell_1} &\leq (1+\gamma) \|\eta_{T^c}^{n+1}\|_{\ell_1} \\ &\leq \frac{\gamma(1+\gamma)}{1-\rho} \left(1 + \frac{1}{K+1-k} \right) \|\eta^n\|_{\ell_1} + (1+\gamma)(C_3 \beta_n + C_4) \sigma_k(z)_{\ell_1}, \end{aligned} \quad (6.17)$$

where $C_4 = C_2\gamma(1 - \rho)^{-1}(1 + 1/(K + 1 - k))$. Since the β_n are bounded, this gives

$$E_{n+1} \leq \mu E_n + B\sigma_k(z)_{\ell_1}.$$

It then follows that if we pick $\tilde{\mu}$ so that $1 > \tilde{\mu} > \mu$, and consider the range of $n > n_0$ such that $E_n \geq (\tilde{\mu} - \mu)^{-1}B\sigma_k(z)_{\ell_1} =: r^*$, then

$$E_{n+1} \leq \tilde{\mu} E_n.$$

Hence we are guaranteed exponential decay of E_n as long as x^n is sufficiently far from its limit. The smallest possible value of r^* corresponds to the case $\tilde{\mu} \approx 1$.

To establish a rate of convergence to a comparably-sized ball centered at z , we consider $\tilde{E}_n = \|x^n - z\|_{\ell_1}$. It then follows that

$$\begin{aligned} \tilde{E}_{n+1} &\leq \|x^{n+1} - x^\epsilon\|_{\ell_1} + \|x^\epsilon - z\|_{\ell_1} \\ &\leq \mu\|x^n - x^\epsilon\|_{\ell_1} + B\sigma_k(z)_{\ell_1} + C_1\sigma_k(z)_{\ell_1} \\ &\leq \mu\|x^n - z\|_{\ell_1} + B\sigma_k(z)_{\ell_1} + C_1(1 + \mu)\sigma_k(z)_{\ell_1} \\ &= \mu\tilde{E}_n + \tilde{B}\sigma_k(z)_{\ell_1}, \end{aligned} \tag{6.18}$$

which shows the claimed exponential decay and also that

$$\limsup_{n \rightarrow \infty} \tilde{E}_n \leq \tilde{B}(1 - \mu)^{-1}\sigma_k(z)_{\ell_1}.$$

■

7 Beyond the convex case: ℓ_τ -minimization for $\tau < 1$

If Φ has the NSP of order K with $\gamma < 1$, then (see §3) ℓ_1 -minimization recovers K -sparse solutions to $\Phi x = y$ for any $y \in \mathbb{R}^m$ that admits such a k -sparse solution, i.e., ℓ_1 -minimization gives also ℓ_0 -minimizers, provided their support has size at most k . In [29], Gribonval and Nielsen showed that in this case, ℓ_1 -minimization also gives the ℓ_τ -minimizers, i.e., ℓ_1 -minimization also solves non-convex optimization problems of the type

$$x^* = \operatorname{argmin}_{z \in \mathcal{F}(y)} \|z\|_{\ell_\tau}^\tau, \text{ for } 0 < \tau < 1. \tag{7.1}$$

Let us first recall the results of [29] that are of most interest to us here, reformulated for our setting and notations.

Lemma 7.1 ([29, Theorem 2]). *Assume that x^* is a K -sparse vector in $\mathcal{F}(y)$ and that $0 < \tau \leq 1$. If*

$$\sum_{i \in T} |\eta_i|^\tau < \sum_{i \in T^c} |\eta_i|^\tau, \text{ or, equivalently, } \sum_{i \in T} |\eta_i|^\tau < \frac{1}{2} \sum_{i=1}^N |\eta_i|^\tau,$$

for all $\eta \in \mathcal{N}$ and for all $T \subset \{1, \dots, N\}$ with $\#T \leq K$, then

$$x^* = \operatorname{argmin}_{z \in \mathcal{F}(y)} \|z\|_{\ell_\tau^N}^\tau.$$

Lemma 7.2 ([29, Theorem 5]). *Let $z \in \mathbb{R}^N$, $0 < \tau_1 \leq \tau_2 \leq 1$, and $K \in \mathbb{N}$. Then*

$$\sup_{T \subset \{1, \dots, N\}, \#T \leq K} \frac{\sum_{i \in T} |z_i|^{\tau_1}}{\sum_{i=1}^N |z_i|^{\tau_1}} \leq \sup_{T \subset \{1, \dots, N\}, \#T \leq K} \frac{\sum_{i \in T} |z_i|^{\tau_2}}{\sum_{i=1}^N |z_i|^{\tau_2}}.$$

Combining these two lemmas with the observations in §3 leads immediately to the following result.

Theorem 7.3 *Fix any $0 < \tau \leq 1$. If Φ satisfies the NSP of order K with constant γ then*

$$\sum_{i \in T} |\eta_i|^\tau < \gamma \sum_{i \in T^c} |\eta_i|^\tau, \quad (7.2)$$

for all $\eta \in \mathcal{N}$ and for all $T \subset \{1, \dots, N\}$ such that $\#T \leq K$.

In addition, if $\gamma < 1$, and if there exists a K -sparse vector in $\mathcal{F}(y)$, then this K -sparse vector is the unique minimizer in $\mathcal{F}(y)$ of $\|\cdot\|_{\ell_\tau}$.

At first sight, these results suggest there is nothing to be gained by carrying out ℓ_τ - rather than ℓ_1 -minimization; in addition sparse recovery via the non-convex problems (7.1) is much harder than the more easily solvable convex relaxation problem of ℓ_1 -minimization.

Yet, we shall show in this section that ℓ_τ -minimization has unexpected benefits, and that it may be both useful *and* practically feasible via an IRLS approach. Before we start, it is expedient to introduce the following definition: we shall say that Φ has the τ -Null Space Property (τ -NSP) of order K with constant $\gamma > 0$ if, for all sets T of cardinality at most K and all $\eta \in \mathcal{N}$,

$$\|\eta_T\|_{\ell_\tau^N}^\tau \leq \gamma \|\eta_{T^c}\|_{\ell_\tau^N}^\tau. \quad (7.3)$$

In what follows we shall construct an IRLS algorithm for ℓ_τ - minimization. We shall see that

- (a) In practice, ℓ_τ -minimization can be carried out by an IRLS algorithm. Hence, the non-convexity does not necessarily make the problem intractable;
- (b) In particular, if Φ satisfies the τ -NSP of order K , and if there exists a k -sparse vector x^* in $\mathcal{F}(y)$, with $k \leq K - \kappa$ for suitable κ given below, then the IRLS algorithm converges to the ℓ_τ -minimizer x^τ , which, therefore, will coincide with x^* ;
- (c) Surprisingly the rate of local convergence of the algorithm is superlinear; the rate is larger for smaller τ , increasing to approach a quadratic regime as $\tau \rightarrow 0$. More precisely, we will show that the local error $E_n := \|x^n - x^*\|_{\ell_\tau^N}^\tau$ satisfies

$$E_{n+1} \leq \mu(\gamma, \tau) E_n^{2-\tau}, \quad (7.4)$$

where $\mu(\gamma, \tau) < 1$ for $\gamma > 0$ sufficiently small. The validity of (7.4) is restricted to x^n in a (small) ball centered at x^* . In particular, if x^0 is close enough to x^* then (7.4) ensures the convergence of the algorithm to the k -sparse solution x^* .

Some of these virtues of ℓ_τ -minimization were recently highlighted by Chartrand and his collaborators [11, 12, 13]. Chartrand and Staneva [13] give a fine analysis of the RIP from which they can conclude that ℓ_τ -minimization not only recovers k -sparse vectors, but that the range of k for which this recovery works is larger for smaller τ . Namely, for random Gaussian matrices, they prove that with high probability on the draw of the matrix sparse recovery by ℓ_τ -minimization works for $k \leq m[c_1(\tau) + \tau c_2(\tau) \log(N/k)]^{-1}$, where $c_1(\tau)$ is bounded and $c_2(\tau)$ decreases to zero as $\tau \rightarrow 0$. In particular, the dependence of the sparsity k on the number N of columns vanishes for $\tau \rightarrow 0$. These bounds give a quantitative estimate of the improvement provided by ℓ_τ -minimization vis a vis ℓ_1 -minimization for which the range of k -sparsity for having exact recovery is clearly smaller (see Figure 8.4 for a numerical illustration).

7.1 Some useful properties of ℓ_τ spaces

We start by listing in one proposition some fundamental and well-known properties of ℓ_τ spaces for $0 < \tau \leq 1$. For further details we refer the reader to, e.g., [19].

Proposition 7.4

(i) Assume $0 < \tau \leq 1$. Then the map $z \mapsto \|z\|_{\ell_\tau^N}$ defines a quasi-norm for $\mathbb{R}^N$, in particular the triangle inequality holds up to a constant, i.e.,

$$\|u + v\|_{\ell_\tau^N} \leq C(\tau) \left(\|u\|_{\ell_\tau^N} + \|v\|_{\ell_\tau^N} \right), \text{ for all } u, v \in \mathbb{R}^N. \quad (7.5)$$

If one considers the τ -th powers of the “ τ -norm”, then one has the so-called “ τ -triangle inequality”:

$$\|u + v\|_{\ell_\tau^N}^\tau \leq \|u\|_{\ell_\tau^N}^\tau + \|v\|_{\ell_\tau^N}^\tau, \text{ for all } u, v \in \mathbb{R}^N. \quad (7.6)$$

(ii) We have, for any $0 < \tau_1 \leq \tau_2 \leq \infty$

$$\|u\|_{\ell_{\tau_2}} \leq \|u\|_{\ell_{\tau_1}}, \text{ for all } u \in \mathbb{R}^N. \quad (7.7)$$

We will refer to this norm estimate by writing the embedding relation $\ell_{\tau_1}^N \hookrightarrow \ell_{\tau_2}^N$.

(iii) (Generalized Hölder inequality) For $0 < \tau \leq 1$ and $0 < p, q < \infty$ such that $\frac{1}{\tau} = \frac{1}{p} + \frac{1}{q}$, and for a positive weight vector $w = (w_i)_{i=1}^N$ we have

$$\|(u_i v_i)_{i=1}^N\|_{\ell_\tau^N(w)} \leq \|u\|_{\ell_p^N(w)} \|v\|_{\ell_q^N(w)}, \text{ for all } u, v \in \mathbb{R}^N, \quad (7.8)$$

where $\|v\|_{\ell_\tau^N(w)} := \left(\sum_{i=1}^N |v_i|^r w_i \right)^{1/r}$, as usual, for $0 < r < \infty$.

For technical reasons, it is often more convenient to employ the τ -triangle inequality (7.6) than (7.5); in this sense, for ℓ_τ -minimization $\|\cdot\|_{\ell_\tau^N}^\tau$ turns out to be more natural as a measure of error than the quasi-norm $\|\cdot\|_{\ell_\tau^N}$.

In order to prove the three claims (a)-(c) listed before the start of this subsection, we also need to generalize to ℓ_τ certain results previously shown only for ℓ_1 . In the following we assume $0 < \tau \leq 1$. We denote by

$$\sigma_k(z)_{\ell_\tau^N} := \sum_{\nu > k} r(z)_\nu^\tau,$$

the error of the best k -term approximation to z with respect to $\|\cdot\|_{\ell_\tau^N}^\tau$. As a straightforward generalization of analogous results valid for the ℓ_1 -norm, we have the following two technical lemmas.

Lemma 7.5 *For any $j \in \{1, \dots, N\}$, we have*

$$|\sigma_j(z)_{\ell_\tau^N} - \sigma_j(z')_{\ell_\tau^N}| \leq \|z - z'\|_{\ell_\tau^N}^\tau,$$

for all $z, z' \in \mathbb{R}^N$. Moreover, for any $J > j$, we have

$$(J - j)r(z)_j^\tau \leq \sigma_j(z)_{\ell_\tau^N} \leq \|z - z'\|_{\ell_\tau^N}^\tau + \sigma_j(z')_{\ell_\tau^N}.$$

Lemma 7.6 *Assume that Φ has the τ -NSP of order K with constant $0 < \gamma < 1$. Then, for any $z, z' \in \mathcal{F}(y)$, we have*

$$\|z' - z\|_{\ell_\tau^N}^\tau \leq \frac{1 + \gamma}{1 - \gamma} \left(\|z'\|_{\ell_\tau^N}^\tau - \|z\|_{\ell_\tau^N}^\tau + 2\sigma_K(z)_{\ell_\tau^N} \right).$$

The proofs of these lemmas are essentially identical to the ones of Lemma 4.1 and Lemma 4.2, except for substituting $\|\cdot\|_{\ell_\tau^N}^\tau$ for $\|\cdot\|_{\ell_1^N}$ and $\sigma_k(\cdot)_{\ell_\tau^N}$ for $\sigma_k(\cdot)_{\ell_1^N}$ respectively.

7.2 An IRLS algorithm for ℓ_τ -minimization

To define an IRLS algorithm promoting ℓ_τ -minimization for a generic $0 < \tau \leq 1$, we first define a τ -dependent functional $\mathcal{J}_\tau$, generalizing $\mathcal{J}$:

$$\mathcal{J}_\tau(z, w, \epsilon) := \frac{\tau}{2} \left[\sum_{j=1}^N z_j^2 w_j + \sum_{j=1}^N \left(\epsilon^2 w_j + \frac{2 - \tau}{\tau} \frac{1}{w_j^{\frac{\tau}{2 - \tau}}} \right) \right], \quad z \in \mathbb{R}^N, w \in \mathbb{R}_+^N, \epsilon \in \mathbb{R}_+. \quad (7.9)$$

The desired algorithm is then defined simply by substituting $\mathcal{J}_\tau$ for $\mathcal{J}$ in Algorithm 1, keeping the same update rule (1.7) for ϵ . In particular we have

$$w_j^{n+1} = \left((x_j^{n+1})^2 + \epsilon_{n+1}^2 \right)^{-\frac{2 - \tau}{2}}, \quad j = 1, \dots, N,$$

and

$$\mathcal{J}_\tau(x^{n+1}, w^{n+1}, \epsilon_{n+1}) = \sum_{j=1}^N \left((x_j^{n+1})^2 + \epsilon_{n+1}^2 \right)^{\frac{\tau}{2}}.$$

Fundamental properties of the algorithm are derived in the same way as before. In particular, the values $\mathcal{J}_\tau(x^n, w^n, \epsilon_n)$ decrease monotonically,

$$\mathcal{J}_\tau(x^{n+1}, w^{n+1}, \epsilon_{n+1}) \leq \mathcal{J}_\tau(x^n, w^n, \epsilon_n), \quad n \geq 0,$$

and the iterates are bounded,

$$\|x^n\|_{\ell_\tau^N}^\tau \leq \mathcal{J}_\tau(x^1, w^0, \epsilon_0) := A_0.$$

As in Lemma 4.4, the weights are uniformly bounded from below, i.e.,

$$w_j^n \geq \tilde{A}_0, \quad j = 1, \dots, N.$$

Moreover, using $\mathcal{J}_\tau$ for $\mathcal{J}$ in Lemma 5.1, we can again prove the *asymptotic regularity* of the iterations, i.e.,

$$\lim_{n \rightarrow \infty} \|x^{n+1} - x^n\|_{\ell_2^N} = 0.$$

The first significant difference with the ℓ_1 -case arises when $\epsilon = \lim_{n \rightarrow \infty} \epsilon_n > 0$. In this latter situation, we need to consider the function

$$f_\epsilon^\tau(z) := \sum_{j=1}^N (z_j^2 + \epsilon^2)^{\frac{\tau}{2}}. \quad (7.10)$$

We denote by $\mathcal{Z}_{\epsilon, \tau}(y)$ its set of minimizers on $\mathcal{F}(y)$ (since $f_{\epsilon, \tau}$ is no longer convex it may have more than one minimizer). Even though every minimizer $z \in \mathcal{Z}_{\epsilon, \tau}(y)$ still satisfies

$$\langle z, \eta \rangle_w = 0, \quad \text{for all } \eta \in \mathcal{N},$$

where $w = w^{\epsilon, \tau, z}$ is defined by $w_j^{\epsilon, \tau, z} = ((z_j)^2 + \epsilon^2)^{\frac{\tau-2}{2}}$, $j = 1, \dots, N$, the converse need no longer be true.

The following theorem summarizes the convergence properties on the algorithm in the case $\tau < 1$.

Theorem 7.7 *Fix $y \in \mathbb{R}^N$. Let K (the same index as in the update rule (1.7)) be chosen so that Φ satisfies the τ -NSP of order K with a constant γ such that $\gamma < 1 - \frac{2}{K+2}$. Let $\tilde{\mathcal{Z}}_{\epsilon, \tau}(y)$ be the set of accumulation points of $(x^n)_{n \in \mathbb{N}}$, and define $\epsilon := \lim_{n \rightarrow \infty} \epsilon_n$. Then, the algorithm has the following properties:*

- (i) *If $\epsilon = 0$, then $\tilde{\mathcal{Z}}_{\epsilon, \tau}(y)$ consists of a single point $\bar{x}$, the $x^{(n)}$ converge to $\bar{x}$, and $\bar{x}$ is an ℓ_τ -minimizer in $\mathcal{F}(y)$ which is also K -sparse.*
- (ii) *If $\epsilon > 0$, then for each $\bar{x} \in \tilde{\mathcal{Z}}_{\epsilon, \tau}(y)$ we have $\langle \bar{x}, \eta \rangle_{w^{\epsilon, \tau, \bar{x}}} = 0$, for all $\eta \in \mathcal{N}$.*
- (iii) *If $z \in \mathcal{F}(y)$ and $\bar{x} \in \tilde{\mathcal{Z}}_{\epsilon, \tau}(y) \cap \mathcal{Z}_{\epsilon, \tau}(y)$, we have*

$$\|z - \bar{x}\|_{\ell_\tau^N}^\tau \leq C_2 \sigma_k(z)_{\ell_\tau^N},$$

for all $k < K - \frac{2\gamma}{1-\gamma}$.

The proof of this theorem uses Lemmas 7.1-7.6 and follows the same arguments as for Theorem 5.3.

Remark 7.8 Unlike Theorem 5.3, Theorem 7.7 does not ensure that the IRLS algorithm converges to the sparsest or to the minimal ℓ_τ -solution. It does provide conditions that are verifiable a posteriori (e.g., $\epsilon = \lim_{n \rightarrow \infty} \epsilon_n = 0$) for such convergence. The reason for this weaker result is the non-convexity of f_ϵ^τ . (In particular, it might happen that $x^{\epsilon, \tau}$ is a local minimizer of f_ϵ^τ , but not a global one, and the estimate in (iii) does not necessarily hold.) Nevertheless, as is often the case for non-convex problems, we can establish a local convergence result that also highlights the rate we can expect for such convergence. This is the content of the following section; it will be followed by numerical results that dovetail nicely with the theoretical results.

7.3 Local superlinear convergence

Throughout this section, we assume that there exists a k -sparse vector x^* in $\mathcal{F}(y)$. We define the error vectors $\eta^n = x^n - x^* \in \mathcal{N}$; we now measure the error by $\|\cdot\|_\tau^\tau$:

$$E_n := \|\eta^n\|_{\ell_N}^\tau.$$

Theorem 7.9 Assume that Φ has the τ -NSP of order K with constant $\gamma \in (0, 1)$ and that $\mathcal{F}(y)$ contains a k sparse vector x^* with $k \leq K$. (Here K is the same as in the definition of ϵ_n in the update rule (1.7) in Algorithm 1.) Suppose that, for a given $0 < \rho < 1$, we have

$$E_{n_0} \leq R^* := [\rho r(x^*)_k]^\tau \quad (7.11)$$

and define

$$\mu := \mu(\rho, K, \gamma, \tau, N) = 2^{1-\tau} \gamma (1 + \gamma) A^\tau \left(1 + \left(\frac{N^{1-\tau}}{K+1-k} \right)^{2-\tau} \right), \quad A := (r(x^*)_k^{1-\tau} (1 - \rho)^{2-\tau})^{-1}.$$

If ρ and γ are sufficiently small so that

$$\mu(R^*)^{1-\tau} = \mu \rho^{\tau(1-\tau)} r(x^*)_k^{\tau(1-\tau)} \leq 1, \quad (7.12)$$

then for all $n \geq n_0$ we have

$$E_{n+1} \leq \mu E_n^{2-\tau}. \quad (7.13)$$

Proof: The proof is by induction on n . We assume that $E_n \leq R^*$ and derive (7.13). As in the proof of Theorem 6.1, we let T denote the support of x^* and so $\#(T) = k$ and $r(x^*)_k$ is the smallest entry in x^* . Following the proof of Theorem 6.1, the first few lines are the same. The first difference is in the following estimate, which holds for $i \in T$ and replaces (6.3),

$$\begin{aligned} \frac{|x_i^*|}{((x_i^n)^2 + \epsilon_n^2)^{1-\tau/2}} &\leq \frac{|x_i^*|}{|x_i^* + \eta_i^n|^{2-\tau}} \leq \frac{|x_i^*|}{(|x_i^*|(1 - \rho))^{2-\tau}} \\ &= \frac{1}{|x_i^*|^{1-\tau} (1 - \rho)^{2-\tau}} \leq A. \end{aligned}$$

Starting with the orthogonality relation (6.2) and using the above inequality and the embedding $\ell_{\tau_1}^N \hookrightarrow \ell_1^N$, we obtain

$$\sum_{i=1}^N |\eta_i^{n+1}|^2 w_i^n \leq A \left(\sum_{i \in T} |\eta_i^{n+1}|^\tau \right)^{1/\tau}.$$

We now apply the τ -NSP to find

$$\|\eta^{n+1}\|_{\ell_2(w^n)}^{2\tau} = \left(\sum_{i=1}^N |\eta_i^{n+1}|^2 w_i^n \right)^\tau \leq \gamma A^\tau \|\eta_{T^c}^{n+1}\|_{\ell_\tau^N}^\tau. \quad (7.14)$$

At the same time, the generalized Hölder inequality (see Proposition 7.4 (iii)) for $p = 2$ and $q = \frac{2\tau}{2-\tau}$, together with the above estimates, yields

$$\begin{aligned} \|\eta_{T^c}^{n+1}\|_{\ell_\tau^N}^{2\tau} &= \|(|\eta_i^{n+1}|(w_i^n)^{-1/\tau})_{i=1}^N\|_{\ell_\tau^N(w^n; T^c)}^{2\tau} \\ &\leq \|\eta^{n+1}\|_{\ell_2(w^n)}^{2\tau} \|((w_i^n)^{-1/\tau})_{i=1}^N\|_{\ell_{2\tau/(2-\tau)}^N(w^n; T^c)}^{2\tau} \\ &\leq \gamma A^\tau \|\eta_{T^c}^{n+1}\|_{\ell_\tau^N}^\tau \|((w_i^n)^{-1/\tau})_{i=1}^N\|_{\ell_{2\tau/(2-\tau)}^N(w^n; T^c)}^{2\tau} \end{aligned}$$

In other words,

$$\|\eta_{T^c}^{n+1}\|_{\ell_\tau^N}^\tau \leq \gamma A^\tau \|((w_i^n)^{-1/\tau})_{i=1}^N\|_{\ell_{2\tau/(2-\tau)}^N(w^n; T^c)}^{2\tau}. \quad (7.15)$$

Let us now estimate the weight term. By the $\frac{\tau}{2}$ -triangle inequality (7.6) we have

$$\begin{aligned} \|((w_i^n)^{-1/\tau})_{i=1}^N\|_{\ell_{2\tau/(2-\tau)}^N(w^n; T^c)}^{2\tau} &= \left(\sum_{i=1}^N (|\eta_i^n|^2 + \epsilon_n^2)^{\frac{\tau}{2}} \right)^{2-\tau} \\ &\leq \left(\sum_{i=1}^N (|\eta_i^n|^\tau + \epsilon_n^\tau) \right)^{2-\tau} = \left(\sum_{i=1}^N |\eta_i^n|^\tau + N \epsilon_n^\tau \right)^{2-\tau} \\ &\leq 2^{1-\tau} \left(\left(\sum_{i=1}^N |\eta_i^n|^\tau \right)^{2-\tau} + N^{2-\tau} \epsilon_n^{\tau(2-\tau)} \right). \end{aligned}$$

Now, an application of Lemma 7.5 gives the following estimates

$$\begin{aligned} N^{2-\tau} \epsilon_n^{\tau(2-\tau)} &= N^{(1-\tau)(2-\tau)} (N^\tau \epsilon_n^\tau)^{2-\tau} \leq N^{(1-\tau)(2-\tau)} (r(x^n)_{K+1}^\tau)^{2-\tau} \\ &\leq \left(\frac{N^{1-\tau}}{K+1-k} \|x^n - x^*\|_{\ell_\tau^N}^\tau \right)^{2-\tau} = \left(\frac{N^{1-\tau}}{K+1-k} \right)^{2-\tau} \left(\|\eta^n\|_{\ell_\tau^N}^\tau \right)^{2-\tau}. \end{aligned}$$

Using these estimates in (7.15) gives

$$\|\eta_{T^c}^{n+1}\|_{\ell_\tau^N}^\tau \leq 2^{1-\tau} \gamma A^\tau \left(1 + \left(\frac{N^{1-\tau}}{K+1-k} \right)^{2-\tau} \right) \left(\|\eta^n\|_{\ell_\tau^N}^\tau \right)^{2-\tau},$$

and (7.13) follows by a further application of the τ -NSP (see (6.5)).

Because of the assumption (7.12), we also have $E_{n+1} \leq R^*$ and so the induction can continue. $\blacksquare$

Remark 7.10 In contrast to the ℓ_1 case, we do not need $\mu < 1$ to ensure that E_n decreases. In fact, all that is needed for the error reduction is $\mu E_n^{1-\tau} < 1$ for some sufficiently large n . In fact, μ could be quite large in cases where the smallest non-zero component of the sparse vector is very small. We have not observed this effect in our examples; we expect that our analysis, although apparently accurate in describing the rate of convergence (see section 8), is too pessimistic in estimating the coefficient μ .

8 Numerical results

In this section we present numerical experiments that illustrate that the bounds derived in the theoretical analysis do manifest themselves in practice.

8.1 Convergence rates

We start with numerical results that confirm the linear rate of convergence of our iteratively re-weighted least square algorithm for ℓ_1 -minimization, and its robust recovery of sparse vectors. In the experiments we used a matrix Φ of dimensions $m \times N$ and Gaussian $N(0, 1/m)$ i.i.d. entries. Such matrices are known to possess (with high probability) the RIP property with optimal bounds [2, 4, 35]. In Figure 8.1 we depict the approximation error to the unique sparsest solution shown in Figure 8.2, and the instantaneous rate of convergence. The numerical results both confirm the expected linear rate of convergence and the robust reconstruction of the sparse vector.

Next, we compare the linear convergence achieved with ℓ_1 -minimization with the superlinear convergence obtained by the iteratively re-weighted least square algorithm promoting ℓ_τ -minimization.

In Figure 8.3 we are interested in the comparison of the rate of convergence when our algorithm is used for different choices of $0 < \tau \leq 1$. For $\tau = 1, .8, .6$ and $.56$, the figure shows the error, as a function of the iteration step n , for the iterative algorithm, with different fixed values of τ . For $\tau = 1$, the rate is linear, as in Figure 8.1. For the smaller values $\tau = .8, .6$ and $.56$ the iterations initially follow the same linear rate; once they are sufficiently close to the sparse solution, the convergence rate speeds up dramatically, suggesting we have entered the region of validity of (7.13). For smaller values of τ numerical experiments do not always lead to convergence: in some cases the algorithm never got to the neighborhood of the solution where convergence is ensured. However, in this case a combination of initial iterations with the ℓ_1 -inspired IRLS (for which we always have convergence) and later iterations with ℓ_τ -inspired IRLS for smaller τ allow again for a very fast convergence to the sparsest solution; this is illustrated in Figure 8.3 for the case $\tau = .5$.

8.2 Enhanced recovery in compressed sensing and relationship with other work

Candès, Wakin, and Boyd [8] showed, by numerical experimentation, that iteratively re-weighted ℓ_1 -minimization, with weights suggested by an ℓ_0 -minimization goal, can enhance the range of sparsity for which perfect reconstruction of a sparse vector “works” in compressed sensing. In experiments

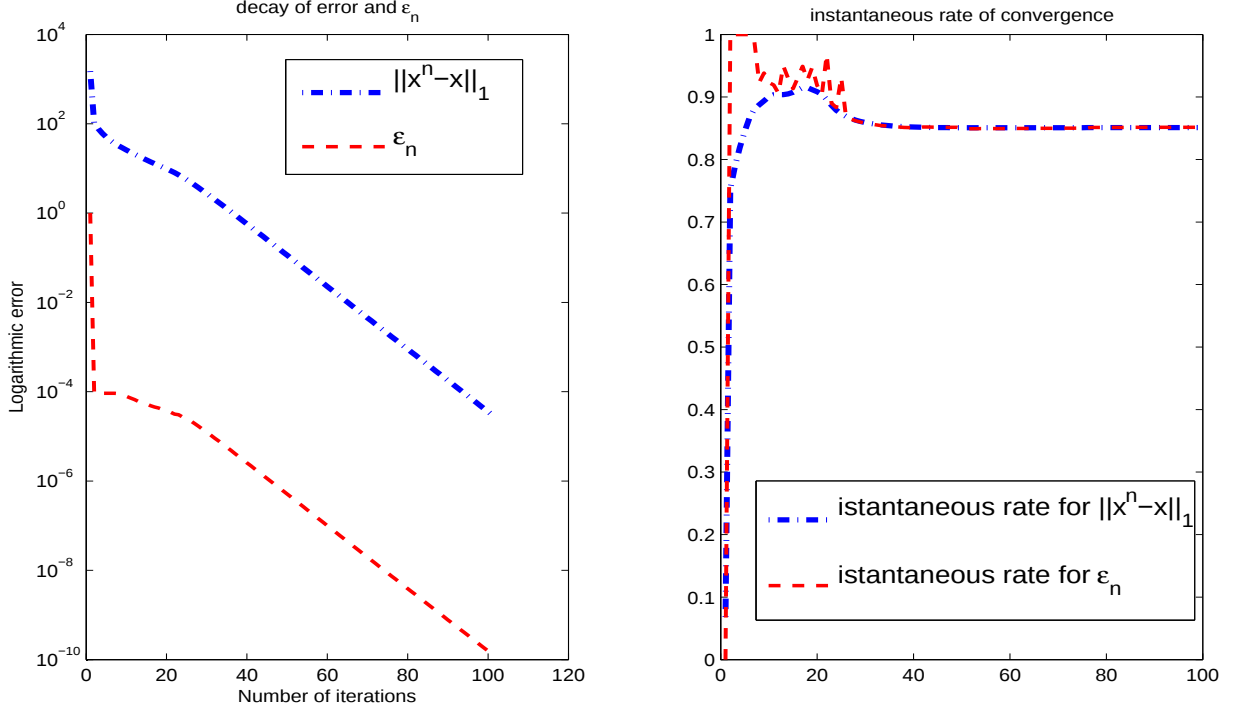


Figure 8.1: An experiment, with a matrix Φ of size 250×1500 with Gaussian $N(0, \frac{1}{250})$ i.i.d. entries, in which recovery is sought of the 45-sparse vector x^* represented in Figure 8.2 from its image $y = \Phi x$. Left: plot of $\log_{10}(\|x^n - x^*\|_{\ell_1})$ as a function of n , where the x^n are generated by Algorithm 1, with ϵ_n defined adaptively, as in (1.7). Note that the scale in the ordinate axis does not report the logarithm $0, -1, -2, \dots$, but the corresponding accuracies $10^0, 10^{-1}, 10^{-2}, \dots$ for $\|x^n - x^*\|_{\ell_1}$. The graph also plots ϵ_n as a function of n . Right: plot of the ratios $\|x^n - x^{n+1}\|_{\ell_1} / \|x^n - x^{n-1}\|_{\ell_1}$, and $(\epsilon_n - \epsilon_{n+1}) / (\epsilon_{n-1} - \epsilon_n)$ for the same examples.

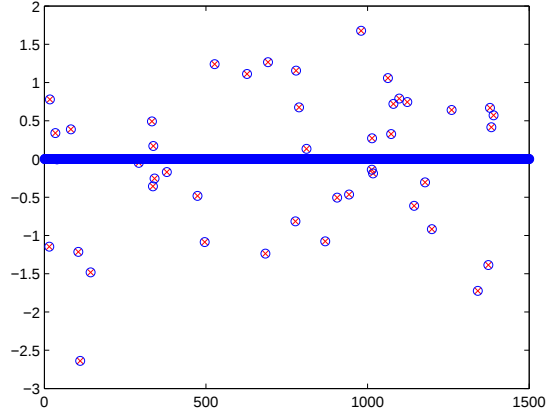


Figure 8.2: The sparse vector used in the example illustrated in Figure 8.1. This vector has dimension 1500, but only 45 non-zero entries.

with iteratively re-weighted ℓ_2 -minimization algorithms, Chartrand and several collaborators observed a similar significant improvement [11, 12, 13, 14, 36]; see in particular [13, Section 4]; we also illustrate this in Figure 8.4. It is to be noted that IRLS algorithms are computationally much less demanding than weighted ℓ_1 -minimization. In addition, there is, as far as we know, no analysis (as yet) for re-weighted ℓ_1 -minimization that is comparable to the detailed theoretical analysis of convergence presented here of our IRLS algorithm, which seems to give a realistic picture of the numerical computations.

9 Acknowledgments

We would like to thank Yu Chen, Michael Overton for various conversations on the topic of this paper, and Rachel Ward for pointing out an improvement of Theorem 5.3.

Ingrid Daubechies gratefully acknowledges partial support by NSF grants DMS-0504924 and DMS-0530865. Ronald DeVore thanks the Courant Institute for supporting an academic year visit when part of this work was done. He also gratefully acknowledges partial support by Office of Naval Research Contracts ONR-N00014-03-1-0051, ONR/DEPSCoR N00014-03-1-0675 and ONR/DEPSCoR N00014-00-1-0470; the Army Research Office Contract DAAD 19-02-1-0028; and the NSF contracts DMS-0221642 and DMS-0200187. Massimo Fornasier acknowledges the financial support provided by the European Union via the Individual Marie Curie fellowship MOIF-CT-2006-039438, and he thanks the Program in Applied and Computational Mathematics at Princeton University for its hospitality during the preparation of this work. Sinan Güntürk has been supported in part by the National Science Foundation Grant CCF-0515187, an Alfred P. Sloan Research Fellowship, and an NYU Goddard Fellowship.

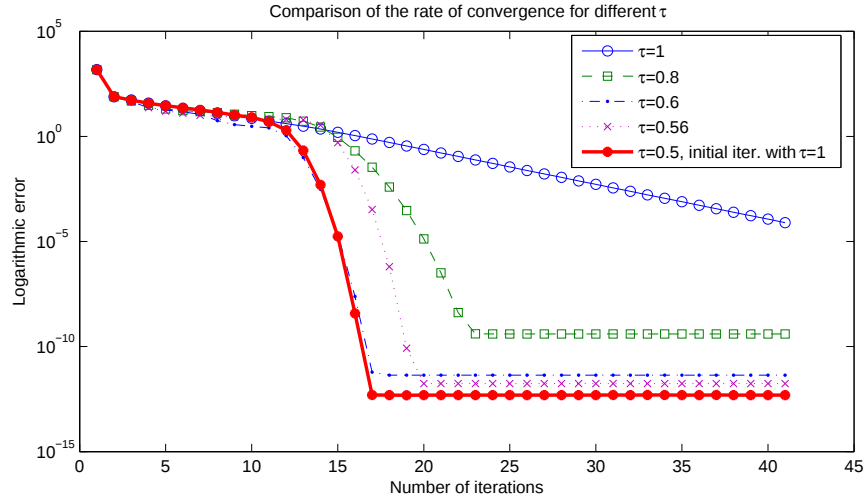


Figure 8.3: We show the decay of logarithmic error, as a function of the number of iterations of the algorithm for different values of τ (1, 0.8, 0.6, 0.56). We show also the results of an experiment in which the initial 10 iterations are performed with $\tau = 1$ and the remaining iterations with $\tau = 0.5$.

References

- [1] R. Baraniuk, *Compressive sensing*, IEEE Signal Processing Magazine, (2007), vol. 24 no. 4, 118–121.
- [2] R. Baraniuk, M. Davenport, R. DeVore, and M. Wakin, *A simple proof of the restricted isometry property for random matrices*, Constr. Approx., to appear
- [3] E. Candès, J. Romberg, and T. Tao, *Stable signal recovery from incomplete and inaccurate measurements*, Comm. Pure Appl. Math. **59** (2006), no. 8, 1207–1223.
- [4] E. Candès and T. Tao, *Near optimal signal recovery from random projections: universal encoding strategies?*, IEEE Trans. Inform. Theory **52** (2006), no. 12, 5406–5425.
- [5] E. J. Candès, *Compressive sampling*, International Congress of Mathematicians. Vol. III, Eur. Math. Soc., Zürich, 2006, pp. 1433–1452.
- [6] E. J. Candès, J. Romberg, and T. Tao, *Exact signal reconstruction from highly incomplete frequency information*, IEEE Trans. Inf. Theory **52** (2006), no. 2, 489–509.
- [7] E. J. Candès and T. Tao, *Decoding by linear programming*, IEEE Trans. Inform. Theory **51** (2005), no. 12, 4203–4215.
- [8] E. J. Candès, M. Wakin, and S. Boyd, *Enhancing sparsity by reweighted ℓ_1 minimization*, (Technical Report, California Institute of Technology).

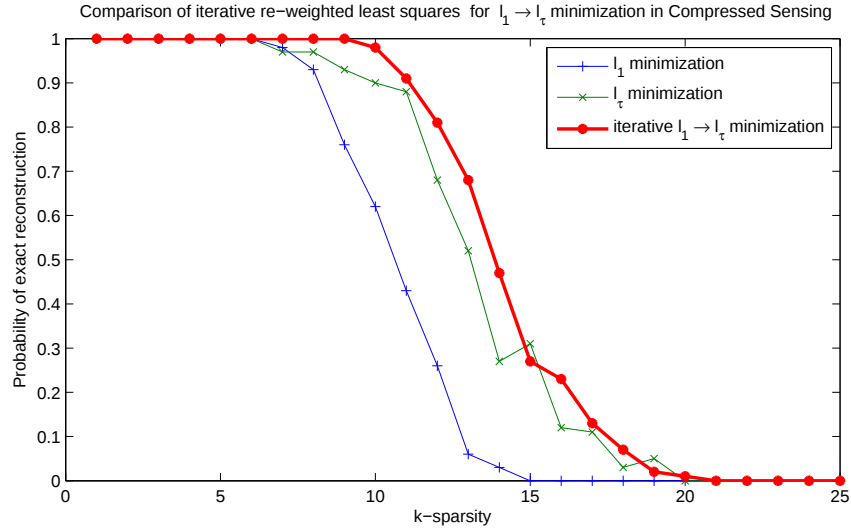


Figure 8.4: The (experimentally determined) probability of successful recovery of a sparse 250-dimensional vector x , with sparsity k , from its image $y = \Phi x$, as a function of k . In these experiments the matrix Φ is 50×250 dimensional, with i.i.d. Gaussian $N(0, \frac{1}{50})$ entries. The matrix is generated once; then, for each sparsity value k shown in the plot, 500 attempts were made, for randomly generated k -sparse vectors x . Two different IRLS algorithms were compared, one with weights inspired by ℓ_1 -minimization and the other with weights that gradually moved from an ℓ_1 - to an ℓ_τ -minimization goal, with final $\tau = 0.5$. We refer to [13, Section 4] for similar experiments (for different values of τ), although realized there by fixing the sparsity *a priori* and randomly generating matrices with an increasing number of measurements m .

- [9] E.J. Candès, *Lecture notes of the IMA course on Compressed Sensing*, (2007), see <http://www.ima.umn.edu/2006-2007/ND6.4-15.07/abstracts.html>.
- [10] A. Chambolle and P.-L. Lions, *Image recovery via total variation minimization and related problems.*, Numer. Math. **76** (1997), no. 2, 167–188.
- [11] R. Chartrand, *Exact reconstructions of sparse signals via nonconvex minimization*, IEEE Signal Process. Lett. **14** (2007), 707–710.
- [12] ———, *Nonconvex compressive sensing and reconstruction of gradient-sparse images: random vs. tomographic Fourier sampling*, preprint (2008).
- [13] R. Chartrand and V. Staneva, *Restricted isometry properties and nonconvex compressive sensing*, preprint (2008).
- [14] R. Chartrand and W. Yin, *Iteratively reweighted algorithms for compressive sensing*, Proceedings of the 33rd International Conference on Acoustics, Speech, and Signal Processing (ICASSP).

- [15] S. S. Chen, D. L. Donoho, and M. A. Saunders, *Atomic decomposition by basis pursuit*, SIAM J. Sci. Comput. **1** (1998), 33–61.
- [16] J. F. Claerbout and F. Muir, *Robust modeling with erratic data*, Geophysics **38** (1973), no. 5, 826–844.
- [17] A. K. Cline, *Rate of convergence of Lawson’s algorithm*, Math. Comp. **26** (1972), 167–176.
- [18] A. Cohen, W. Dahmen, and R. DeVore, *Compressed sensing and best k -term approximation*, J. Amer. Math. Soc., to appear.
- [19] R. DeVore, *Nonlinear approximation*, Acta Numerica **7** (1998), 51–150.
- [20] D. L. Donoho, *Compressed sensing*, IEEE Trans. Inf. Theory **52** (2006), no. 4, 1289–1306.
- [21] D. L. Donoho, *High-dimensional centrally symmetric polytopes with neighborliness proportional to dimension*, Discrete Comput. Geom. **35** (2006), no. 4, 617–652.
- [22] D. L. Donoho and B. F. Logan, *Signal recovery and the large sieve*, SIAM J. Appl. Math. **52** (1992), no. 2, 557–591.
- [23] D. L. Donoho and P. B. Stark, *Uncertainty principles and signal recovery*, SIAM J. Appl. Math. **49** (1989), no. 3, 906–931.
- [24] D. L. Donoho and J. Tanner, *Sparse nonnegative solutions of underdetermined linear equations by linear programming*, Proc. Nat. Acad. Sci. **102** (2005), no. 27, 9446–9451.
- [25] ———, *Counting faces of randomly-projected polytopes when the projection radically lowers dimension*, preprint (2006).
- [26] D. L. Donoho and Y. Tsaig, *Fast solution of ℓ_1 -norm minimization problems when the solution may be sparse*, preprint (2006).
- [27] M. Fornasier and R. March, *Restoration of color images by vector valued BV functions and variational calculus*, SIAM J. Appl. Math. **68** (2007), no. 2, 437–460.
- [28] I.F. Gorodnitsky and B.D. Rao. *Sparse signal reconstruction from limited data using FOCUSS: a recursive weighted norm minimization algorithm*, IEEE Transactions on Signal Processing, **45** (1997), no. 3, 600–616.
- [29] R. Gribonval and M. Nielsen, *Highly sparse representations from dictionaries are unique and independent of the sparseness measure*, Appl. Comput. Harmon. Anal. **22** (2007), no. 3, 335–355.
- [30] C. L. Lawson, *Contributions to the Theory of Linear Least Maximum Approximation*, Ph.D. thesis, University of California, Los Angeles, 1961.

- [31] Y. Li, *A globally convergent method for l_p problems*, SIAM J. Optim. **3** (1993), no. 3, 609–629.
- [32] M.S. O'Brien, A.N. Sinclair and S.M. Kramer, *Recovery of asparse spike time series by ℓ_1 norm deconvolution*, IEEE Trans. Signal Processing, **42** (1994), 3353–3365.
- [33] M. R. Osborne, *Finite Algorithms in Optimization and Data Analysis*, John Wiley & Sons Ltd., Chichester, 1985.
- [34] A. M. Pinkus, *On L_1 -Approximation*, Cambridge Tracts in Mathematics, vol. 93, Cambridge University Press, Cambridge, 1989.
- [35] M. Rudelson and R. Vershynin, *On sparse reconstruction from Fourier and Gaussian measurements*, Communications on Pure and Applied Mathematics, to appear.
- [36] R. Saab, R. Chartrand, and O. Yilmaz, *Stable sparse approximations via nonconvex optimization*, Proceedings of the 33rd International Conference on Acoustics, Speech, and Signal Processing (ICASSP).
- [37] F. Santosa and W. W. Symes, *Linear inversion of band-limited reflection seismograms*, SIAM J. Sci. Stat. Comput. **7** (1986), no. 4, 1307–1330.
- [38] H. L. Taylor, S. C. Banks, and J. F. McCoy, *Deconvolution with the ℓ_1 norm*, Geophysics **44** (1979), no. 1, 39–52.
- [39] R. Tibshirani, *Regression shrinkage and selection via the lasso*, J. Roy. Statist. Soc. Ser. B **58** (1996), no. 1, 267–288.