

Decoding by Linear Programming

Emmanuel J. Candes and Terence Tao

Abstract—This paper considers a natural error correcting problem with real valued input/output. We wish to recover an input vector $f \in \mathbf{R}^n$ from corrupted measurements $y = Af + e$. Here, A is an m by n (coding) matrix and e is an arbitrary and unknown vector of errors. Is it possible to recover f exactly from the data y ?

We prove that under suitable conditions on the coding matrix A , the input f is the unique solution to the ℓ_1 -minimization problem ($\|x\|_{\ell_1} := \sum_i |x_i|$)

$$\min_{g \in \mathbf{R}^n} \|y - Ag\|_{\ell_1}$$

provided that the support of the vector of errors is not too large, $\|e\|_{\ell_0} := |\{i : e_i \neq 0\}| \leq \rho \cdot m$ for some $\rho > 0$. In short, f can be recovered exactly by solving a simple convex optimization problem (which one can recast as a linear program). In addition, numerical experiments suggest that this recovery procedure works unreasonably well; f is recovered exactly even in situations where a significant fraction of the output is corrupted.

This work is related to the problem of finding sparse solutions to vastly underdetermined systems of linear equations. There are also significant connections with the problem of recovering signals from highly incomplete measurements. In fact, the results introduced in this paper improve on our earlier work. Finally, underlying the success of ℓ_1 is a crucial property we call the uniform uncertainty principle that we shall describe in detail.

Index Terms—Basis pursuit, decoding of (random) linear codes, duality in optimization, Gaussian random matrices, ℓ_1 minimization, linear codes, linear programming, principal angles, restricted orthonormality, singular values of random matrices, sparse solutions to underdetermined systems.

I. INTRODUCTION

A. Decoding of Linear Codes

THIS paper considers the model problem of recovering an input vector $f \in \mathbf{R}^n$ from corrupted measurements $y = Af + e$. Here, A is an m by n matrix (we will assume throughout the paper that $m > n$), and e is an arbitrary and unknown vector of errors. The problem we consider is whether it is possible to recover f exactly from the data y . And if so, how?

Our problem has of course the flavor of error correction problems which arise in coding theory as we may think of A as a

linear code; a linear code is a given collection of codewords which are vectors $a_1, \dots, a_n \in \mathbf{R}^m$ —the columns of the matrix A . We would like to emphasize, however, that there is a clear distinction between our real-valued setting and the finite alphabet one which is more common in the information theory literature. Given a vector $f \in \mathbf{R}^n$ (the “plaintext”) we can then generate a vector Af in $\mathbf{R}^m$ (the “ciphertext”); if A has full rank, then one can clearly recover the plaintext f from the ciphertext Af . But now we suppose that the ciphertext Af is corrupted by an arbitrary vector $e \in \mathbf{R}^m$ giving rise to the corrupted ciphertext $Af + e$. The question is then: given the coding matrix A and $Af + e$, can one recover f exactly?

As is well known, if the fraction of the corrupted entries is too large, then of course we have no hope of reconstructing f from $Af + e$; for instance, assume $m = 2n$ and consider two distinct plaintexts f, f' and form a vector $g \in \mathbf{R}^m$ by concatenating the first half of Af together with the second half of Af' . Then $g = Af + e = Af' + e'$ where both e and e' are supported on sets of size at most $n = m/2$. This simple example shows that accurate decoding is impossible when the size of the support of the error vector is greater or equal to a half of that of the output Af . Therefore, a common assumption in the literature is to assume that only a small fraction of the entries are actually damaged

$$\|e\|_{\ell_0} := |\{i : e_i \neq 0\}| \leq \rho \cdot m. \quad (1.1)$$

For which values of ρ can we hope to reconstruct e with practical algorithms? That is, with algorithms whose complexity is at most polynomial in the length m of the code A ?

To reconstruct f , note that it is obviously sufficient to reconstruct the vector e since knowledge of $Af + e$ together with e gives Af , and consequently f , since A has full rank. Our approach is then as follows. We construct a matrix which annihilates the $m \times n$ matrix A on the left, i.e., such that $FA = 0$. This can be done in an obvious fashion by taking a matrix F whose kernel is the range of A in $\mathbf{R}^m$, which is an n -dimensional subspace (e.g., F could be the orthogonal projection onto the cokernel of A). We then apply F to the output $y = Af + e$ and obtain

$$\tilde{y} = F(Af + e) = Fe \quad (1.2)$$

since $FA = 0$. Therefore, the decoding problem is reduced to that of reconstructing a *sparse* vector e from the observations Fe (by sparse, we mean that only a fraction of the entries of e are nonzero). Therefore, the overarching theme is that of the sparse reconstruction problem, which has recently attracted a lot of attention as we are about to see.

Manuscript received February 2005; revised September 2, 2005. The work of E. J. Candes is supported in part by the National Science Foundation under Grants DMS 01-40698 (FRG) and ACI-0204932 (ITR), and by an A. P. Sloan Fellowship. The work of T. Tao is supported in part by a grant from the Packard Foundation.

E. J. Candes is with the Department of Applied and Computational Mathematics, California Institute of Technology, Pasadena, CA 91125 USA (e-mail: emmanuel@acm.caltech.edu).

T. Tao is with the Department of Mathematics, University of California, Los Angeles, CA 90095 USA (e-mail: tao@math.ucla.edu).

Communicated by M. P. Fossorier, Associate Editor for Coding Techniques. Digital Object Identifier 10.1109/TIT.2005.858979

B. Sparse Solutions to Underdetermined Systems

Finding sparse solutions to underdetermined systems of linear equations is in general NP-hard. For example, the sparsest solution is given by

$$(P_0) \quad \min_{d \in \mathbf{R}^m} \|d\|_{\ell_0} \quad \text{subject to} \quad Fd = \tilde{y} (= Fe) \quad (1.3)$$

and to the best of our knowledge, solving this problem essentially requires exhaustive searches over all subsets of columns of F , a procedure which clearly is combinatorial in nature and has exponential complexity.

Formally, given an integer $p \times m$ matrix F and an integer vector b , the problem of deciding whether there is a vector with rational entries such that $Fd = b$, and with fewer than a fixed number of nonzero entries is NP-complete, see [2, Problem MP5]. In fact, the ℓ_0 -minimization problem contains the subset-sum problem. Assume, for instance, that $m = 2k$ and $p = k + 1$, and consider the following set of $2k$ vectors in $\mathbf{R}^{k+1}$:

$$\begin{array}{cccc} e_1 & e_2 & \dots & e_k \\ e_1 + a_1 e_{k+1} & e_2 + a_2 e_{k+1} & \dots & e_k + a_k e_{k+1} \end{array}$$

where $e_1, \dots, e_{k+1}$ are the usual basis vectors and $a_1, \dots, a_k$ are integers. Now let α be another integer and consider the problem of deciding whether

$$e_1 + \dots + e_k + \alpha e_{k+1}$$

is a k -sparse linear combination of the above $2k$ vectors ($k - 1$ -sparse is impossible). This is exactly the subset-sum problem, i.e., whether one can write α as a sum of a subset of $a_1, \dots, a_k$. It is well known that this is an NP-complete problem.

This computational intractability has recently led researchers to develop alternatives to (P_0) , and a frequently discussed approach considers a similar program in the ℓ_1 -norm which goes by the name of *basis pursuit* [3]

$$(P_1) \quad \min_{x \in \mathbf{R}^m} \|d\|_{\ell_1}, \quad Fd = \tilde{y} \quad (1.4)$$

where we recall that $\|d\|_{\ell_1} = \sum_{i=1}^m |d_i|$. Unlike the ℓ_0 -norm which enumerates the nonzero coordinates, the ℓ_1 -norm is convex. It is also well known [4] that (P_1) can be recast as a linear program (LP).

Motivated by the problem of finding sparse decompositions of special signals in the field of mathematical signal processing and following upon the ground breaking work of Donoho and Huo [5], a series of beautiful articles [6]–[9] showed exact equivalence between the two programs (P_0) and (P_1) . In a nutshell, this work shows that for $m/2$ by m matrices F obtained by concatenation of two orthonormal bases, the solution to both (P_0) and (P_1) are unique and identical provided that in the most favorable case, the vector e has at most $0.914\sqrt{m/2}$ nonzero entries. This is of little practical use here since we are interested in procedures that might recover a signal when a constant fraction of the output is unreliable.

Using very different ideas and together with Romberg [10], the authors proved that the equivalence holds with overwhelming probability for various types of random matrices

provided that the number of nonzero entries in the vector e be of the order of $m/\log m$ [1], [11]. In the special case where F is an $m/2$ by m random matrix with independent standard normal entries, [12] proved that the number of nonzero entries may be as large as $\rho \cdot m$, where $\rho > 0$ is some very small and unspecified positive constant independent of m .

C. Innovations

This paper introduces the concept of a restrictedly almost orthonormal system—a collection of vectors which behaves like an almost orthonormal system but only for sparse linear combinations. Thinking about these vectors as the columns of the matrix F , we show that this condition allows for the exact reconstruction of sparse linear combination of these vectors, i.e., e . Our results are significantly different than those mentioned above as they are deterministic and do not involve any kind of randomization, although they can of course be specialized to random matrices. For instance, we shall see that a Gaussian matrix with independent entries sampled from the standard normal distribution is restrictedly almost orthonormal with overwhelming probability, and that minimizing the ℓ_1 -norm recovers sparse decompositions with a number of nonzero entries of size $\rho_0 \cdot m$; we shall actually give numerical values for ρ_0 .

We presented the connection with sparse solutions to underdetermined systems of linear equations merely for pedagogical reasons. There is a more direct approach. To recover f from corrupted data $y = Af + e$, we consider solving the following ℓ_1 -minimization problem:

$$(P'_1) \quad \min_{g \in \mathbf{R}^n} \|y - Ag\|_{\ell_1}. \quad (1.5)$$

Now f is the unique solution of (P'_1) if and only if e is the unique solution of (P_1) . In other words, (P_1) and (P'_1) are equivalent programs. To see why this is true, observe on the one hand that since $y = Af + e$, we may decompose g as $g = f + h$ so that

$$(P'_1) \Leftrightarrow \min_{h \in \mathbf{R}^n} \|e - Ah\|_{\ell_1}.$$

On the other hand, the constraint $Fx = Fe$ means that $x = e - Ah$ for some $h \in \mathbf{R}^n$ and, therefore,

$$\begin{aligned} (P_1) &\Leftrightarrow \min_{h \in \mathbf{R}^n} \|x\|_{\ell_1}, \quad x = e - Ah \\ &\Leftrightarrow \min_{h \in \mathbf{R}^n} \|e - Ah\|_{\ell_1} \end{aligned}$$

which proves the claim.

The program (P'_1) may also be re-expressed as an LP—hence, the title of this paper. Indeed, the ℓ_1 -minimization problem is equivalent to

$$\min 1^T t, \quad -t \leq y - Ag \leq t \quad (1.6)$$

where the optimization variables are $t \in \mathbf{R}^m$ and $g \in \mathbf{R}^n$ (as is standard, the generalized vector inequality $x \leq y$ means that $x_i \leq y_i$ for all i). As a result, (P'_1) is an LP with inequality constraints and can be solved efficiently using standard optimization algorithms, see [13].

D. Restricted Isometries

In the remainder of this paper, it will be convenient to use some linear algebra notations. We denote by $(v_j)_{j \in J} \in \mathbb{R}^p$ the columns of the matrix F and by H the linear subspace spanned by these vectors. Further, for any $T \subseteq J$, we let F_T be the submatrix with column indices $j \in T$ so that

$$F_T c = \sum_{j \in T} c_j v_j \in H.$$

To introduce the notion of almost orthonormal system, we first observe that if the columns of F are sufficiently “degenerate,” the recovery problem cannot be solved. In particular, if there exists a nontrivial sparse linear combination $\sum_{j \in T} c_j v_j = 0$ of the v_j which sums to zero, and $T = T_1 \cup T_2$ is any partition of T into two disjoint sets, then the vector y

$$y := \sum_{j \in T_1} c_j v_j = \sum_{j \in T_2} (-c_j) v_j$$

has two distinct sparse representations. On the other hand, linear dependencies $\sum_{j \in J} c_j v_j = 0$ which involve a large number of nonzero coefficients c_j , as opposed to a sparse set of coefficients, do not present an obvious obstruction to sparse recovery. At the other extreme, if the $(v_j)_{j \in J}$ are an orthonormal system, then the recovery problem is easily solved by setting $c_j = \langle f, v_j \rangle$.

The main result of this paper is that if we impose a “restricted orthonormality hypothesis,” which is far weaker than assuming orthonormality, then (P_1) solves the recovery problem, even if the $(v_j)_{j \in J}$ are highly linearly dependent (for instance, it is possible for $m := |J|$ to be much larger than the dimension of the span of the v_j 's). To make this quantitative we introduce the following definition.

Definition 1.1 (Restricted Isometry Constants): Let F be the matrix with the finite collection of vectors $(v_j)_{j \in J} \in \mathbb{R}^p$ as columns. For every integer $1 \leq S \leq |J|$, we define the S -restricted isometry constants δ_S to be the smallest quantity such that F_T obeys

$$(1 - \delta_S) \|c\|^2 \leq \|F_T c\|^2 \leq (1 + \delta_S) \|c\|^2 \quad (1.7)$$

for all subsets $T \subseteq J$ of cardinality at most S , and all real coefficients $(c_j)_{j \in T}$. Similarly, we define the S, S' -restricted orthogonality constants $\theta_{S, S'}$ for $S + S' \leq |J|$ to be the smallest quantity such that

$$|\langle F_T c, F_{T'} c' \rangle| \leq \theta_{S, S'} \cdot \|c\| \|c'\| \quad (1.8)$$

holds for all disjoint sets $T, T' \subseteq J$ of cardinality $|T| \leq S$ and $|T'| \leq S'$.

The numbers δ_S and $\theta_{S, S'}$ measure how close the vectors v_j are to behaving like an orthonormal system, but only when restricting attention to sparse linear combinations involving no more than S vectors. These numbers are clearly nondecreasing

in S, S' . For $S = 1$, the value δ_1 only conveys magnitude information about the vectors v_j ; indeed δ_1 is the best constant such that

$$1 - \delta_1 \leq \|v_j\|^2 \leq 1 + \delta_1, \quad \text{for all } j \in J. \quad (1.9)$$

In particular, $\delta_1 = 0$ if and only if all of the v_j 's have unit length. Section II-C establishes that the higher δ_S control the orthogonality numbers $\theta_{S, S'}$.

Lemma 1.1: We have $\theta_{S, S'} \leq \delta_{S+S'} \leq \theta_{S, S'} + \max(\delta_S, \delta_{S'})$ for all S, S' .

To see the relevance of the restricted isometry numbers δ_S to the sparse recovery problem, consider the following simple observation.

Lemma 1.2: Suppose that $S \geq 1$ is such that $\delta_{2S} < 1$, and let $T \subseteq J$ be such that $|T| \leq S$. Let $f := Fc$ where c is an arbitrary vector supported on T . Then c is the unique minimizer to

$$(P_0) \quad \min \|d\|_{\ell_0}, \quad Fd = f$$

so that c can be reconstructed from knowledge of the vector f (and the v_j 's).

Proof: We prove that there is a unique c with $\|c\|_{\ell_0} \leq S$ and obeying $f = \sum_{j \in J} c_j v_j$. Suppose for the sake of contradiction that f had two distinct sparse representations $f = Fc = Fc'$ where c and c' were supported on sets obeying $|T|, |T'| \leq S$. Then

$$F(c - c') = 0.$$

By construction $c - c'$ is supported on $T \cup T'$ of size less or equal to $2S$. Taking norms of both sides and applying (1.7) and the hypothesis $\delta_{2S} < 1$ we conclude that $\|c - c'\|^2 = 0$, contradicting the hypothesis that the two representations were distinct. $\square$

E. Main Results

Note that the previous lemma is an abstract existence argument which shows what might theoretically be possible, but does not supply any efficient algorithm to recover T and c_j from f and $(v_j)_{j \in J}$ other than by brute-force search—as discussed earlier. In contrast, our main theorem shows that, by imposing slightly stronger conditions on δ_{2S} , the ℓ_1 -minimization program (P_1) recovers f exactly.

Theorem 1.3: Suppose that $S \geq 1$ is such that

$$\delta_S + \theta_{S, S} + \theta_{S, 2S} < 1 \quad (1.10)$$

and let c be a real vector supported on a set $T \subseteq J$ obeying $|T| \leq S$. Put $f := Fc$. Then c is the unique minimizer to

$$(P_1) \quad \min \|d\|_{\ell_1}, \quad Fd = f.$$

Note from Lemma 1.1 that (1.10) implies $\delta_{2S} < 1$, and is in turn implied by $\delta_S + \delta_{2S} + \delta_{3S} < 1$. Thus, the condition (1.10)

is roughly “three times as strict” as the condition required for Lemma 1.2.

Theorem 1.3 is inspired by our previous work [1], see also [10], [11], but unlike those earlier results, our results here are *deterministic*, and thus do not have a nonzero probability of failure, provided of course one can ensure that the system $(v_j)_{j \in J}$ verifies the condition (1.10). By virtue of the previous discussion, we have the companion result.

Theorem 1.4: Suppose F is such that $FA = 0$ and let $S \geq 1$ be a number obeying the hypothesis of Theorem 1.3. Set $y = Af + e$, where e is a real vector supported on a set of size at most S . Then f is the unique minimizer to

$$(P'_1) \quad \min_{g \in \mathbf{R}^n} \|y - Ag\|_{\ell_1}.$$

F. Gaussian Random Matrices

An important question is then to find matrices with good restricted isometry constants, i.e., such that (1.10) holds for large values of S . Indeed, such matrices will tolerate a larger fraction of output in error while still allowing exact recovery of the original input by linear programming. How to construct such matrices might be delicate. In Section III, however, we will argue that generic matrices, namely, samples from the Gaussian unitary ensemble obey (1.10) for relatively large values of S .

Theorem 1.5: Assume $p \leq m$ and let F be a p by m matrix whose entries are independent and identically distributed (i.i.d.) Gaussian with mean zero and variance $1/p$. Then the condition of Theorem 1.3 holds with overwhelming probability provided that $r = S/m$ is small enough so that

$$r < r^*(p, m)$$

where $r^*(p, m)$ is given in (3.23). (By “with overwhelming probability,” we mean with probability decaying exponentially in m .) In the limit of large samples, r^* only depends upon the ratio, and numerical evaluations show that the condition holds for $r \leq 3.6 \cdot 10^{-4}$ in the case where $p/m = 3/4$, $r \leq 3.2 \cdot 10^{-4}$ when $p/m = 2/3$, and $r \leq 2.3 \cdot 10^{-4}$ when $p/m = 1/2$.

In other words, Gaussian matrices are a class of matrices for which one can solve an underdetermined systems of linear equations by minimizing ℓ_1 provided, of course, the input vector has fewer than $\rho \cdot m$ nonzero entries with $\rho > 0$. We mentioned earlier that this result is similar to [12]. What is new here is that by using a very different machinery, one can obtain explicit numerical values which were not available before.

In the context of error correcting, the consequence is that a fraction of the output may be corrupted by *arbitrary* errors and yet, solving a convex problem would still recover f exactly—a rather unexpected feat.

Corollary 1.6: Suppose A is an n by m Gaussian matrix and set $p = m - n$. Under the hypotheses of Theorem 1.5, the solution to (P'_1) is unique and equal to f .

This is an immediate consequence of Theorem 1.5. The only thing we need to argue is why we may think of the annihilator

F (such that $FA = 0$) as a matrix with independent Gaussian entries. Observe that the range of A is a random space of dimension n embedded in $\mathbf{R}^m$ so that the data $\tilde{y} = Fe$ is the projection of e on a random space of dimension p . The range of a p by m matrix with independent Gaussian entries precisely is a random subspace of dimension p , which justifies the claim.

We would like to point out that the numerical bounds we derived in this paper are overly pessimistic. We are confident that finer arguments and perhaps new ideas will allow to derive versions of Theorem 1.5 with better bounds. The discussion section will enumerate several possibilities for improvement.

G. A Notion of Capacity

We now develop a notion of “capacity,” that is, an upper bound on the support size of the error vector beyond which recovery is information-theoretically impossible. Define the *cospark* of the matrix A as

$$\text{cospark}(A) := \min_{h \in \mathbf{R}^n: h \neq 0} \|Ah\|_{\ell_0}. \quad (1.11)$$

We call this number the *cospark* because it is related to another quantity others have called the *spark*. The spark of a matrix F is the minimal number of columns from F that are linearly dependent [6], [7]. In other words

$$\text{spark}(F) := \min_{x \neq 0} \|x\|_{\ell_0}, \quad \text{subject to } Fx = 0.$$

Suppose A has full rank and assume that F is any full rank matrix such that $FA = 0$, then $\text{cospark}(A) = \text{spark}(F)$, which simply follows from

$$\min_{x \neq 0} \|x\|_{\ell_0}, \quad \text{s. t. } Fx = 0 \quad \Leftrightarrow \quad \min_{h \neq 0} \|Ah\|_{\ell_0}$$

since $Fx = 0$ means that x is in the range of A and, therefore, of the form Ah where h is nonzero since x does not vanish.

Lemma 1.7: Suppose the fraction of errors ρ obeys

$$\rho < \frac{\text{cospark}(A)}{2m}; \quad (1.12)$$

then, exact decoding is achieved by minimizing $\|y - Ag\|_{\ell_0}$. Conversely, if $\rho \geq \text{cospark}(A)/2m$, then accurate decoding is not possible.

Proof: Suppose first that the number of errors obeys $\|e\|_{\ell_0} < \text{cospark}(A)/2$ and let F be of full rank with $FA = 0$. Since minimizing $\|y - Ag\|_{\ell_0}$ is equivalent to finding the shortest d so that $Fd = Fe$, and that the proof of Lemma 1.2 assures that e is the unique vector with at most $\text{cospark}(A)/2$ nonzero entries, the recovery is exact.

The converse is immediate. Let $h \neq 0$ be a vector such that $\|Ah\|_{\ell_0} = \text{cospark}(A)$ and write $Ah := e - e'$ where both $\|e\|_{\ell_0}$ and $\|e'\|_{\ell_0}$ are less or equal to $\text{cospark}(A)/2$. Let f be any vector in $\mathbf{R}^n$ and set $f' = f + h$. Then

$$y = Af' + e' = Af + e$$

and, therefore, f and f' are indistinguishable. $\square$

Corollary 1.8: Suppose the entries of A are independent normal with mean 0 and variance 1. Then with probability 1, $\text{cospark}(A) = m - n + 1$. As a consequence, one can

theoretically recover any plaintext if and only if the fraction error obeys

$$\rho \leq \frac{1 - n/m}{2}. \quad (1.13)$$

In other words, given a fraction of error ρ , accurate decoding is possible provided that the “rate” n/m obeys $n/m \leq 2\rho - 1$. This may help to establish a link between this work and other information literature.

Proof: It is clear that $\text{cospark}(A) \leq m - n + 1$ as one can find an $h \neq 0$ so that the first $n - 1$ entries of Ah vanish. Further, all the n by n submatrices of A are invertible with probability 1. This follows from the fact that each n by n submatrix is invertible with probability 1, and that there is a finite number of them. With probability one, therefore, it is impossible to find $h \neq 0$ such that Ah has n vanishing entries. This gives $\text{cospark}(A) = m - n + 1$. $\square$

H. Organization of the Paper

The paper is organized as follows. Section II proves our main claim, namely, Theorem 1.3 (and hence Theorem 1.4) while Section III introduces elements from random matrix theory to establish Theorem 1.5. In Section IV, we present numerical experiments which suggest that in practice, (P'_1) works unreasonably well and recovers the f exactly from $y = Af + e$ provided that the fraction of the corrupted entries be less than about 17% in the case where $m = 2n$ and less than about 34% in the case where $m = 4n$. Section V explores the consequences of our results for the recovery of signals from highly incomplete data and ties our findings with some of our earlier work. Finally, we conclude with a short discussion section whose main purpose is to outline areas for improvement.

II. PROOF OF MAIN RESULTS

Our main result, namely, Theorem 1.3 is proved by duality. As we will see in Section II-A, c is the unique minimizer if the matrix F_T has full rank and if one can find a vector w with the two properties

- i) $\langle w, v_j \rangle = \text{sgn}(c_j)$, for all $j \in T$;
- ii) and $|\langle w, v_j \rangle| < 1$, for all $j \notin T$;

where $\text{sgn}(c_j)$ is the sign of c_j ($\text{sgn}(c_j) = 0$, for $c_j = 0$). The two conditions above say that a specific dual program is feasible and is called the exact reconstruction property in [1], see also [10]. The reader might find it helpful to see how these conditions arise and we first informally sketch the argument (a rigorous proof follows in Section II-A). Suppose we wish to minimize $J(d)$ subject to $Fd = z$, and that J is differentiable. Together with $Fd = z$, the classical Lagrange multiplier optimality condition (see [13]) asserts that there exists w (w is a Lagrange multiplier) such that

$$\nabla J(d) - F^T w = 0.$$

In a nutshell, for c to be a solution, we must have $F^T w = \nabla J(c)$. In our case $J(c) = \sum_j |c_j|$, and so $(\nabla J(c))_j = \text{sgn}(c_j)$ for $c_j \neq 0$, but J is otherwise not smooth at zero. In this case,

we ask that $F^T w$ be a subgradient of J at the point c [13], which here means

$$(F^T w)_j = \text{sgn}(c_j), \text{ if } c_j \neq 0, \text{ and } |(F^T w)_j| \leq 1, \text{ otherwise.}$$

Since by definition $(F^T w)_j = \langle w, v_j \rangle$, conditions i) and ii) are now natural.

A. Proof of Theorem 1.3

Observe first that standard convex arguments give that there exists at least one minimizer $d = (d_j)_{j \in J}$ to the problem (P_1) . We need to prove that $d = c$. Since c obeys the constraints of this problem, d obeys

$$\|d\|_{\ell_1} \leq \|c\|_{\ell_1} = \sum_{j \in T} |c_j|. \quad (2.14)$$

Now take a w obeying properties i) and ii) (see the remark following Lemma 2.2). Using the fact that the inner product $\langle w, v_j \rangle$ is equal to the sign of c on T and has absolute value strictly less than one on the complement, we then compute

$$\begin{aligned} \|d\|_{\ell_1} &= \sum_{j \in T} |c_j + (d_j - c_j)| + \sum_{j \notin T} |d_j| \\ &\geq \sum_{j \in T} \text{sgn}(c_j)(c_j + (d_j - c_j)) + \sum_{j \notin T} d_j \langle w, v_j \rangle \\ &= \sum_{j \in T} |c_j| + \sum_{j \in T} (d_j - c_j) \langle w, v_j \rangle + \sum_{j \notin T} d_j \langle w, v_j \rangle \\ &= \sum_{j \in T} |c_j| + \langle w, \sum_{j \in J} d_j v_j - \sum_{j \in T} c_j \rangle \\ &= \sum_{j \in T} |c_j| + \langle w, f - f \rangle \\ &= \sum_{j \in T} |c_j|. \end{aligned}$$

Comparing this with (2.14) we see that all the inequalities in the above computation must in fact be equality. Since $|\langle w, v_j \rangle|$ was strictly less than 1 for all $j \notin T$, this in particular forces $d_j = 0$ for all $j \notin T$. Thus,

$$\sum_{j \in T} (d_j - c_j) v_j = f - f = 0.$$

Applying (1.7) (and noting from hypothesis that $\delta_S < 1$) we conclude that $d_j = c_j$ for all $j \in T$. Thus, $d = c$ as claimed.

For $|T| \leq S$ with S obeying the hypothesis of Theorem 1.3, F_T has full rank since $\delta_S < 1$ and thus, the proof simply consists in constructing a dual vector w ; this is the object of the next section. This concludes the proof of our theorem.

B. Exact Reconstruction Property

We now examine the sparse reconstruction property and begin with Lemma 2.1 which establishes that the coefficients $\langle w, v_j \rangle$, $j \notin T$, are suitably bounded except for an exceptional set. Lemma 2.2 strengthens this results by eliminating the exceptional set.

Lemma 2.1 (Dual Reconstruction Property, ℓ_2 Version): Let $S, S' \geq 1$ be such that $\delta_S < 1$, and c be a real vector supported

on $T \subset J$ such that $|T| \leq S$. Then there exists a vector $w \in H$ such that $\langle w, v_j \rangle = c_j$ for all $j \in T$. Furthermore, there is an “exceptional set” $E \subset J$ which is disjoint from T , of size at most

$$|E| \leq S' \quad (2.15)$$

and with the properties

$$|\langle w, v_j \rangle| \leq \frac{\theta_{S,S'}}{(1 - \delta_S)\sqrt{S'}} \cdot \|c\|, \quad \text{for all } j \notin T \cup E$$

and

$$\left(\sum_{j \in E} |\langle w, v_j \rangle|^2 \right)^{1/2} \leq \frac{\theta_S}{1 - \delta_S} \cdot \|c\|$$

where $\theta_S := \theta_{S,2S}$ for short. In addition, $\|w\| \leq K \cdot \|c\|$ for some constant $K > 0$ only depending upon δ_S .

Proof: Recall that $F_T : \ell_2(T) \rightarrow H$ is the linear transformation $F_T c_T := \sum_{j \in T} c_j v_j$ where $c_T := (c_j)_{j \in T}$ (we use the subscript T in c_T to emphasize that the input is a $|T|$ -dimensional vector), and let F_T^* be the adjoint transformation

$$F_T^* w := (\langle w, v_j \rangle)_{j \in T}.$$

Property (1.7) gives

$$1 - \delta_S \leq \lambda_{\min}(F_T^* F_T) \leq \lambda_{\max}(F_T^* F_T) \leq 1 + \delta_S$$

where $\lambda_{\min}$ and $\lambda_{\max}$ are the minimum and maximum eigenvalues of the positive-definite operator $F_T^* F_T$. In particular, since $\delta_S < 1$, we see that $F_T^* F_T$ is invertible with

$$\|(F_T^* F_T)^{-1}\| \leq \frac{1}{1 - \delta_S}. \quad (2.16)$$

Also note that $\|F_T(F_T^* F_T)^{-1}\| \leq \sqrt{1 + \delta_S}/(1 - \delta_S)$ and set $w \in H$ to be the vector

$$w := F_T(F_T^* F_T)^{-1} c_T;$$

it is then clear that $F_T^* w = c_T$, i.e., $\langle w, v_j \rangle = c_j$ for all $j \in T$. In addition, $\|w\| \leq K \cdot \|c_T\|$ with $K = \sqrt{1 + \delta_S}/(1 - \delta_S)$. Finally, if T' is any set in J disjoint from T with $|T'| \leq S'$ and $d_{T'} = (d_j)_{j \in T'}$ is any sequence of real numbers, then (1.8) and (2.16) give

$$\begin{aligned} |\langle F_T^* w, d_{T'} \rangle_{\ell_2(T')}| &= |\langle w, F_{T'} d_{T'} \rangle_{\ell_2(T')}| \\ &= \left| \left\langle \sum_{j \in T} ((F_T^* F_T)^{-1} c_T)_j v_j, \sum_{j \in T'} d_j v_j \right\rangle \right| \\ &\leq \theta_{S,S'} \cdot \|(F_T^* F_T)^{-1} c_T\| \cdot \|d_{T'}\| \\ &\leq \frac{\theta_{S,S'}}{1 - \delta_S} \|c_T\| \cdot \|d_{T'}\|; \end{aligned}$$

since $d_{T'}$ was arbitrary, it follows from the variational definition of the ℓ_2 norm, $\|u\|_{\ell_2} = \sup_v |\langle u, v \rangle| / \|v\|_{\ell_2}$, that

$$\|F_T^* w\|_{\ell_2(T')} \leq \frac{\theta_{S,S'}}{1 - \delta_S} \|c_T\|.$$

In other words

$$\left(\sum_{j \in T'} |\langle w, v_j \rangle|^2 \right)^{1/2} \leq \frac{\theta_{S,S'}}{1 - \delta_S} \|c_T\| \quad (2.17)$$

whenever $T' \subset J \setminus T$ and $|T'| \leq S'$. In particular, if we set

$$E := \{j \in J \setminus T : |\langle w, v_j \rangle| > \frac{\theta_{S,S'}}{(1 - \delta_S)\sqrt{S'}} \cdot \|c_T\|\}$$

then $|E|$ must obey $|E| \leq S'$, since otherwise we could contradict (2.17) by taking a subset T' of E of cardinality S' . The claims now follow. $\square$

We now derive a solution with better control on the sup norm of $|\langle w, v_j \rangle|$ outside of T , by iterating away the exceptional set E (while keeping the values on T fixed). The proof of the following lemma is given in the Appendix.

Lemma 2.2 (Dual Reconstruction Property, ℓ_∞ Version): Let $S \geq 1$ be such that $\delta_S + \theta_{S,2S} < 1$, and c be a real vector supported on $T \subset J$ obeying $|T| \leq S$. Then there exists a vector $w \in H$ such that $\langle w, v_j \rangle = c_j$ for all $j \in T$. Furthermore, w obeys

$$|\langle w, v_j \rangle| \leq \frac{\theta_S}{(1 - \delta_S - \theta_{S,2S})\sqrt{S}} \cdot \|c\|, \quad \text{for all } j \notin T. \quad (2.18)$$

Lemma 2.2 actually solves the dual recovery problem. Indeed, our result states that one can find a vector $w \in H$ obeying both properties i) and ii) stated at the beginning of the section. To see why ii) holds, observe that $\|\text{sgn}(c)\| = \sqrt{|T|} \leq \sqrt{S}$ and, therefore, (2.18) gives for all $j \notin T$

$$|\langle w, v_j \rangle| \leq \frac{\theta_{S,S}}{(1 - \delta_S - \theta_{S,2S})} \cdot \sqrt{\frac{|T|}{S}} \leq \frac{\theta_{S,S}}{(1 - \delta_S - \theta_{S,2S})} < 1$$

provided that $\delta_S + \theta_{S,S} + \theta_{S,2S} < 1$.

It is likely that one may push the condition $\delta_S + \theta_{S,S} + \theta_{S,2S} < 1$ a little further. The key idea is as follows. Each vector w_n in the iteration scheme used to prove Lemma 2.2 was designed to annihilate the influence of w_{n-1} on the exceptional set T_{n-1} . But annihilation is too strong of a goal. It would be just as suitable to design w_n to moderate the influence of w_{n-1} enough so that the inner product with elements in T_{n-1} is small rather than zero. However, we have not pursued such refinements as the arguments would become considerably more complicated than the calculations presented here.

C. Approximate Orthogonality

Lemma 1.1 gives control of the size of the principal angle between subspaces of dimension S and S' , respectively. This is useful because it allows to guarantee exact reconstruction from the knowledge of the δ numbers only.

Proof Proof of Lemma 1.1: We first show that $\theta_{S,S'} \leq \delta_{S+S'}$. By homogeneity it will suffice to show that

$$|\langle \sum_{j \in T} c_j v_j, \sum_{j' \in T'} c'_{j'} v_{j'} \rangle| \leq \delta_{S+S'}$$

whenever $|T| \leq S$, $|T'| \leq S'$, T, T' are disjoint, and

$$\sum_{j \in T} |c_j|^2 = \sum_{j' \in T'} |c'_{j'}|^2 = 1.$$

Now (1.7) gives

$$2(1 - \delta_{S+S'}) \leq \left\| \sum_{j \in T} c_j v_j + \sum_{j' \in T'} c'_{j'} v_{j'} \right\|^2 \leq 2(1 + \delta_{S+S'})$$

together with

$$2(1 - \delta_{S+S'}) \leq \left\| \sum_{j \in T} c_j v_j - \sum_{j' \in T'} c_{j'} v_{j'} \right\|^2 \leq 2(1 + \delta_{S+S'})$$

and the claim now follows from the parallelogram identity

$$\langle f, g \rangle = \frac{\|f + g\|^2 - \|f - g\|^2}{4}.$$

It remains to show that $\delta_{S+S'} \leq \theta_S + \delta_S$. Again by homogeneity, it suffices to establish that

$$|\langle \sum_{j \in \tilde{T}} c_j v_j, \sum_{j' \in \tilde{T}} c_{j'} v_{j'} \rangle - 1| \leq (\delta_S + \theta_S)$$

whenever $|\tilde{T}| \leq S + S'$ and $\sum_{j \in \tilde{T}} |c_j|^2 = 1$. To prove this property, we partition $\tilde{T}$ as $\tilde{T} = T \cup T'$ where $|T| \leq S$ and $|T'| \leq S'$ and write $\sum_{j \in T} |c_j|^2 = \alpha$ and $\sum_{j \in T'} |c_j|^2 = 1 - \alpha$. Equation (1.7) together with (1.8) give

$$(1 - \delta_S)\alpha \leq \langle \sum_{j \in T} c_j v_j, \sum_{j' \in T} c_{j'} v_{j'} \rangle \leq (1 + \delta_S)\alpha,$$

$$|\langle \sum_{j \in T} c_j v_j, \sum_{j' \in T'} c_{j'} v_{j'} \rangle| \leq \theta_{S,S'} \alpha^{1/2} (1 - \alpha)^{1/2}$$

and

$$(1 - \delta_{S'})(1 - \alpha) \leq \langle \sum_{j \in T'} c_j v_j, \sum_{j' \in T'} c_{j'} v_{j'} \rangle \leq (1 + \delta_{S'})(1 - \alpha).$$

Hence,

$$\begin{aligned} |\langle \sum_{j \in \tilde{T}} c_j v_j, \sum_{j' \in \tilde{T}} c_{j'} v_{j'} \rangle - 1| &\leq \delta_S \alpha + \delta_{S'} (1 - \alpha) + 2\theta_{S,S'} \alpha^{1/2} (1 - \alpha)^{1/2} \\ &\leq \max(\delta_S, \delta_{S'}) + \theta_S \end{aligned}$$

as claimed. (We note that it is possible to optimize this bound a little further but will not do so here.) $\square$

III. GAUSSIAN RANDOM MATRICES

In this section, we argue that with overwhelming probability, Gaussian random matrices have “good” isometry constants. Consider a p by m matrix F whose entries are i.i.d. Gaussian with mean zero and variance $1/p$ and let T be a subset of the columns. We wish to study the extremal eigenvalues of $F_T^* F_T$. Following upon the work of Marchenko and Pastur [14], Geman [15], and Silverstein [16] (see also [17]) proved that

$$\begin{aligned} \lambda_{\min}(F_T^* F_T) &\rightarrow (1 - \sqrt{\gamma})^2 \quad \text{a.s.} \\ \lambda_{\max}(F_T^* F_T) &\rightarrow (1 + \sqrt{\gamma})^2 \quad \text{a.s.} \end{aligned}$$

in the limit where p and $|T| \rightarrow \infty$ with

$$|T|/p \rightarrow \gamma \leq 1.$$

In other words, this says that loosely speaking and in the limit of large p , the restricted isometry constant $\delta(F_T^* F_T)$ for a fixed T behaves like

$$1 - \delta(F_T^* F_T) \leq \lambda_{\min}(F_T^* F_T) \leq \lambda_{\max}(F_T^* F_T) \leq 1 + \delta(F_T^* F_T)$$

where

$$\delta(F_T^* F_T) \approx 2\sqrt{|T|/p} + |T|/p.$$

Restricted isometry constants must hold for all sets T of cardinality less or equal to S , and we shall make use of concentra-

tion inequalities to develop such a uniform bound. Note that for $T' \subset T$, we obviously have $\lambda_{\min}(F_T^* F_T) \leq \lambda_{\min}(F_{T'}^* F_{T'})$ and $\lambda_{\max}(F_{T'}^* F_{T'}) \leq \lambda_{\max}(F_T^* F_T)$ and, therefore, attention may be restricted to matrices of size S . Now, there are large deviation results about the singular values of F_T [18]. For example, letting $\sigma_{\max}(F_T)$ (resp., $\sigma_{\min}$) be the largest singular value of F_T so that $\sigma_{\max}^2(F_T) = \lambda_{\max}(F_T^* F_T)$ (resp., $\sigma_{\min}^2(F_T) = \lambda_{\min}(F_T^* F_T)$), Ledoux [19] applies the concentration inequality for Gaussian measures, and for each fixed $t > 0$, obtains the deviation bounds

$$\mathbf{P}\left(\sigma_{\max}(F_T) > 1 + \sqrt{|T|/p} + o(1) + t\right) \leq e^{-pt^2/2} \quad (3.19)$$

$$\mathbf{P}\left(\sigma_{\min}(F_T) < 1 - \sqrt{|T|/p} + o(1) - t\right) \leq e^{-pt^2/2}, \quad (3.20)$$

here the scaling of interest has $|T|/p$ constant, and $o(1)$ is a small term tending to zero as $p \rightarrow \infty$ and which can be calculated explicitly, see [20]. For example, this last reference shows that one can select $o(1)$ in (3.19) as $\frac{1}{2p^{1/3}} \cdot \gamma^{1/6} (1 + \sqrt{\gamma})^{2/3}$.

Lemma 3.1: Put $r = S/m$ and set

$$f(r) := \sqrt{m/p} \cdot \left(\sqrt{r} + \sqrt{2H(r)} \right) = \sqrt{S/p} + \sqrt{2H(r)m/p}$$

where H is the entropy function $H(q) := -q \log q - (1 - q) \log(1 - q)$ defined for $0 < q < 1$. For each $\epsilon > 0$, the restricted isometry constant δ_S of a p by m Gaussian matrix F obeys (for m and p large enough)

$$\mathbf{P}\left(\sqrt{1 + \delta_S} > 1 + (1 + \epsilon)f(r)\right) \leq 2e^{-mH(r) \cdot \epsilon/2}. \quad (3.21)$$

Proof: As discussed above, we may restrict our attention to sets $|T|$ such that $|T| = S$. Set $\lambda = (1 + \epsilon)f(r)$ for short and let $\bar{\sigma} := \max \sigma_{\max}(F_T)$ where the maximum is over all sets T of size S , and similarly $\underline{\sigma} := \min \sigma_{\min}(F_T)$. Then $1 + \delta_S = \max(\bar{\sigma}^2, 2 - \underline{\sigma}^2)$ and

$$\mathbf{P}(\sqrt{1 + \delta_S} > 1 + \lambda) \leq \mathbf{P}(\bar{\sigma} > 1 + \lambda) + \mathbf{P}(\sqrt{2 - \underline{\sigma}^2} > 1 + \lambda).$$

It will suffice to prove that

$$\mathbf{P}(\bar{\sigma} > 1 + \lambda) \leq e^{-mH(r) \cdot \epsilon/2}$$

and

$$\mathbf{P}(\underline{\sigma} < 1 - \lambda) \leq e^{-mH(r) \cdot \epsilon/2}$$

as for $\underline{\sigma} \geq 1 - \lambda$ we have $\sqrt{2 - \underline{\sigma}^2} \leq 1 + \lambda$. Both statements are proved in exactly the same way and we only include that for $\bar{\sigma}$.

Denote by η_p the $o(1)$ -term appearing in (3.19). Observe that

$$\begin{aligned} \mathbf{P}(\bar{\sigma} > 1 + \sqrt{S/p} + \eta_p + t) &\leq |\{T : |T| = S\}| \mathbf{P}(\sigma_{\max}(F_T) > 1 + \sqrt{S/p} + \eta_p + t) \\ &\leq \binom{m}{S} e^{-pt^2/2}. \end{aligned}$$

From Stirling's approximation $\log m! = m \log m - m + O(\log m)$ we have the well-known formula

$$\log \binom{m}{S} = mH(r) + O(\log m)$$

which gives

$$\log \mathbf{P}(\bar{\sigma} > 1 + \sqrt{S/p} + \eta_p + t) \leq mH(r) + O(\log m) - pt^2/2.$$

We have

$$1 + (1 + \epsilon)f(r) \geq 1 + \sqrt{S/p} + (1 + \epsilon)\sqrt{2mH(r)/p}$$

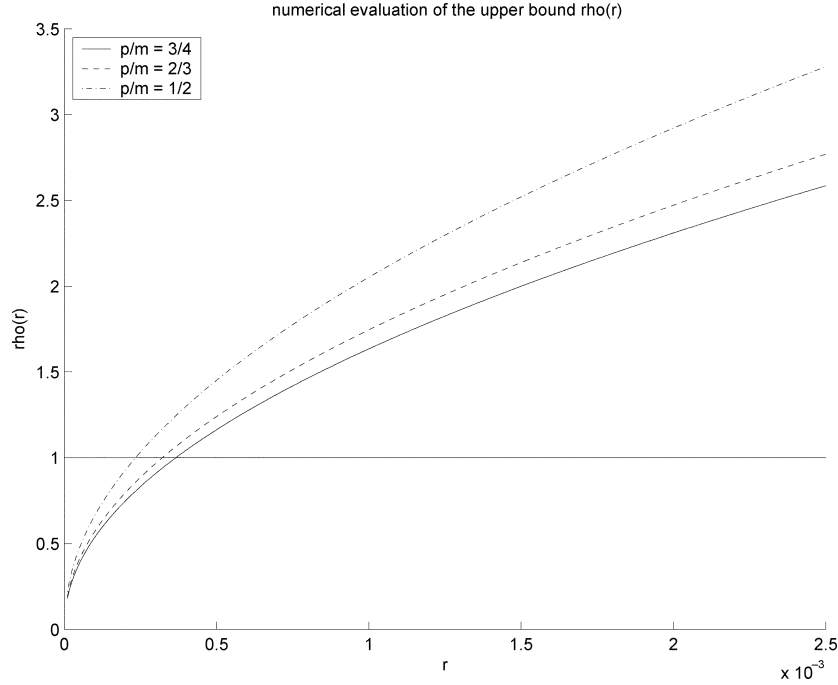


Fig. 1. Behavior of the upper bound $\rho_{p/m}(r)$ for three values of the ratio p/m , namely, $p/m = 3/4, 2/3, 1/2$.

and suppose m and p are large enough so that

$$\eta_p \leq \epsilon/2 \cdot \sqrt{2mH(r)/p}.$$

Then

$$1 + (1 + \epsilon)f(r) \geq 1 + \sqrt{S/p} + \eta_p + (1 + \epsilon/2)\sqrt{2mH(r)/p}$$

and, therefore,

$$\begin{aligned} \log \mathbf{P}(\bar{\sigma} > 1 + (1 + \epsilon)f(r)) \\ \leq mH(r) + O(\log m) - \frac{p(1 + \epsilon/2)^2 2mH(r)}{2p} \\ \leq O(\log m) - mH(r)\epsilon. \end{aligned}$$

For m sufficiently large, the term $O(\log m)$ is less than $mH(r)\epsilon/2$, which concludes the proof. $\square$

Ignoring the ϵ 's, Lemma 3.1 states that with overwhelming probability

$$\delta_S < -1 + [1 + f(r)]^2. \quad (3.22)$$

A similar conclusion holds for δ_{2S} and δ_{3S} and, therefore, we established that with very high probability

$$\delta_S + \delta_{2S} + \delta_{3S} < \rho_{p/m}(r) \quad (3.23)$$

with

$$\rho_{p/m}(r) = \sum_{j=1}^3 -1 + [1 + f(jr)]^2.$$

In conclusion, Lemma 1.1 shows that the hypothesis of our main theorem holds provided that the ratio $r = S/m$ is small so that $\rho_{p/m}(r) < 1$. In other words, in the limit of large samples, S/m

may be taken as any value obeying $\rho_{p/m}(S/m) < 1$ which we used to give numerical values in Theorem 1.5. Fig. 1 graphs the function $\rho_{p/m}(r)$ for several values of the ratio p/m .

IV. NUMERICAL EXPERIMENTS

This section investigates the practical ability of ℓ_1 to recover an object $f \in \mathbf{R}^n$ from corrupted data $y = Af + e$, $y \in \mathbf{R}^m$ (or equivalently to recover the sparse vector of errors $e \in \mathbf{R}^m$ from the underdetermined system of equations $Fe = z \in \mathbf{R}^{m-n}$). The goal here is to evaluate empirically the location of the break-point as to get an accurate sense of the performance one might expect in practice. In order to do this, we performed a series of experiments designed as follows:

- 1) select n (the size of the input signal) and m so that with the same notations as before, A is an m by n matrix; sample A with independent Gaussian entries;
- 2) select S as a percentage of m ;
- 3) select a support set T of size $|T| = S$ uniformly at random, and sample a vector e on T with i.i.d. Gaussian entries;¹
- 4) make $\tilde{y} = Af + e$ (the choice of f does not matter as is clear from the discussion and here, f is also selected at random), solve (P'_1) and obtain f^* ;
- 5) compare f to f^* ;
- 6) repeat 100 times for each S and A ;
- 7) repeat for various sizes of n and m .

The results are presented in Figs. 2 and 3. Fig. 2 examines the situation in which the length of the code is twice that of the input vector $m = 2n$, for $m = 512$ and $m = 1024$. Our experiments show that one recovers the input vector all the time

¹Just as in [11], the results presented here do not seem to depend on the actual distribution used to sample the errors.

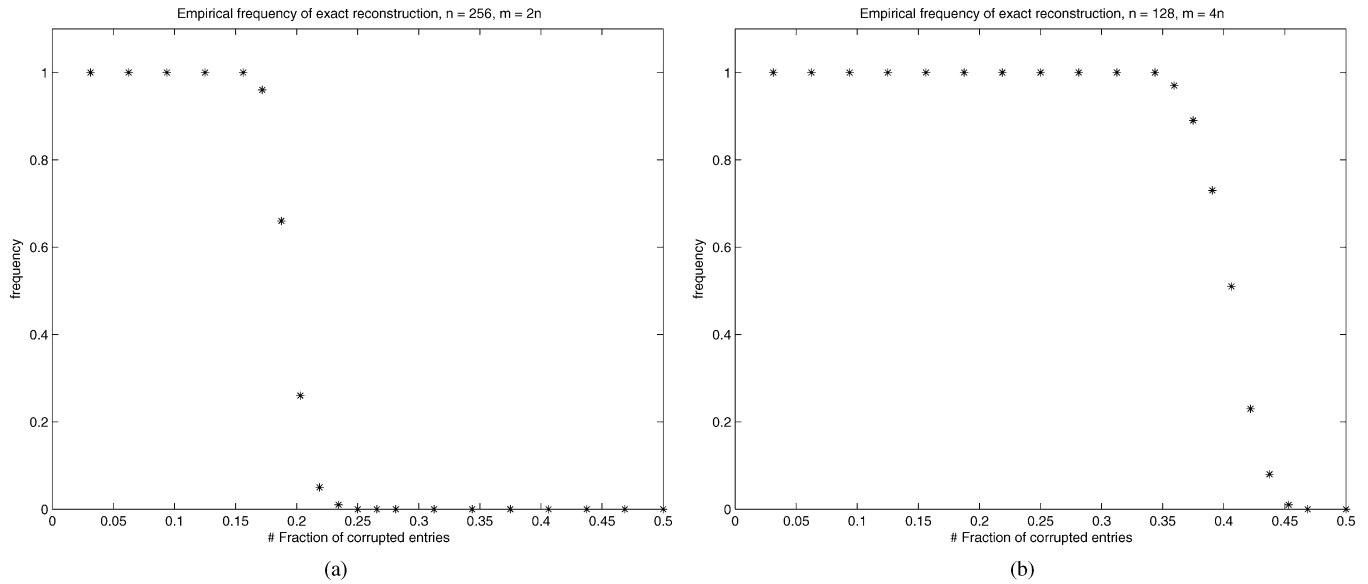


Fig. 2. ℓ_1 -recovery of an input signal from $y = Af + e$ with A an m by n matrix with independent Gaussian entries. In this experiment, we “oversample” the input signal by a factor 2 so that $m = 2n$. (a) Success rate of (P_1) for $m = 512$. (b) Success rate of (P_1) for $m = 1024$. Observe the similar pattern and cutoff point. In these experiments, exact recovery occurs as long as about 17% or less of the entries are corrupted.

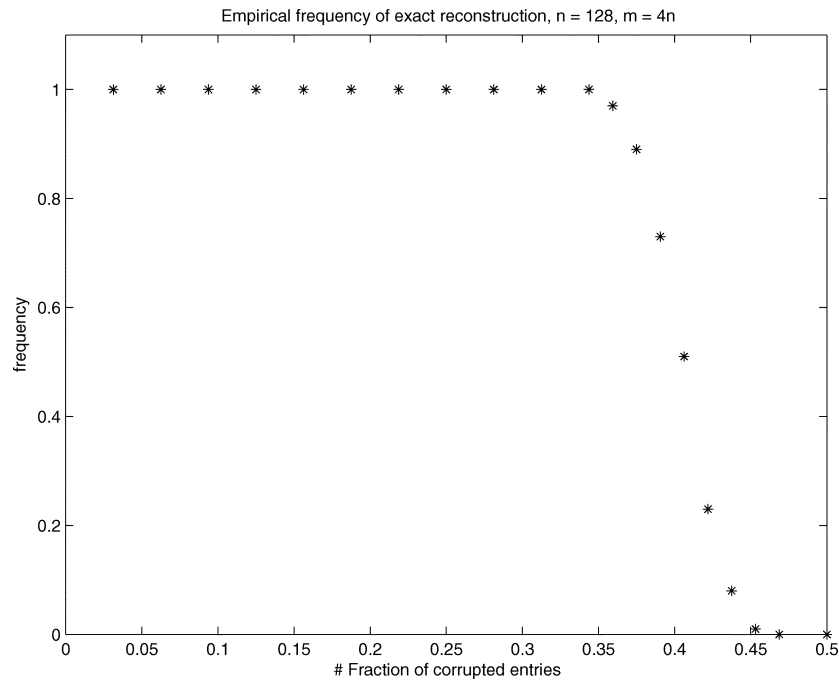


Fig. 3. ℓ_1 -recovery of an input signal from $y = Af + e$ with A an m by n matrix with independent Gaussian entries. In this experiment, we “oversample” the input signal by a factor 4 so that $m = 4n$. In these experiments, exact recovery occurs as long as about 34% or less of the entries are corrupted.

as long as the fraction of the corrupted entries is below 17%. This holds for $m = 512$ (Fig. 2(a)) and $m = 1024$ (Fig. 2(b)). In Fig. 3, we investigate how these results change as the length of the codewords increases compared to the length of the input, and examine the situation in which $m = 4n$, with $m = 512$. Our experiments show that one recovers the input vector all the time as long as the fraction of the corrupted entries is below 34%.

It might be tempting to compare the empirical cutoff observed in the figures with the theoretical limit calculated in Corollary 1.8. For example, for $n/m = 1/4$, Fig. 3 suggests a cutoff near 35% while Corollary 1.8 gives a theoretical limit of $\rho =$

$(1 - 1/4)/2 = 37.5\%$. One needs to apply caution, however, as these two quantities refer to two distinct situations. On the one hand, the limit $\rho = (1 - n/m)/2$ guarantees that *all* signals may be decoded exactly. On the other hand, the numerical simulations examine the case where both f and the error vector e are chosen randomly, and does not search for the worst possible pair. In this sense, the simulations are not really “testing” Theorem 1.5. Under a probability model for the error pattern e , say, it is clear that information theoretically, one can tolerate a larger fraction of errors if one wishes to be able to decode accurately only for *most* as opposed to *all* error vectors.

V. OPTIMAL SIGNAL RECOVERY

Our recent work [1] developed a set of ideas showing that it is surprisingly possible to reconstruct interesting classes of signals accurately from highly incomplete measurements. The results in this paper are inspired and improve upon this earlier work and we now elaborate on this connection. Suppose we wish to reconstruct an object α in $\mathbf{R}^m$ from the K linear measurements

$$y_k = \langle \alpha, \phi_k \rangle \quad k = 1, \dots, K \quad \text{or} \quad y = F\alpha \quad (5.24)$$

with ϕ_k , the k th row of the matrix F . Of special interest is the vastly underdetermined case $K \ll N$, where there are many more unknowns than observations. We choose to formulate the problem abstractly but for concreteness, we might think of α as the coefficients $\alpha = \Psi^* f$ of a digital signal or image f in some nice orthobasis, e.g., a wavelet basis, so that the information about the signal is of the form $y = F\alpha = F\Psi^* f$.

Suppose now that the object of interest is *compressible* in the sense that the reordered entries of α decay like a power law; concretely, suppose that the entries of α , rearranged in decreasing order of magnitude, $|\alpha|_{(1)} \geq |\alpha|_{(2)} \geq \dots \geq |\alpha|_{(m)}$, obey

$$|\alpha|_{(k)} \leq B \cdot k^{-s} \quad (5.25)$$

for some $s \geq 1$. We will denote by $\mathcal{F}_s(B)$ the class of all signals $\alpha \in \mathbf{R}^m$ obeying (5.25). The claim is that it is possible to reconstruct compressible signals from only a small number of random measurements.

Theorem 5.1: Let F be the measurement matrix as in (5.24) and consider the solution $\alpha^\#$ to

$$\min_{\tilde{\alpha} \in \mathbf{R}^m} \|\tilde{\alpha}\|_{\ell_1}, \quad \text{subject to } F\tilde{\alpha} = y. \quad (5.26)$$

Let $S \leq K$ such that $\delta_S + 2\theta_S + \theta_{S,2S} < 1$ and set $\lambda = K/S$. Then $\alpha^\#$ obeys

$$\sup_{\alpha \in \mathcal{F}_s(B)} \|\alpha - \alpha^\#\| \leq C \cdot (K/\lambda)^{-(s-1/2)}. \quad (5.27)$$

To appreciate the content of the theorem, suppose one would have available an *oracle* letting us know which coefficients α_k , $1 \leq k \leq m$, are large (e.g., in the scenario we considered earlier, the oracle would tell us which wavelet coefficients of f are large). Then we would acquire information about the K largest coefficients and obtain a truncated version α_K of α obeying

$$\|\alpha - \alpha_K\| \geq c \cdot K^{-(s-1/2)}$$

for generic elements taken from $\mathcal{F}_s(B)$. Now (5.27) says that not knowing anything about the location of the largest coefficients, one can essentially obtain the same approximation error by nonadaptive sampling, provided the number of measurements is increased by a factor λ . The larger S , the smaller the oversampling factor, and hence, the connection with the decoding problem. Such considerations make clear that Theorem 5.1 supplies a very concrete methodology for recovering a compressible object from limited measurements and as such, it may have a significant bearing on many fields of science and technology. We refer the reader to [1] and [21] for a discussion of its implications.

Suppose, for example, that F is a Gaussian random matrix as in Section III. We will assume the same special normalization so that the variance of each individual entry is equal to $1/K$. Calculations identical to those from Section III give that with overwhelming probability, F obeys the hypothesis of the theorem provided that

$$S \leq K/\lambda, \quad \lambda = \rho \cdot \log(m/K)$$

for some positive constant $\rho > 0$. Now consider the statement of the theorem; there is a way to invoke linear programming and obtain a reconstruction based upon $O(K \log(m/K))$ measurements only, which is at least as good as that one would achieve by knowing all the information about f and selecting the K largest coefficients. In fact, this is an optimal statement as (5.27) correctly identifies the minimum number of measurements needed to obtain a given precision. In short, it is impossible to obtain a precision of about $K^{-(s-1/2)}$ with fewer than $K \log(m/K)$ measurements, see [1], [21].

Theorem 5.1 is stronger than our former result, namely, Theorem 1.4 in [1]. To see why this is true, recall the former claim: [1] introduced two conditions, the uniform uncertainty principle (UUP) and the exact reconstruction principle (ERP). In a nutshell, a *random* matrix F obeys the UUP with oversampling factor λ if F obeys

$$\delta_S \leq 1/2, \quad S = \rho \cdot K/\lambda \quad (5.28)$$

with probability at least $1 - O(N^{-\gamma/\rho})$ for some fixed positive constant $\gamma > 0$. Second, a measurement matrix F obeys the ERP with oversampling factor λ if for each fixed subset T of size $|T| \leq S$ (5.28) and each “sign” vector c defined on T , there exists with the same overwhelmingly large probability a vector $w \in H$ with the following properties:

- i) $\langle w, v_j \rangle = c_j$, for all $j \in T$;
- ii) and $|\langle w, v_j \rangle| \leq \frac{1}{2}$, for all j not in T .

Note that these are the conditions listed at the beginning of Section II except for the $1/2$ factor on the complement of T . Fix $\alpha \in \mathcal{F}_s(B)$. [1] argued that if a random matrix obeyed the UUP and the ERP both with oversampling factor λ , then

$$\|\alpha - \alpha^\#\| \leq C \cdot (K/\lambda)^{-(s-1/2)}$$

with inequality holding with the same probability as before. Against this background, several comments are now in order.

- First, the new statement is more general as it applies to all matrices, not just random matrices.
- Second, whereas our previous statement argued that for each $\alpha \in \mathbf{R}^m$, one would be able—with high probability—to reconstruct α accurately, it did not say anything about the worst case error for a fixed measurement matrix F . This is an instance where the order of the quantifiers plays a role. Do we need different F ’s for different objects? Theorem 5.1 answers this question unambiguously; the same F will provide an optimal reconstruction for *all* the objects in the class.
- Third, Theorem 5.1 says that the ERP condition is redundant, and hence the hypothesis may be easier to check in practice. In addition, eliminating the ERP isolates the real

reason for success as it ties everything down to the UUP. In short, the ability to recover an object from limited measurements depends on how close F is to an orthonormal system, but only when restricting attention to sparse linear combinations of columns.

We will not prove this theorem as this is a minor modification of that of Theorem 1.4 in the aforementioned reference. The key point is to observe that if F obeys the hypothesis of our theorem, then by definition F obeys the UUP with probability one, but F also obeys the ERP, again with probability one, as this is the content of Lemma 2.2. Hence, both the UUP and ERP hold and therefore, the conclusion of Theorem 5.1 follows. (The fact that the ERP actually holds for all sign vectors of size less than S is the reason why (5.27) holds uniformly over all elements taken from $\mathcal{F}_s(B)$, see [1].)

VI. DISCUSSION

A. Connections With Other Works

In our linear programming model, the plaintext and ciphertext had real-valued components. Another intensively studied model occurs when the plaintext and ciphertext take values in the finite field $F_2 := \{0, 1\}$. In recent work of Feldman *et al.* [22]–[25], linear programming methods (based on relaxing the space of codewords to a convex polytope) were developed to establish a polynomial-time decoder which can correct a constant fraction of errors, and also achieve the information-theoretic capacity of the code. Our methods are restricted to real-valued texts, and the work cited above requires texts in F_2 . Also, our error analysis is deterministic and is thus guaranteed to correct arbitrary errors provided that they are sufficiently sparse. In summary, there does not seem to be any explicit known connection with this line of work but it would perhaps be of future interest to explore if there is one.

The ideas presented in this paper may be adapted to recover input vectors taking values from a finite alphabet. We hope to report on work in progress in a followup paper.

B. Improvements

There is little doubt that more elaborate arguments will yield versions of Theorem 1.5 with tighter bounds. Immediately following the proof of Lemma 2.2, we already remarked that one might slightly improve the condition $\delta_S + \theta_{S,S} + \theta_{S,2S} < 1$ at the expense of considerable complications. More to the point, we must admit that we used well-established tools from random matrix theory and it is likely that more sophisticated ideas might be deployed successfully. We now discuss some of these.

Our main hypothesis reads $\delta_S + \theta_{S,S} + \theta_{S,2S} < 1$ but in order to reduce the problem to the study of those δ numbers (and use known results), our analysis actually relied upon the more stringent condition $\delta_S + \delta_{2S} + \delta_{3S} < 1$ instead, since

$$\delta_S + \theta_{S,S} + \theta_{S,2S} \leq \delta_S + \delta_{2S} + \delta_{3S}.$$

This introduces a gap. Consider a fixed set T of size $|T| = S$. Using the notations of that Section III, we argued that

$$\delta(F_T) \approx 2\sqrt{S/p} + S/p$$

and developed a large deviation bound to quantify the departure from the right-hand side. Now let T and T' be two disjoint sets of respective sizes S and S' and consider $\theta(F_T, F_{T'})$: $\theta(F_T, F_{T'})$ is the cosine of the principal angle between the two random subspaces spanned by the columns of F_T and $F_{T'}$ respectively; formally

$$\theta(F_T, F_{T'}) = \sup \langle u, u' \rangle, \\ u \in \text{span}(F_T), u' \in \text{span}(F_{T'}), \|u\| = \|u'\| = 1.$$

We remark that this quantity plays an important analysis in statistical analysis because of its use to test the significance of correlations between two sets of measurements, compare the literature on canonical correlation analysis [26]. Among other things, it is known [27] that

$$\theta(F_T, F_{T'}) \rightarrow \sqrt{\gamma(1-\gamma')} + \sqrt{\gamma'(1-\gamma)} \quad \text{a.s.}$$

as $p \rightarrow \infty$ with $S/p \rightarrow \gamma$ and $S'/p \rightarrow \gamma'$. In other words, whereas we used the limiting behaviors

$$\delta(F_{2T}) \rightarrow 2\sqrt{2\gamma} + 2\gamma, \quad \delta(F_{3T}) \rightarrow 2\sqrt{3\gamma} + 3\gamma$$

there is a chance one might employ instead

$$\theta(F_T, F_{T'}) \rightarrow 2\sqrt{\gamma(1-\gamma)}$$

for $|T| = |T'| = S$, and

$$\theta(F_T, F_{T'}) \rightarrow \sqrt{\gamma(1-2\gamma)} + \sqrt{2\gamma(1-\gamma)}$$

for $|T'| = 2|T| = 2S$, which are better estimates. Just as in Section III, one might then look for concentration inequalities transforming this limiting behavior into corresponding large deviation inequalities. We are aware of very recent work of Johnstone and his colleagues [28] which might be here of substantial help.

Finally, tighter large deviation bounds might exist together with more clever strategies to derive uniform bounds (valid for all T of size less than S) from individual bounds (valid for a single T). With this in mind, it is interesting to note that our approach hits a limit as

$$\liminf_{S \rightarrow \infty, S/m \rightarrow r} \delta_S + \theta_{S,S} + \theta_{S,2S} \geq J(m/p \cdot r) \quad (6.29)$$

where $J(r) := 2\sqrt{r} + r + (2 + \sqrt{2})\sqrt{r(1-r)} + \sqrt{r(1-2r)}$. Since $J(r)$ is greater than 1 if and only if $r > 2.36$, one would certainly need new ideas to improve Theorem 1.5 beyond the cutoff point in the range of about 2%. The lower limit (6.29) is probably not sharp since it does not explicitly take into account the ratio between m and p ; at best, it might serve as an indication of the limiting behavior when the ration p/m is not too small.

C. Other Coding Matrices

This paper introduced general results stating that it is possible to correct for errors by ℓ_1 -minimization. We then explained how the results specialize in the case where the coding matrix A is sampled from the Gaussian ensemble. It is clear, however,

that one could use other matrices and still obtain similar results; namely, that (P'_1) recovers f exactly provided that the number of corrupted entries does not exceed $\rho \cdot m$. In fact, our previous work suggests that partial Fourier matrices would enjoy similar properties [1], [11]. Other candidates might be the so-called *noiselets* of Coifman, Geshwind, and Meyer [29]. These alternatives might be of great practical interest because they would come with fast algorithms for applying A or A^* to an arbitrary vector g and, hence, speed up the computations to find the ℓ_1 -minimizer.

APPENDIX PROOF OF LEMMA 2.2

We may normalize $\sum_{j \in T} |c_j|^2 = \sqrt{S}$. Write $T_0 := T$. Using Lemma 2.1, we can find a vector $w_1 \in H$ and a set $T_1 \subseteq J$ such that

$$\begin{aligned} T_0 \cap T_1 &= \emptyset \\ |T_1| &\leq S \\ \langle w_1, v_j \rangle &= c_j, \quad \text{for all } j \in T_0 \\ |\langle w_1, v_j \rangle| &\leq \frac{\theta_{S,S'}}{(1-\delta_S)}, \quad \text{for all } j \notin T_0 \cup T_1 \\ \left(\sum_{j \in T_1} |\langle w_1, v_j \rangle|^2 \right)^{1/2} &\leq \frac{\theta_S}{1-\delta_S} \sqrt{S} \\ \|w_1\| &\leq K. \end{aligned}$$

Applying Lemma 2.1 iteratively gives a sequence of vectors $w_{n+1} \in H$ and sets $T_{n+1} \subseteq J$ for all $n \geq 1$ with the properties

$$\begin{aligned} T_n \cap (T_0 \cup T_{n+1}) &= \emptyset \\ |T_{n+1}| &\leq S \\ \langle w_{n+1}, v_j \rangle &= \langle w_n, v_j \rangle, \quad \text{for all } j \in T_n \\ \langle w_{n+1}, v_j \rangle &= 0, \quad \text{for all } j \in T_0 \\ |\langle w_{n+1}, v_j \rangle| &\leq \frac{\theta_S}{1-\delta_S} \left(\frac{\theta_{S,2S}}{1-\delta_S} \right)^n, \\ &\quad \text{for all } j \notin T_0 \cup T_n \cup T_{n+1} \\ \left(\sum_{j \in T_{n+1}} |\langle w_{n+1}, v_j \rangle|^2 \right)^{1/2} &\leq \frac{\theta_S}{1-\delta_S} \left(\frac{\theta_{S,2S}}{1-\delta_S} \right)^n \sqrt{S} \\ \|w_{n+1}\| &\leq \left(\frac{\theta_S}{1-\delta_S} \right)^{n-1} K. \end{aligned}$$

By hypothesis, we have $\frac{\theta_{S,2S}}{1-\delta_S} < 1$. Thus, if we set

$$w := \sum_{n=1}^{\infty} (-1)^{n-1} w_n$$

then the series is absolutely convergent and, therefore, w is a well-defined vector in H . We now study the coefficients

$$\langle w, v_j \rangle = \sum_{n=1}^{\infty} (-1)^{n-1} \langle w_n, v_j \rangle \quad (1.30)$$

for $j \in J$.

Consider first $j \in T_0$, it follows from the construction that $\langle w_1, v_j \rangle = c_j$ and $\langle w_n, v_j \rangle = 0$ for all $n \geq 2$, and hence,

$$\langle w, v_j \rangle = c_j, \quad \text{for all } j \in T_0.$$

Second, fix j with $j \notin T_0$ and let $I_j := \{n \geq 1 : j \in T_n\}$. Since T_n and T_{n+1} are disjoint, we see that the integers in the set I_j are spaced at least two apart. Now if $n \in I_j$, then by definition $j \in T_n$ and, therefore,

$$\langle w_{n+1}, v_j \rangle = \langle w_n, v_j \rangle.$$

In other words, the n and $n+1$ terms in (1.30) cancel each other out. Thus, we have

$$\langle w, v_j \rangle = \sum_{n \geq 1, n-1 \notin I_j} (-1)^{n-1} \langle w_n, v_j \rangle.$$

On the other hand, if $n, n-1 \notin I_j$, and $n \neq 0$, then $j \notin T_n \cap T_{n-1}$ and

$$|\langle w_n, v_j \rangle| \leq \frac{\theta_{S,S}}{1-\delta_S} \left(\frac{\theta_{S,2S}}{1-\delta_S} \right)^{n-1}$$

which by the triangle inequality and the geometric series formula gives

$$\left| \sum_{n \geq 1, n-1 \notin I_j} (-1)^{n-1} \langle w_n, v_j \rangle \right| \leq \frac{\theta_{S,S}}{1-\delta_S - \theta_{S,2S}}.$$

In conclusion

$$|\langle w, v_j \rangle - 1_{0 \in I_j} \langle w_0, v_j \rangle| \leq \frac{\theta_{S,S}}{1-\delta_S - \theta_{S,2S}}$$

and since $|T| \leq S$, the claim follows.

ACKNOWLEDGMENT

The authors wish to thank Rafail Ostrovsky for pointing out possible connections between their earlier work and the decoding problem. E. J. Candes would also like to acknowledge inspiring conversations with Leonard Schulmann and Justin Romberg.

REFERENCES

- [1] E. J. Candes and T. Tao, "Near optimal signal recovery from random projections: Universal encoding strategies?," *IEEE Trans. Inf. Theory*. Available on the AxiXiv preprint server: math.CA/0410542, submitted for publication.
- [2] M. R. Garey and D. S. Johnson, *Computers and Intractability: A Guide to the Theory of NP-Completeness*. San Francisco, CA: Freeman, 1979.
- [3] S. S. Chen, D. L. Donoho, and M. A. Saunders, "Atomic decomposition by basis pursuit," *SIAM J. Sci. Comput.*, vol. 20, pp. 33–61, 1998.
- [4] P. Bloomfield and W. Steiger, *Least Absolute Deviations: Theory, Applications, and Algorithms*. Boston, MA: Birkhäuser, 1983.
- [5] D. L. Donoho and X. Huo, "Uncertainty principles and ideal atomic decomposition," *IEEE Trans. Inf. Theory*, vol. 47, no. 7, pp. 2845–2862, Nov. 2001.
- [6] R. Gribonval and M. Nielsen, "Sparse representations in unions of bases," *IEEE Trans. Inf. Theory*, vol. 49, no. 12, pp. 3320–3325, Dec. 2003.
- [7] D. L. Donoho and M. Elad, "Optimally sparse representation in general (nonorthogonal) dictionaries via ℓ_1 minimization," in *Proc. Nat. Acad. Sci. USA*, vol. 100, 2003, pp. 2197–2202.

- [8] M. Elad and A. M. Bruckstein, "A generalized uncertainty principle and sparse representation in pairs of $\mathbf{R}^N$ bases," *IEEE Trans. Inf. Theory*, vol. 48, no. 9, pp. 2558–2567, Sep. 2002.
- [9] J. A. Tropp, "Greed is good: Algorithmic results for sparse approximation," *IEEE Trans. Inf. Theory*, vol. 50, no. 10, pp. 2231–2242, Oct. 2004.
- [10] E. J. Candes, J. Romberg, and T. Tao, "Robust uncertainty principles: Exact signal reconstruction from highly incomplete frequency information," *IEEE Trans. Inf. Theory*. Available on the ArXiv preprint server: math.GM/0409186, submitted for publication.
- [11] E. J. Candes and J. Romberg, "Quantitative robust uncertainty principles and optimally sparse decompositions," *Found. Comput. Math.*. Available on the ArXiv preprint server: math.CA/0411273, to be published.
- [12] D. L. Donoho, "For most large underdetermined systems of linear equations the minimal ℓ_1 -norm solution is also the sparsest solution," unpublished manuscript, Sep. 2004.
- [13] S. Boyd and L. Vandenberghe, *Convex Optimization*. Cambridge, U.K.: Cambridge Univ. Press, 2004.
- [14] V. A. Marchenko and L. A. Pastur, "Distribution of eigenvalues in certain sets of random matrices" (in Russian), *Mat. Sbornik (N.S.)*, vol. 72, pp. 407–535, 1967.
- [15] S. Geman, "A limit theorem for the norm of random matrices," *Ann. Probab.*, vol. 8, pp. 252–261, 1980.
- [16] J. W. Silverstein, "The smallest eigenvalue of a large dimensional Wishart matrix," *Ann. Probab.*, vol. 13, pp. 1364–1368, 1985.
- [17] Z. D. Bai and Y. Q. Yin, "Limit of the smallest eigenvalue of a large-dimensional sample covariance matrix," *Ann. Probab.*, vol. 21, pp. 1275–1294, 1993.
- [18] S. J. Szarek, "Condition numbers of random matrices," *J. Complexity*, vol. 7, pp. 131–149, 1991.
- [19] M. Ledoux, *The Concentration of Measure Phenomenon*. Providence, RI: Amer. Math. Soc., 2001, vol. 89, Mathematical Surveys and Monographs.
- [20] N. El Karoui, "New results about random covariance matrices and statistical applications," Ph.D. dissertation, Stanford Univ., Stanford, CA, 2004.
- [21] D. L. Donoho, "Compressed sensing," unpublished manuscript, Sep. 2004.
- [22] J. Feldman, "Decoding error-correcting codes via linear programming," Ph.D. dissertation, MIT, Cambridge, MA, 2003.
- [23] —, "LP decoding achieves capacity," presented at the ACM-SIAM Symp. Discrete Algorithms (SODA), Vancouver, BC, Canada, Jan. 2005.
- [24] J. Feldman, T. Malkin, C. Stein, R. A. Servedio, and M. J. Wainwright, "LP decoding corrects a constant fraction of errors," in *Proc. IEEE Int. Symp. Information Theory (ISIT)*, Chicago, IL, Jun./Jul. 2004, p. 68.
- [25] J. Feldman, M. J. Wainwright, and D. R. Karger, "Using linear programming to decode binary linear codes," *IEEE Trans. Inf. Theory*, vol. 51, no. 3, pp. 954–972, Mar. 2005.
- [26] R. J. Muirhead, *Aspects of Multivariate Statistical Theory*. New York: Wiley, 1982, Wiley Series in Probability and Mathematical Statistics.
- [27] K. W. Wachter, "The limiting empirical measure of multiple discriminant ratios," *Ann. Statist.*, vol. 8, pp. 937–957, 1980.
- [28] I. M. Johnstone, "Large covariance matrices (Third 2004 Wald Lecture)," presented at the 6th World Congress of the Bernoulli Society and 67th Annual Meeting of the IMS, Barcelona, Spain, 2004.
- [29] R. Coifman, F. Geshwind, and Y. Meyer, "Noiselets," *Appl. Comput. Harmon. Anal.*, vol. 10, pp. 27–44, 2001.