

Introduction to Compressive Sensing

Collection Editor:

Marco F. Duarte

Introduction to Compressive Sensing

Collection Editor:

Marco F. Duarte

Authors:

Mark A. Davenport
Ronald DeVore
Marco F. Duarte
Chinmay Hegde
Jason Laska

Michael A Lexa
Shriram Sarvotham
Mona Sheikh
Wotao Yin

Online:

< <http://cnx.org/content/col11355/1.2/> >

C O N N E X I O N S

Rice University, Houston, Texas

This selection and arrangement of content as a collection is copyrighted by Marco F. Duarte. It is licensed under the Creative Commons Attribution 3.0 license (<http://creativecommons.org/licenses/by/3.0/>).

Collection structure revised: August 29, 2011

PDF generated: September 23, 2011

For copyright and attribution information for the modules contained in this collection, see p. 106.

Table of Contents

1 Analog Sampling Theory

1.1	The Shannon-Whitaker Sampling Theorem	1
1.2	Stable Signal Representations	3
1.3	Optimal Encoding	5
1.4	Kolmogorov Entropy	7
1.5	Optimal Encoding of Bandlimited Signals	8

2 Sparsity and Compressibility

2.1	Introduction to vector spaces	11
2.2	Bases and frames	13
2.3	Sparse representations	14
2.4	Compressible signals	17

3 Compressive Sensing

3.1	Sensing matrix design	21
3.2	Null space conditions	22
3.3	The restricted isometry property	24
3.4	The RIP and the NSP	28
3.5	Matrices that satisfy the RIP	30
3.6	Coherence	32
3.7	Sub-Gaussian random variables	33
3.8	Concentration of measure for sub-Gaussian random variables	35
3.9	Proof of the RIP for sub-Gaussian matrices	38

4 ℓ_1 -norm minimization

4.1	Signal recovery via ℓ_1 -norm minimization	43
4.2	Noise-free signal recovery	45
4.3	Signal recovery in noise	47
4.4	Instance-optimal guarantees revisited	51
4.5	The cross-polytope and phase transitions	53
4.6	Sparse recovery algorithms	54
4.7	Convex optimization-based methods	55
4.8	Greedy algorithms	58
4.9	Combinatorial algorithms	62
4.10	Bayesian methods	64

5 Applications of Compressive Sensing

5.1	Linear regression and model selection	67
5.2	Sparse error correction	67
5.3	Group testing and data stream algorithms	68
5.4	Compressive medical imaging	69
5.5	Analog-to-information conversion	69
5.6	Single-pixel camera	72
5.7	Hyperspectral imaging	76
5.8	Compressive processing of manifold-modeled data	81
5.9	Inference using compressive measurements	85
5.10	Compressive sensor networks	88
5.11	Genomic sensing	90

Glossary	92
Bibliography	93

Index	104
Attributions	106

Chapter 1

Analog Sampling Theory

1.1 The Shannon-Whitaker Sampling Theorem¹

The classical theory behind the encoding analog signals into bit streams and decoding bit streams back into signals, rests on a famous sampling theorem which is typically referred to as the Shannon-Whitaker Sampling Theorem. In this course, this sampling theory will serve as a benchmark to which we shall compare the new theory of compressed sensing.

To introduce the Shannon-Whitaker theory, we first define the class of bandlimited signals. A bandlimited signal is a signal whose Fourier transform only has finite support. We shall denote this class as B_A and define it in the following way:

$$B_A := \{f \in L_2(\mathbb{R}) : \hat{f}(\omega) = 0, |\omega| \geq A\pi\}. \quad (1.1)$$

Here, the Fourier transform of f is defined by

$$\hat{f}(\omega) := \frac{1}{\sqrt{2\pi}} \int_{\mathbb{R}} f(t) e^{-i\omega t} dt. \quad (1.2)$$

This formula holds for any $f \in L_1$ and extends easily to $f \in L_2$ via limits. The inversion of the Fourier transform is given by

$$f(t) := \frac{1}{\sqrt{2\pi}} \int_{\mathbb{R}} \hat{f}(\omega) e^{i\omega t} d\omega. \quad (1.3)$$

Theorem 1.1: Shannon-Whitaker Sampling Theorem

If $f \in B_A$, then f can be uniquely determined by the uniformly spaced samples $f\left(\frac{n}{A}\right)$ and in fact, is given by

$$f(t) = \sum_{n \in \mathbb{Z}} f\left(\frac{n}{A}\right) \operatorname{sinc}(\pi(A t - n)), \quad (1.4)$$

where $\operatorname{sinc}(t) = \frac{\sin t}{t}$.

Proof:

It is enough to consider $A = 1$, since all other cases can be reduced to this through a simple change of variables. Because $f \in B_{A=1}$, the Fourier inversion formula takes the form

$$f(t) = \frac{1}{\sqrt{2\pi}} \int_{-\pi}^{\pi} \hat{f}(\omega) e^{i\omega t} d\omega. \quad (1.5)$$

¹This content is available online at <<http://cnx.org/content/m15146/1.2/>>.

Define $F(\omega)$ as the 2π periodization of $\hat{f}$,

$$F(\omega) := \sum_{n \in \mathbb{Z}} \hat{f}(\omega - 2n\pi). \quad (1.6)$$

Because $F(\omega)$ is periodic, it admits a Fourier series representation

$$F(\omega) = \sum_{n \in \mathbb{Z}} c_n e^{-in\omega}, \quad (1.7)$$

where the Fourier coefficients c_n given by

$$\begin{aligned} c_n &= \frac{1}{2\pi} \int_{-\pi}^{\pi} F(\omega) e^{in\omega} d\omega \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{f}(\omega) e^{in\omega} d\omega. \end{aligned} \quad (1.8)$$

By comparing ((1.8)) with ((1.5)), we conclude that

$$c_n = \frac{1}{\sqrt{2\pi}} f(n). \quad (1.9)$$

Therefore by plugging ((1.9)) back into ((1.6)), we have that

$$F(\omega) = \frac{1}{\sqrt{2\pi}} \sum_{n \in \mathbb{Z}} f(n) e^{-in\omega}. \quad (1.10)$$

Now, because

$$\hat{f}(\omega) = F(\omega) \chi_{[-\pi, \pi]} = \frac{1}{\sqrt{2\pi}} \sum_{n \in \mathbb{Z}} f(n) e^{-in\omega} \chi_{[-\pi, \pi]}, \quad (1.11)$$

and because of the facts that

$$\begin{aligned} \mathcal{F}(\chi_{[-\pi, \pi]}) &= \frac{1}{\sqrt{2\pi}} \text{sinc}(\pi\omega) \quad \text{and} \\ \mathcal{F}(g(t-n)) &= e^{-in\omega} \mathcal{F}(g(t)), \end{aligned} \quad (1.12)$$

we conclude

$$f(t) = \sum_{n \in \mathbb{Z}} f(n) \text{sinc}(\pi(t-n)). \quad (1.13)$$

Comments:

1. (Good news) The set $\{\text{sinc}(\pi(t-n))\}_{n \in \mathbb{Z}}$ is an orthogonal system and therefore, has the property that the L_2 norm of the function and its Fourier coefficients are related by,

$$\|f\|_{L_2}^2 = 2\pi \sum_{n \in \mathbb{Z}} |f(n)|^2 \quad (1.14)$$

2. (Bad news) The representation of f in terms of sinc functions is not a stable representation, i.e.

$$\sum_{n \in \mathbb{Z}} |\text{sinc}(\pi(t-n))| \approx \sum_{n \in \mathbb{Z}} \frac{1}{|t-n|+1} \rightarrow \text{divergences} \quad (1.15)$$

1.2 Stable Signal Representations²

To fix the instability of the Shannon representation, we assume that the signal is slightly more bandlimited than before

$$\hat{f}(\omega) = 0 \quad \text{for} \quad |\omega| \geq \pi - \delta, \quad \delta > 0, \quad (1.16)$$

and instead of using $\chi_{[-\pi, \pi]}$, we multiply by another function $\hat{g}(\omega)$ which is very similar in form to the characteristic function, but decays at its boundaries in a smoother fashion (i.e. it has more derivatives). A candidate function $\hat{g}$ is sketched in Figure 1.1.



Figure 1.1: Sketch of $\hat{g}$.

Now, it is a property of the Fourier transform that an increased smoothness in one domain translates into a faster decay in the other. Thus, we can fix our instability problem, by choosing $\hat{g}$ so that $\hat{g}$ is smooth and $\hat{g}(\omega) = 1$, $|\omega| \leq \pi - \delta$ and $\hat{g} = 0$, $|\omega| > \pi$. By choosing the smoothness of g suitably large, we can, for any given $m \geq 1$, choose g to satisfy

$$|g(t)| \leq \frac{C}{(|t| + 1)^m} \quad (1.17)$$

for some constant $C > 0$.

Using such a $\hat{g}$, we can rewrite () as

$$\hat{f}(\omega) = F(\omega) \hat{g}(\omega) = \frac{1}{\sqrt{2\pi}} \sum_{n \in \mathbb{Z}} f(n) e^{-in\omega} \hat{g}(\omega). \quad (1.18)$$

Thus, we have the new representation

$$f(t) = \sum_{n \in \mathbb{Z}} f(n) g(t - n), \quad (1.19)$$

where we gain stability from our additional assumption that the signal is bandlimited on $[-\pi - \delta, \pi - \delta]$.

²This content is available online at <<http://cnx.org/content/m15144/1.1/>>.

Does this assumption really hurt? No, not really because if our signal is really bandlimited to $[-\pi, \pi]$ and not $[-\pi - \delta, \pi - \delta]$, we can always take a slightly larger bandwidth, say $[-\lambda\pi, \lambda\pi]$ where λ is a little larger than one, and carry out the same analysis as above. Doing so, would only mean slightly oversampling the signal (small cost).

Recall that in the end we want to convert analog signals into bit streams. Thus far, we have the two representations

$$\begin{aligned} f(t) &= \sum_{n \in \mathbb{Z}} f(n) \operatorname{sinc}(\pi(t - n)), \\ f(t) &= \sum_{n \in \mathbb{Z}} f\left(\frac{n}{\lambda}\right) g(\lambda t - n). \end{aligned} \quad (1.20)$$

Shannon's Theorem tells us that if $f \in B_A$, we should sample f at the Nyquist rate A (which is twice the support of $\hat{f}$) and then take the binary representation of the samples. Our more stable representation says to slightly oversample f and then convert to a binary representation. Both representations offer perfect reconstruction, although in the more stable representation, one is straddled with the additional task of choosing an appropriate λ .

In practical situations, we shall be interested in approximating f on an interval $[-T, T]$ for some $T > 0$ and not for all time. Questions we still want to answer include

1. How many bits do we need to represent f in $B_{A=1}$ on some interval $[-T, T]$ in the norm $L_\infty[-T, T]$?
2. Using this methodology, what is the optimal way of encoding?
3. How is the optimal encoding implemented?

Towards this end, we define

$$B_A := \{f \in L_2(\mathbb{R}) : |\hat{f}(\omega)| = 0, |\omega| \geq A\pi\}. \quad (1.21)$$

Then for any $f \in B_A$, we can write

$$f = \sum_n f\left(\frac{n}{A}\right) \cdot \operatorname{sinc} \pi(A t - n). \quad (1.22)$$

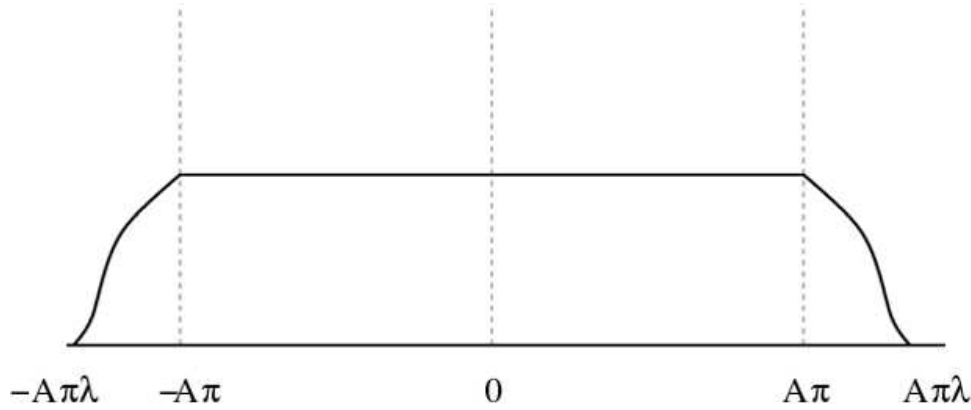


Figure 1.2: Fourier transform of $g_\lambda(\cdot)$.

In other words, samples at $0, \pm \frac{1}{A}, \pm \frac{2}{A}, \dots$ are sufficient to reconstruct f . Recall also that $\operatorname{sinc}(x) = \frac{\sin(x)}{x}$ decays poorly (leading to numerical instability). We can overcome this problem by slight over-sampling. Say we over-sample by a factor $\lambda > 1$. Then, we can write

$$f = \sum_n f\left(\frac{n}{\lambda A}\right) g_\lambda(\lambda A t - n). \quad (1.23)$$

Hence we need samples at $0, \pm \frac{1}{\lambda A}, \pm \frac{2}{\lambda A}$, etc. What is the advantage? Sampling more often than necessary buys us stability because we now have a choice for $g_\lambda(\cdot)$. If we choose $g_\lambda(\cdot)$ infinitely differentiable whose Fourier transform looks as shown in Figure 1.2 we can obtain

$$|g_\lambda(t)| \leq \frac{c_{\lambda,k}}{(1+|t|)^k}, \quad k = 1, 2, \dots \quad (1.24)$$

and therefore $g_\lambda(\cdot)$ decays very fast. In other words, a sample's influence is felt only locally. Note however, that over-sampling generates basis functions that are redundant (linearly dependent), unlike the integer translates of the sinc $(\cdot)$ function.

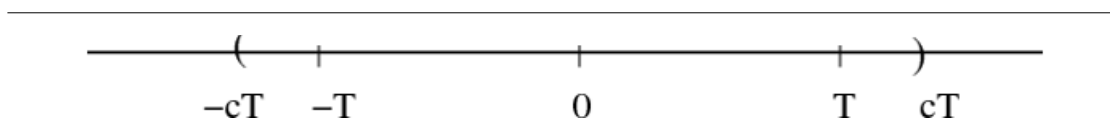


Figure 1.3: To reconstruct signals in $[-T, T]$, the sampling interval is $[-cT, cT]$.

If we restrict our reconstruction to t in the interval $[-T, T]$, we will only need samples only from $[-cT, cT]$, for $c > 1$ (see Figure 1.3), because the distant samples will have little effect on the reconstruction in $[-T, T]$.

1.3 Optimal Encoding³

We shall consider now the encoding of signals on $[-T, T]$ where $T > 0$ is fixed. Ultimately we shall be interested in encoding classes of bandlimited signals like the class B_A . However, we begin the story by considering the more general setting of encoding the elements of any given compact subset K of a normed linear space X . One can determine the best encoding of K by what is known as the Kolmogorov entropy of K in X .

To begin, let us consider an encoder-decoder pair (E, D) . E maps K to a finite stream of bits. D maps a stream of bits to a signal in X . This is illustrated in Figure 1.4. Note that many functions can be mapped onto the same bitstream.

³This content is available online at <http://cnx.org/content/m15139/1.1/>.

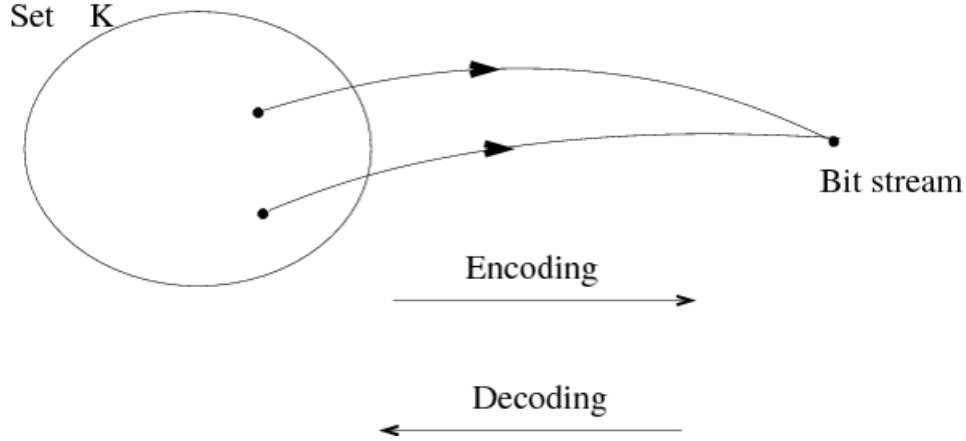


Figure 1.4: Illustration of encoding and decoding.

Define the distortion d for this encoder-decoder by

$$d(K, E, D, X) := \sup_{f \in K} \|f - D(Ef)\|_{\underline{X}}. \quad (1.25)$$

Let $n(K, E) = \sup_{f \in K} \#Ef$ where $\#Ef$ is the number of bits in the bitstream Ef . Thus n is the maximum length of the bitstreams for the various $f \in K$. There are two ways we can define optimal encoding:

1. Prescribe ε , the maximum distortion that we are willing to tolerate. For this ε , find the smallest $n_\varepsilon(K, X) := \inf_{(E, D)} \{n(K, E) : d(K, E, D, X) \leq \varepsilon\}$. This is the smallest bit budget under which we could encode all elements of K to distortion ε .
2. Prescribe N : find the smallest distortion $d(K, E, D, X)$ over all E, D with $n(K, E) \leq N$. This is the best encoding performance possible with a prescribed bit budget.

There is a simple mathematical solution to these two encoding problems based on the notion of Kolmogorov Entropy.

1.4 Kolmogorov Entropy⁴

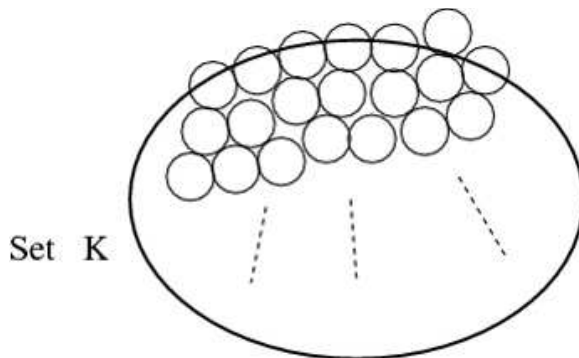


Figure 1.5: Coverings of K by balls of radius ε .

Given $\varepsilon > 0$, and the compact set K , consider all coverings of K by balls of radius ε , as shown in Figure 1.5. In other words,

$$K \subseteq \bigcup_{i=1}^N b(f_i, \varepsilon). \quad (1.26)$$

Let $N_\varepsilon := \inf \{N : \text{over all such covers}\}$. $N_\varepsilon(K)$ is called the covering number of K . Since it depends on X and K , we write it as $N_\varepsilon = N_\varepsilon(K, X)$.

Rule 1.1: Kolmogorov entropy

The Kolmogorov entropy, denoted by $H_\varepsilon(K, X)$, of the compact set K in X is defined as the logarithm of the covering number:

$$H_\varepsilon(K, X) = \log N_\varepsilon(K, X). \quad (1.27)$$

The Kolmogorov entropy solves our problem of optimal encoding in the sense of the following theorem.

Theorem 1.2:

For any compact set $K \subset X$, we have $n_\varepsilon(K, X) = \lceil H_\varepsilon(K, X) \rceil$, where $\lceil \cdot \rceil$ is the ceiling function.

Proof:

Sketch: We can define an encoder-decoder as follows To encode: Say $f \in K$. Just specify which ball it is covered by. Because the number of balls is $N_\varepsilon(K, \overline{X})$, we need at most $\lceil \log N_\varepsilon(K, \overline{X}) \rceil$ bits to specify any such ball.

To decode: Just take the center of the ball specified by the bitstream.

It is now easy to see that this encoder-decoder pair is optimal in either of the senses given above.

□

The above encoder is not practical. However, the Kolmogorov entropy tells us the best performance we can expect from any encoder-decoder pair. Kolmogorov entropy is defined in the deterministic setting. It is the analogue of the Shannon entropy which is defined in a stochastic setting.

⁴This content is available online at <http://cnx.org/content/m15137/1.1/>.

1.5 Optimal Encoding of Bandlimited Signals⁵

We now turn back to the encoding of signals. We are interested in encoding the set

$$B_A(M) = \{f \in B_A : |f(t)| \leq M, t \in \mathbb{R}\} \quad (1.28)$$

where M is arbitrary but fixed. We shall restrict our discussion to the case where distortion is measured in $L_\infty[-T, T]$ where $T > 0$ is arbitrary but fixed. Then, $B_A(M)$ is a compact subset of L_∞ : $B_A(M) \subseteq L_\infty[-T, T]$.

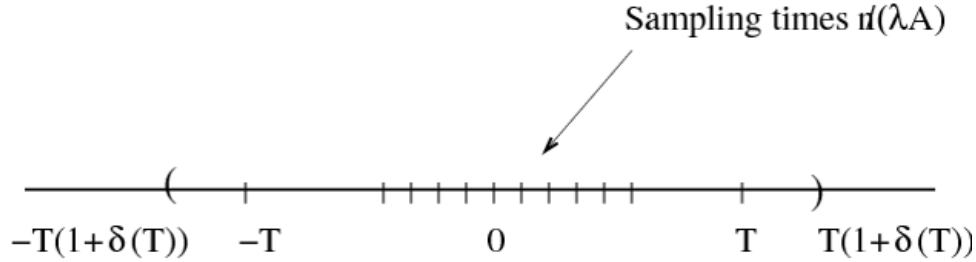


Figure 1.6: Sample points $\frac{n}{\lambda A}$ are chosen in the interval $[-T(1+\delta), T(1+\delta)]$.

We shall sketch how one can construct an asymptotically optimal encoder/decoder for B_A . The details for this construction can be found in [1].

We know $f(\omega) = 0$ for $|\omega| \geq A\pi$, and $|f| \leq M$. How can we encode f in practice? We begin by choosing $\lambda = \lambda(T) > 1$ (see Figure 1.6) which will represent a slight oversampling factor we shall utilize. Given a target distortion $\varepsilon > 0$, we choose k so that $2^{-k-1} < \varepsilon \leq 2^{-k}$. Given f , we shall encode f by first taking samples $f(\frac{n}{\lambda A})$ for $\frac{n}{\lambda A} \in [-T(1+\delta), T(1+\delta)]$ where $\delta(T) > 0$. In other words, we sample f on a slightly larger interval than $[-T, T]$. For each sample $f(\frac{n}{\lambda A})$, we shall use the first $k + k_0(T)$ bits of its binary expansion. In other words, our encoder takes f and the samples $f(\frac{n}{\lambda A})$ and then assigns to $f(\frac{n}{\lambda A})$ the first $k + k_0(T)$ bits of this number.

To decode, the receiver would take the bits and construct the approximation $\bar{f}(\frac{n}{\lambda A})$ to $f(\frac{n}{\lambda A})$ from the bits provided. Notice that we have the accuracy

$$\left| f\left(\frac{n}{\lambda A}\right) - \bar{f}\left(\frac{n}{\lambda A}\right) \right| \leq 2^{-k-k_0} \cdot M. \quad (1.29)$$

We utilize the function g_λ satisfying (1.28) to define

$$\bar{f}(t) = \sum_{n \in N_T} \bar{f}\left(\frac{n}{\lambda A}\right) g_\lambda(\lambda A t - n), \quad (1.30)$$

where

$$N_T := \{n : -T(1+\delta) \leq \frac{n}{\lambda A} \leq T(1+\delta)\}. \quad (1.31)$$

We then have

$$\begin{aligned} |f(t) - \bar{f}(t)| &\leq \sum_{n \in N_T} \left| f\left(\frac{n}{\lambda A}\right) - \bar{f}\left(\frac{n}{\lambda A}\right) \right| \cdot |g_\lambda(\lambda A t - n)| \\ &\quad + \sum_{|\frac{n}{\lambda A}| > T(1+\delta)} \left| f\left(\frac{n}{\lambda A}\right) \right| \cdot |g_\lambda(\lambda A t - n)| \end{aligned} \quad (1.32)$$

⁵This content is available online at <http://cnx.org/content/m15140/1.2/>.

The term $|f(\frac{n}{\lambda A}) - \bar{f}(\frac{n}{\lambda A})|$ that appears in the first summation in ((1.32)) is bounded by $M \cdot 2^{-k-k_0}$. The term $|f(\frac{n}{\lambda A})|$ that appears in the second summation in the same equation is bounded by M . Therefore,

$$|f(t) - \bar{f}(t)| \leq \sum_{n \in N_T} M \cdot 2^{-k-k_0} \cdot |g_\lambda(\lambda At - n)| + \sum_{|\frac{n}{\lambda A}| > T(1+\delta)} M \cdot |g_\lambda(\lambda At - n)| =: S_1 + S_2 \quad (1.33)$$

We can estimate S_1 by

$$\begin{aligned} S_1 &= \sum_{n \in N_T} M \cdot 2^{-k-k_0} \cdot |g_\lambda(\lambda At - n)| \\ &\leq M \cdot 2^{-k-k_0} \cdot \sum_n |g_\lambda(\lambda At - n)| \\ &\leq M \cdot C_0(\lambda) \cdot 2^{-k-k_0} \quad (\text{because } g(\cdot) \text{ decays fast}) \end{aligned} \quad (1.34)$$

Therefore, if we choose k_0 sufficiently large, then $S_1 \leq M \cdot C_0(\lambda) \cdot 2^{-k-k_0} \leq \frac{\varepsilon}{2}$. The second summation S_2 can also be bounded by $\varepsilon/2$ by using the fast decay of the function g_λ (see ()).

To make the encoder/decoder specific we need to precisely define δ and λ . It turns out that the best choices (in terms of bit rate performance on the class B_A) depend on T . But $\delta_T \rightarrow 0$ and $\lambda_T \rightarrow 1$ as $T \rightarrow \infty$. Recall that Shannon sampling requires $2T\lambda A$ samples. Since our encoder/decoder uses $k + k_0$ bits per sample, the total number of bits is $(k + k_0) \cdot 2\lambda AT(1 + \delta)$, and so coding will require roughly k bits per Shannon sample.

This encoder/decoder can be proven to be optimal in the sense of averaged performance as we shall now describe. The average of performance of optimal encoding is defined by

$$\lim_{T \rightarrow \infty} \frac{n_\varepsilon(B_A(M), L_\infty[-T, T])}{2T} \quad (1.35)$$

If we replace the optimal bit rate n_ε in ((1.35)) by the number of bits required by our encoder/decoder then the resulting limit will be the same as that in ((1.35)).

In summary, to encode band limited signals on an interval $[-T, T]$, an optimal strategy is to sample at a slightly higher rate than Nyquist and on a slightly large interval than $[-T, T]$. Each sample should then be quantized by using the binary expansion of the sample. In this way, for an investment of k bits per Nyquist rate sample, we get a distortion of 2^{-k} .

To get a feel for the number of bits required by such an encoder, let us say $A = 10^6$ (signals band limited to 1Mhz). Say $T = 24$ hours $\approx 10^5$ seconds, and $k = 10$ bits. Then, $A \cdot k \cdot 2T = 10^6 \cdot 10 \cdot 10^5 = 10^{12}$ bits. This is too BIG!

The above encoding is known as Pulse Coded Modulation (PCM). In practice, people frequently use another encoder called Sigma-Delta Modulation. Instead of oversampling just slightly, Sigma Delta over samples a lot and then assign only one (or a few) bits per sample.

Why is Sigma-Delta preferred to PCM in practice? There are two reasons commonly given:

1. Getting accurate samples, quantization, etc. is not practical because of noise. For better accuracy, we need more expensive hardware.
2. Noise shaping. In Sigma-Delta, the distortion is higher but the distortion is spread over frequencies outside of the desired range.

In PCM, the distortion decays exponentially (like 2^{-k}), whereas for Sigma-Delta, the distortion decays like a polynomial (like $\frac{1}{k^m}$). Although the distortion decays faster in PCM, the distortion in Sigma-Delta is spread outside the desired frequency range.

Chapter 2

Sparsity and Compressibility

2.1 Introduction to vector spaces¹

For much of its history, signal processing has focused on signals produced by physical systems. Many natural and man-made systems can be modeled as linear. Thus, it is natural to consider signal models that complement this kind of linear structure. This notion has been incorporated into modern signal processing by modeling signals as *vectors* living in an appropriate *vector space*. This captures the linear structure that we often desire, namely that if we add two signals together then we obtain a new, physically meaningful signal. Moreover, vector spaces allow us to apply intuitions and tools from geometry in $\mathbb{R}^3$, such as lengths, distances, and angles, to describe and compare signals of interest. This is useful even when our signals live in high-dimensional or infinite-dimensional spaces.

Throughout this course², we will treat signals as real-valued functions having domains that are either continuous or discrete, and either infinite or finite. These assumptions will be made clear as necessary in each chapter. In this course, we will assume that the reader is relatively comfortable with the key concepts in vector spaces. We now provide only a brief review of some of the key concepts in vector spaces that will be required in developing the theory of compressive sensing³ (CS). For a more thorough review of vector spaces see this introductory course in Digital Signal Processing⁴.

We will typically be concerned with *normed vector spaces*, i.e., vector spaces endowed with a *norm*. In the case of a discrete, finite domain, we can view our signals as vectors in an N -dimensional Euclidean space, denoted by $\mathbb{R}^N$. When dealing with vectors in $\mathbb{R}^N$, we will make frequent use of the ℓ_p norms, which are defined for $p \in [1, \infty]$ as

$$\|x\|_p = \begin{cases} \left(\sum_{i=1}^N |x_i|^p \right)^{\frac{1}{p}}, & p \in [1, \infty); \\ \max_{i=1,2,\dots,N} |x_i|, & p = \infty. \end{cases} \quad (2.1)$$

In Euclidean space we can also consider the standard *inner product* in $\mathbb{R}^N$, which we denote

$$\langle x, z \rangle = z^T x = \sum_{i=1}^N x_i z_i. \quad (2.2)$$

This inner product leads to the ℓ_2 norm: $\|x\|_2 = \sqrt{\langle x, x \rangle}$.

In some contexts it is useful to extend the notion of ℓ_p norms to the case where $p < 1$. In this case, the “norm” defined in (2.1) fails to satisfy the triangle inequality, so it is actually a quasinorm. We will also

¹This content is available online at <http://cnx.org/content/m37167/1.6/>.

²*An Introduction to Compressive Sensing* <http://cnx.org/content/col11133/latest/>

³“Introduction to compressive sensing” <http://cnx.org/content/m37172/latest/>

⁴*Digital Signal Processing* <http://cnx.org/content/col11172/latest/>

make frequent use of the notation $\|x\|_0 := |\text{supp}(x)|$, where $\text{supp}(x) = \{i : x_i \neq 0\}$ denotes the support of x and $|\text{supp}(x)|$ denotes the cardinality of $\text{supp}(x)$. Note that $\|\cdot\|_0$ is not even a quasinorm, but one can easily show that

$$\lim_{p \rightarrow 0} \|x\|_p = |\text{supp}(x)|, \quad (2.3)$$

justifying this choice of notation. The ℓ_p (quasi-)norms have notably different properties for different values of p . To illustrate this, in Figure 2.1 we show the unit sphere, i.e., $\{x : \|x\|_p = 1\}$, induced by each of these norms in $\mathbb{R}^2$. Note that for $p < 1$ the corresponding unit sphere is nonconvex (reflecting the quasinorm's violation of the triangle inequality).

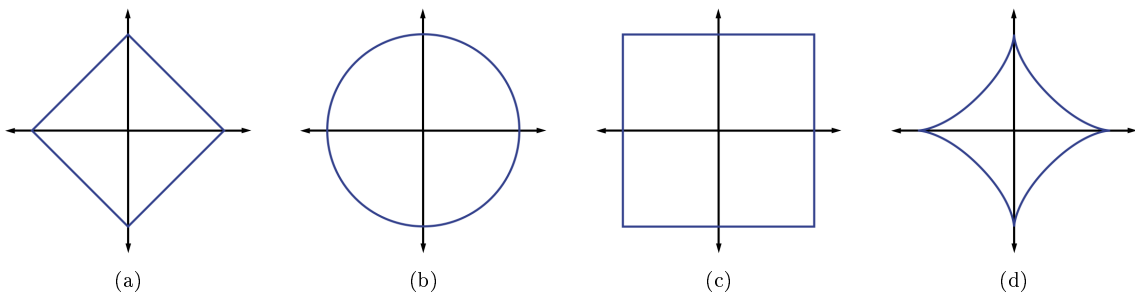


Figure 2.1: Unit spheres in $\mathbb{R}^2$ for the ℓ_p norms with $p = 1, 2, \infty$, and for the ℓ_p quasinorm with $p = \frac{1}{2}$. (a) Unit sphere for ℓ_1 norm (b) Unit sphere for ℓ_2 norm (c) Unit sphere for ℓ_∞ norm (d) Unit sphere for ℓ_p quasinorm

We typically use norms as a measure of the strength of a signal, or the size of an error. For example, suppose we are given a signal $x \in \mathbb{R}^2$ and wish to approximate it using a point in a one-dimensional affine space A . If we measure the approximation error using an ℓ_p norm, then our task is to find the $\hat{x} \in A$ that minimizes $\|x - \hat{x}\|_p$. The choice of p will have a significant effect on the properties of the resulting approximation error. An example is illustrated in Figure 2.2. To compute the closest point in A to x using each ℓ_p norm, we can imagine growing an ℓ_p sphere centered on x until it intersects with A . This will be the point $\hat{x} \in A$ that is closest to x in the corresponding ℓ_p norm. We observe that larger p tends to spread out the error more evenly among the two coefficients, while smaller p leads to an error that is more unevenly distributed and tends to be sparse. This intuition generalizes to higher dimensions, and plays an important role in the development of CS theory.

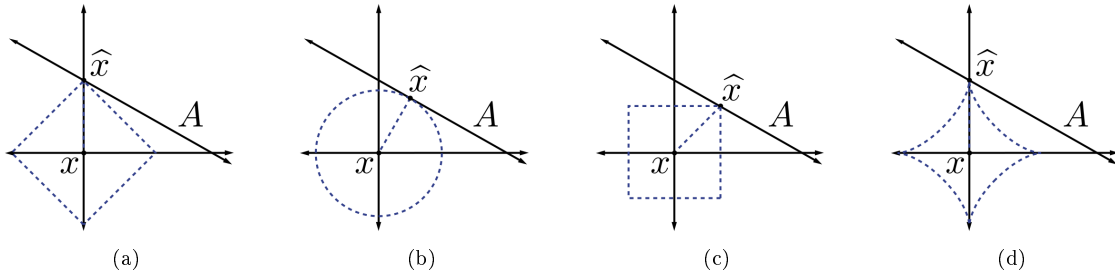


Figure 2.2: Best approximation of a point in $\mathbb{R}^2$ by a one-dimensional subspace using the ℓ_p norms for $p = 1, 2, \infty$, and the ℓ_p quasinorm with $p = \frac{1}{2}$. (a) Approximation in ℓ_1 norm (b) Approximation in ℓ_2 norm (c) Approximation in ℓ_∞ norm (d) Approximation in ℓ_p quasinorm

2.2 Bases and frames⁵

A set $\Psi = \{\psi_i\}_{i \in \mathcal{I}}$ is called a basis for a finite-dimensional vector space (Section 2.1) $\mathcal{V}$ if the vectors in the set span $\mathcal{V}$ and are linearly independent. This implies that each vector in the space can be represented as a linear combination of this (smaller, except in the trivial case) set of basis vectors in a unique fashion. Furthermore, the coefficients of this linear combination can be found by the inner product of the signal and a dual set of vectors. In discrete settings, we will only consider real finite-dimensional Hilbert spaces where $\mathcal{V} = \mathbb{R}^N$ and $\mathcal{I} = \{1, \dots, N\}$.

Mathematically, any signal $x \in \mathbb{R}^N$ may be expressed as,

$$x = \sum_{i \in \mathcal{I}} a_i \tilde{\psi}_i, \quad (2.4)$$

where our coefficients are computed as $a_i = \langle x, \psi_i \rangle$ and $\{\tilde{\psi}_i\}_{i \in \mathcal{I}}$ are the vectors that constitute our dual basis. Another way to denote our basis and its dual is by how they operate on x . Here, we call our dual basis $\tilde{\Psi}$ our *synthesis basis* (used to reconstruct our signal by (2.4)) and Ψ is our *analysis basis*.

An orthonormal basis (ONB) is defined as a set of vectors $\Psi = \{\psi_i\}_{i \in \mathcal{I}}$ that form a basis and whose elements are orthogonal and unit norm. In other words, $\langle \psi_i, \psi_j \rangle = 0$ if $i \neq j$ and one otherwise. In the case of an ONB, the *synthesis basis* is simply the Hermitian adjoint of *analysis basis* ($\tilde{\Psi} = \Psi^T$).

It is often useful to generalize the concept of a basis to allow for sets of possibly linearly dependent vectors, resulting in what is known as a *frame*. More formally, a frame is a set of vectors $\{\Psi_i\}_{i=1}^n$ in $\mathbb{R}^d$, $d < n$ corresponding to a matrix $\Psi \in \mathbb{R}^{d \times n}$, such that for all vectors $x \in \mathbb{R}^d$,

$$A\|x\|_2^2 \leq \|\Psi^T x\|_2^2 \leq B\|x\|_2^2 \quad (2.5)$$

with $0 < A \leq B < \infty$. Note that the condition $A > 0$ implies that the rows of Ψ must be linearly independent. When A is chosen as the largest possible value and B as the smallest for these inequalities to hold, then we call them the (*optimal*) *frame bounds*. If A and B can be chosen as $A = B$, then the frame is called *A-tight*, and if $A = B = 1$, then Ψ is a *Parseval frame*. A frame is called *equal-norm*, if there exists some $\lambda > 0$ such that $\|\Psi_i\|_2 = \lambda$ for all $i = 1, \dots, N$, and it is *unit-norm* if $\lambda = 1$. Note also that while the concept of a frame is very general and can be defined in infinite-dimensional spaces, in the case where Ψ is a $d \times N$ matrix A and B simply correspond to the smallest and largest eigenvalues of $\Psi\Psi^T$, respectively.

⁵This content is available online at <http://cnx.org/content/m37165/1.6/>.

Frames can provide richer representations of data due to their redundancy: for a given signal x , there exist infinitely many coefficient vectors α such that $x = \Psi\alpha$. In particular, each choice of a dual frame $\tilde{\Psi}$ provides a different choice of a coefficient vector α . More formally, any frame satisfying

$$\Psi\tilde{\Psi}^T = \tilde{\Psi}\Psi^T = I \quad (2.6)$$

is called an (alternate) dual frame. The particular choice $\tilde{\Psi} = (\Psi\Psi^T)^{-1}\Psi$ is referred to as the *canonical dual frame*. It is also known as the Moore-Penrose pseudoinverse. Note that since $A > 0$ requires Ψ to have linearly independent rows, we ensure that $\Psi\Psi^T$ is invertible, so that $\tilde{\Psi}$ is well-defined. Thus, one way to obtain a set of feasible coefficients is via

$$\alpha_d = \Psi^T(\Psi\Psi^T)^{-1}x. \quad (2.7)$$

One can show that this sequence is the smallest coefficient sequence in ℓ_2 norm, i.e., $\|\alpha_d\|_2 \leq \|\alpha\|_2$ for all α such that $x = \Psi\alpha$.

Finally, note that in the sparse approximation (Section 2.3) literature, it is also common for a basis or frame to be referred to as a *dictionary* or *overcomplete dictionary* respectively, with the dictionary elements being called *atoms*.

2.3 Sparse representations⁶

Transforming a signal to a new basis or frame (Section 2.2) may allow us to represent a signal more concisely. The resulting compression is useful for reducing data storage and data transmission, which can be quite expensive in some applications. Hence, one might wish to simply transmit the analysis coefficients obtained in our basis or frame expansion instead of its high-dimensional correlate. In cases where the number of non-zero coefficients is small, we say that we have a sparse representation. Sparse signal models allow us to achieve high rates of compression and in the case of compressive sensing⁷, we may use the knowledge that our signal is sparse in a known basis or frame to recover our original signal from a small number of measurements. For sparse data, only the non-zero coefficients need to be stored or transmitted in many cases; the rest can be assumed to be zero).

Mathematically, we say that a signal x is K -sparse when it has at most K nonzeros, i.e., $\|x\|_0 \leq K$. We let

$$\Sigma_K = \{x : \|x\|_0 \leq K\} \quad (2.8)$$

denote the set of all K -sparse signals. Typically, we will be dealing with signals that are not themselves sparse, but which admit a sparse representation in some basis Ψ . In this case we will still refer to x as being K -sparse, with the understanding that we can express x as $x = \Psi\alpha$ where $\|\alpha\|_0 \leq K$.

Sparsity has long been exploited in signal processing and approximation theory for tasks such as compression [50], [145], [164] and denoising [54], and in statistics and learning theory as a method for avoiding overfitting [180]. Sparsity also figures prominently in the theory of statistical estimation and model selection [97], [167], in the study of the human visual system [142], and has been exploited heavily in image processing tasks, since the multiscale wavelet transform [124] provides nearly sparse representations for natural images. Below, we briefly describe some one-dimensional (1-D) and two-dimensional (2-D) examples.

2.3.1 1-D signal models

We will now present an example of three basis expansions that yield different levels of sparsity for the same signal. A simple periodic signal is sampled and represented as a periodic train of weighted impulses (see Figure 2.3). One can interpret sampling as a basis expansion where our elements in our basis are impulses

⁶This content is available online at <http://cnx.org/content/m37168/1.5/>.

⁷"Introduction to compressive sensing" <http://cnx.org/content/m37172/latest/>

placed at periodic points along the time axis. We know that in this case, our dual basis consists of sinc functions used to reconstruct our signal from discrete-time samples. This representation contains many non-zero coefficients, and due to the signal's periodicity, there are many redundant measurements. Representing the signal in the Fourier basis, on the other hand, requires only two non-zero basis vectors, scaled appropriately at the positive and negative frequencies (see Figure 2.3). Driving the number of coefficients needed even lower, we may apply the discrete cosine transform (DCT) to our signal, thereby requiring only a single non-zero coefficient in our expansion (see Figure 2.3). The DCT equation is $X_k = \sum_{n=0}^{N-1} x_n \cos\left(\frac{\pi}{N}\left(n + \frac{1}{2}\right)k\right)$ with $k = 0, \dots, N-1$, x_n the input signal, and N the length of the transform.

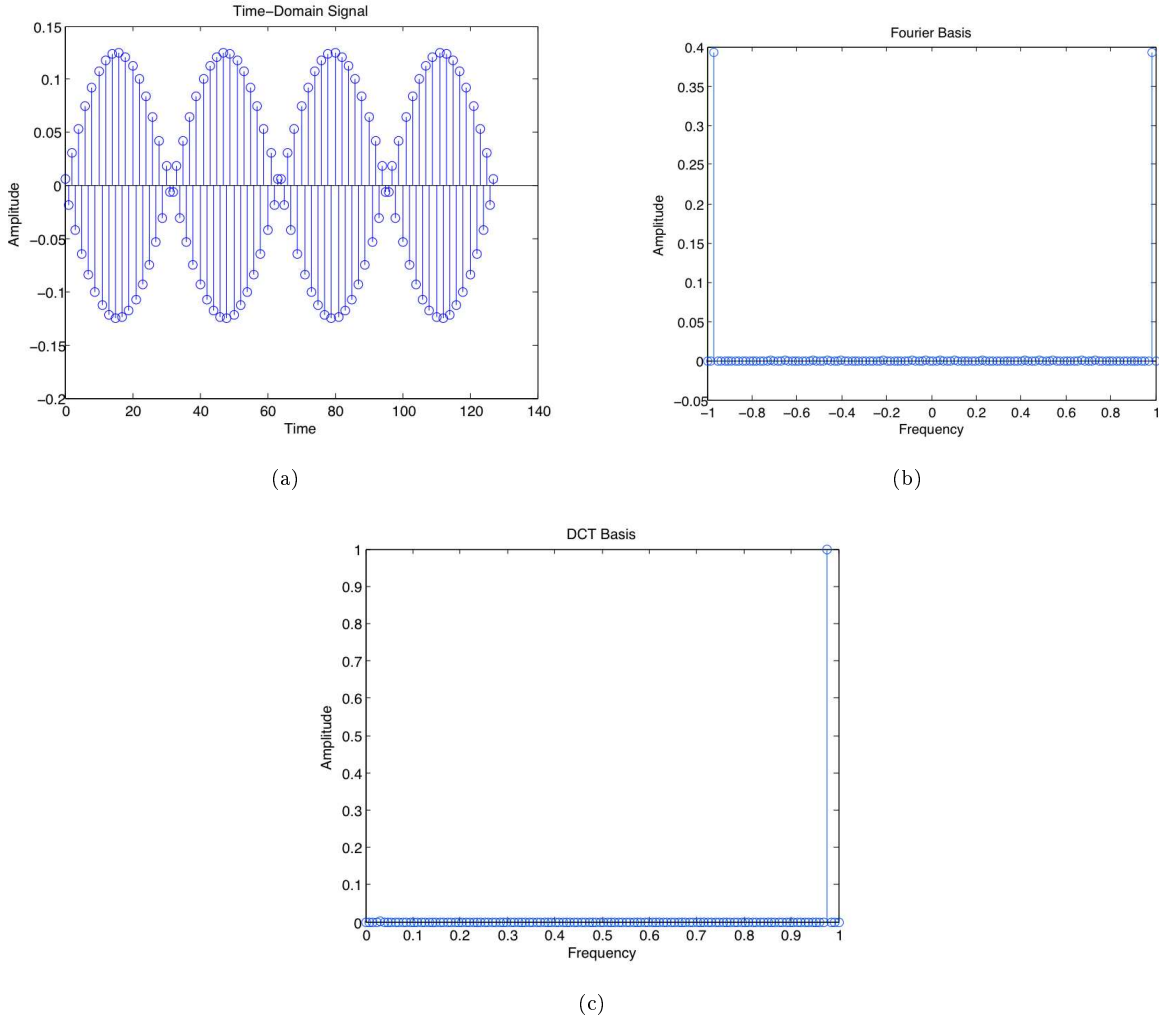


Figure 2.3: Cosine signal in three representations: (a) Train of impulses (b) Fourier basis (c) DCT basis

2.3.2 2-D signal models

This same concept can be extended to 2-D signals as well. For instance, a binary picture of a nighttime sky is sparse in the standard pixel domain because most of the pixels are zero-valued black pixels. Likewise, natural images are characterized by large smooth or textured regions and relatively few sharp edges. Signals with this structure are known to be very nearly sparse when represented using a multiscale wavelet transform [124]. The wavelet transform consists of recursively dividing the image into its low- and high-frequency components. The lowest frequency components provide a coarse scale approximation of the image, while the higher frequency components fill in the detail and resolve edges. What we see when we compute a wavelet transform of a typical natural image, as shown in Figure 2.4, is that most coefficients are very small. Hence, we can obtain a good approximation of the signal by setting the small coefficients to zero, or *thresholding* the coefficients, to obtain a K -sparse representation. When measuring the approximation error using an ℓ_p norm, this procedure yields the *best K -term approximation* of the original signal, i.e., the best approximation of the signal using only K basis elements.⁸

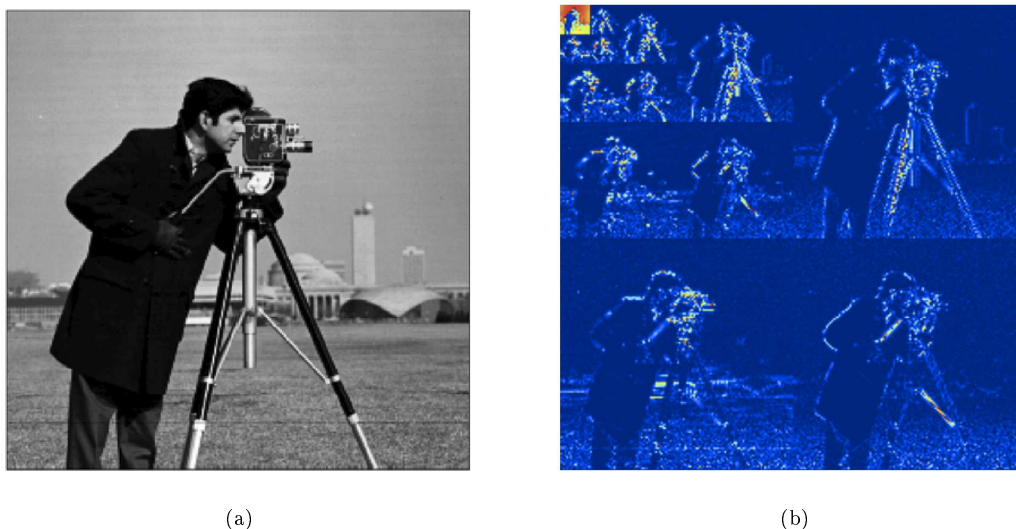


Figure 2.4: Sparse representation of an image via a multiscale wavelet transform. (a) Original image. (b) Wavelet representation. Large coefficients are represented by light pixels, while small coefficients are represented by dark pixels. Observe that most of the wavelet coefficients are close to zero.

Sparsity results through this decomposition because in most natural images most pixel values vary little from their neighbors. Areas with little contrast difference can be represent with low frequency wavelets. Low frequency wavelets are created through stretching a mother wavelet and thus expanding it in space. On the other hand, discontinuities, or edges in the picture, require high frequency wavelets, which are created through compacting a mother wavelet. At the same time, the transitions are generally limited in space, mimicking the properties of the high frequency compacted wavelet. See "Compressible signals" (Section 2.4) for an example.

⁸Thresholding yields the best K -term approximation of a signal with respect to an orthonormal basis. When redundant frames are used, we must rely on sparse approximation algorithms like those described later in this course [75], [124].

2.4 Compressible signals⁹

2.4.1 Compressibility and K -term approximation

An important assumption used in the context of compressive sensing¹⁰ (CS) is that signals exhibit a degree of structure. So far the only structure we have considered is sparsity (Section 2.3), i.e., the number of non-zero values the signal has when representation in an orthonormal basis (Section 2.2) Ψ . The signal is considered sparse if it has only a few nonzero values in comparison with its overall length.

Few structured signals are truly sparse; rather they are compressible. A signal is *compressible* if its sorted coefficient magnitudes in Ψ decay rapidly. To consider this mathematically, let x be a signal which is compressible in the basis Ψ :

$$x = \Psi\alpha, \quad (2.9)$$

where α are the coefficients of x in the basis Ψ . If x is compressible, then the magnitudes of the sorted coefficients α_s observe a power law decay:

$$|\alpha_s| \leq C_1 s^{-q}, s = 1, 2, \dots \quad (2.10)$$

We define a signal as being compressible if it obeys this power law decay. The larger q is, the faster the magnitudes decay, and the more compressible a signal is. Figure 2.5 shows images that are compressible in different bases.

⁹This content is available online at <<http://cnx.org/content/m37166/1.5/>>.

¹⁰"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

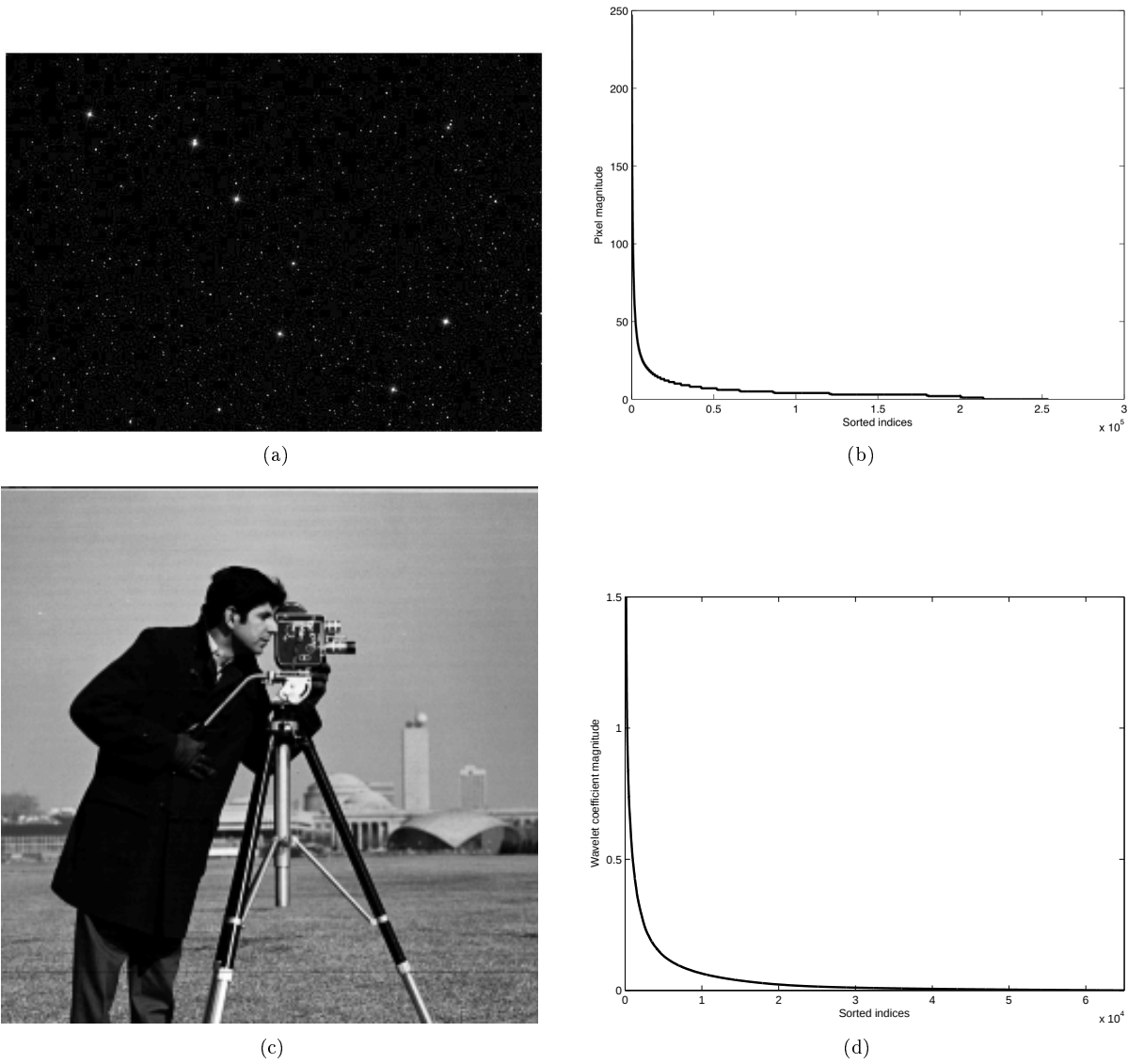


Figure 2.5: The image in the upper left is a signal that is compressible in space. When the pixel values are sorted from largest to smallest, there is a sharp descent. The image in the lower left is not compressible in space, but it is compressible in wavelets since its wavelet coefficients exhibit a power law decay.

Because the magnitudes of their coefficients decay so rapidly, compressible signals can be represented well by $K \ll N$ coefficients. The best K -term approximation of a signal is the one in which the K largest coefficients are kept, with the rest being zero. The error between the true signal and its K term approximation is denoted the K -term approximation error $\sigma_K(x)$, defined as

$$\sigma_K(x) = \arg \min_{\alpha \in \Sigma_K} \|x - \Psi\alpha\|_2. \quad (2.11)$$

For compressible signals, we can establish a bound with power law decay as follows:

$$\sigma_K(x) \leq C_2 K^{1/2-s}. \quad (2.12)$$

In fact, one can show that $\sigma_K(x)_2$ will decay as K^{-r} if and only if the sorted coefficients α_i decay as $i^{-r+1/2}$ [51]. Figure 2.6 shows an image and its K -term approximation.



(a)



(b)

Figure 2.6: Sparse approximation of a natural image. (a) Original image. (b) Approximation of image obtained by keeping only the largest 10% of the wavelet coefficients. Because natural images are compressible in a wavelet domain, approximating this image in terms of its largest wavelet coefficients maintains good fidelity.

2.4.2 Compressibility and ℓ_p spaces

A signal's compressibility is related to the ℓ_p space to which the signal belongs. An infinite sequence $x(n)$ is an element of an ℓ_p space for a particular value of p if and only if its ℓ_p norm is finite:

$$\|x\|_p = \left(\sum_i |x_i|^p \right)^{\frac{1}{p}} < \infty. \quad (2.13)$$

The smaller p is, the faster the sequence's values must decay in order to converge so that the norm is bounded. In the limiting case of $p = 0$, the "norm" is actually a pseudo-norm and counts the number of non-zero values. As p decreases, the size of its corresponding ℓ_p space also decreases. Figure 2.7 shows various ℓ_p unit balls (all sequences whose ℓ_p norm is 1) in 3 dimensions.

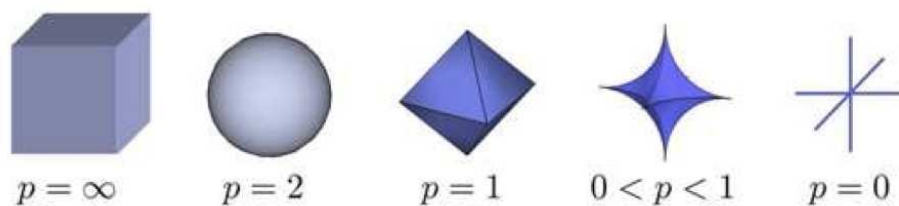


Figure 2.7: As the value of p decreases, the size of the corresponding ℓ_p space also decreases. This can be seen visually when comparing the size of the spaces of signals, in three dimensions, for which the ℓ_p norm is less than or equal to one. The volume of these ℓ_p “balls” decreases with p .

Suppose that a signal is sampled infinitely finely, and call it $x[n]$. In order for this sequence to have a bounded ℓ_p norm, its coefficients must have a power-law rate of decay with $q > 1/p$. Therefore a signal which is in an ℓ_p space with $p \leq 1$ obeys a power law decay, and is therefore compressible.

Chapter 3

Compressive Sensing

3.1 Sensing matrix design¹

In order to make the discussion more concrete, we will restrict our attention to the standard finite-dimensional compressive sensing² (CS) model. Specifically, given a signal $x \in \mathbb{R}^N$, we consider measurement systems that acquire M linear measurements. We can represent this process mathematically as

$$y = \Phi x, \tag{3.1}$$

where Φ is an $M \times N$ matrix and $y \in \mathbb{R}^M$. The matrix Φ represents a *dimensionality reduction*, i.e., it maps $\mathbb{R}^N$, where N is generally large, into $\mathbb{R}^M$, where M is typically much smaller than N . Note that in the standard CS framework we assume that the measurements are *non-adaptive*, meaning that the rows of Φ are fixed in advance and do not depend on the previously acquired measurements. In certain settings adaptive measurement schemes can lead to significant performance gains.

Note that although the standard CS framework assumes that x is a finite-length vector with a discrete-valued index (such as time or space), in practice we will often be interested in designing measurement systems for acquiring continuously-indexed signals such as continuous-time signals or images. For now we will simply think of x as a finite-length window of Nyquist-rate samples, and we temporarily ignore the issue of how to directly acquire compressive measurements without first sampling at the Nyquist rate.

There are two main theoretical questions in CS. First, how should we design the sensing matrix Φ to ensure that it preserves the information in the signal x ? Second, how can we recover the original signal x from measurements y ? In the case where our data is sparse (Section 2.3) or compressible (Section 2.4), we will see that we can design matrices Φ with $M \ll N$ that ensure that we will be able to recover (Section 4.1) the original signal accurately and efficiently using a variety of practical algorithms (Section 4.6).

We begin in this part of the course³ by first addressing the question of how to design the sensing matrix Φ . Rather than directly proposing a design procedure, we instead consider a number of desirable properties that we might wish Φ to have (including the null space property (Section 3.2), the restricted isometry property (Section 3.3), and bounded coherence (Section 3.6)). We then provide some important examples of matrix constructions (Section 3.5) that satisfy these properties.

¹This content is available online at <http://cnx.org/content/m37169/1.6/>.

²"Introduction to compressive sensing" <http://cnx.org/content/m37172/latest/>

³*An Introduction to Compressive Sensing* <http://cnx.org/content/col11133/latest/>

3.2 Null space conditions⁴

A natural place to begin in establishing conditions on Φ in the context of designing a sensing matrix (Section 3.1) is by considering the null space of Φ , denoted

$$\mathcal{N}(\Phi) = \{z : \Phi z = 0\}. \quad (3.2)$$

If we wish to be able to recover *all* sparse (Section 2.3) signals x from the measurements Φx , then it is immediately clear that for any pair of distinct vectors $x, x' \in \Sigma_K = \{x : \|x\|_0 \leq K\}$, we must have $\Phi x \neq \Phi x'$, since otherwise it would be impossible to distinguish x from x' based solely on the measurements y . More formally, by observing that if $\Phi x = \Phi x'$ then $\Phi(x - x') = 0$ with $x - x' \in \Sigma_{2K}$, we see that Φ uniquely represents all $x \in \Sigma_K$ if and only if $\mathcal{N}(\Phi)$ contains no vectors in Σ_{2K} . There are many equivalent ways of characterizing this property; one of the most common is known as the *spark* [59].

3.2.1 The spark

Definition 3.1:

The spark of a given matrix Φ is the smallest number of columns of Φ that are linearly dependent.

This definition allows us to pose the following straightforward guarantee.

Theorem 3.1: (Corollary 1 of [59])

For any vector $y \in \mathbb{R}^M$, there exists at most one signal $x \in \Sigma_K$ such that $y = \Phi x$ if and only if $\text{spark}(\Phi) > 2K$.

Proof:

We first assume that, for any $y \in \mathbb{R}^M$, there exists at most one signal $x \in \Sigma_K$ such that $y = \Phi x$. Now suppose for the sake of a contradiction that $\text{spark}(\Phi) \leq 2K$. This means that there exists some set of at most $2K$ columns that are linearly independent, which in turn implies that there exists an $h \in \mathcal{N}(\Phi)$ such that $h \in \Sigma_{2K}$. In this case, since $h \in \Sigma_{2K}$ we can write $h = x - x'$, where $x, x' \in \Sigma_K$. Thus, since $h \in \mathcal{N}(\Phi)$ we have that $\Phi(x - x') = 0$ and hence $\Phi x = \Phi x'$. But this contradicts our assumption that there exists at most one signal $x \in \Sigma_K$ such that $y = \Phi x$. Therefore, we must have that $\text{spark}(\Phi) > 2K$.

Now suppose that $\text{spark}(\Phi) > 2K$. Assume that for some y there exist $x, x' \in \Sigma_K$ such that $y = \Phi x = \Phi x'$. We therefore have that $\Phi(x - x') = 0$. Letting $h = x - x'$, we can write this as $\Phi h = 0$. Since $\text{spark}(\Phi) > 2K$, all sets of up to $2K$ columns of Φ are linearly independent, and therefore $h = 0$. This in turn implies $x = x'$, proving the theorem.

It is easy to see that $\text{spark}(\Phi) \in [2, M + 1]$. Therefore, Theorem 3.1, (Corollary 1 of [59]), p. 22 yields the requirement $M \geq 2K$.

3.2.2 The null space property

When dealing with *exactly* sparse vectors, the spark provides a complete characterization of when sparse recovery is possible. However, when dealing with *approximately* sparse (Section 2.4) signals we must introduce somewhat more restrictive conditions on the null space of Φ [36]. Roughly speaking, we must also ensure that $\mathcal{N}(\Phi)$ does not contain any vectors that are too compressible in addition to vectors that are sparse. In order to state the formal definition we define the following notation that will prove to be useful throughout much of this course⁵. Suppose that $\Lambda \subset \{1, 2, \dots, N\}$ is a subset of indices and let $\Lambda^c = \{1, 2, \dots, N\} \setminus \Lambda$. By x_Λ we typically mean the length N vector obtained by setting the entries of x indexed by Λ^c to zero.

⁴This content is available online at <http://cnx.org/content/m37170/1.6/>.

⁵An Introduction to Compressive Sensing <http://cnx.org/content/col11133/latest/>

Similarly, by Φ_Λ we typically mean the $M \times N$ matrix obtained by setting the columns of Φ indexed by Λ^c to zero.⁶

Definition 3.2:

A matrix Φ satisfies the *null space property* (NSP) of order K if there exists a constant $C > 0$ such that,

$$\|h_\Lambda\|_2 \leq C \frac{\|h_{\Lambda^c}\|_1}{\sqrt{K}} \quad (3.3)$$

holds for all $h \in \mathcal{N}(\Phi)$ and for all Λ such that $|\Lambda| \leq K$.

The NSP quantifies the notion that vectors in the null space of Φ should not be too concentrated on a small subset of indices. For example, if a vector h is exactly K -sparse, then there exists a Λ such that $\|h_{\Lambda^c}\|_1 = 0$ and hence (3.3) implies that $h_\Lambda = 0$ as well. Thus, if a matrix Φ satisfies the NSP then the only K -sparse vector in $\mathcal{N}(\Phi)$ is $h = 0$.

To fully illustrate the implications of the NSP in the context of sparse recovery, we now briefly discuss how we will measure the performance of sparse recovery algorithms when dealing with general non-sparse x . Towards this end, let $\Delta : \mathbb{R}^M \rightarrow \mathbb{R}^N$ represent our specific recovery method. We will focus primarily on guarantees of the form

$$\|\Delta(\Phi x) - x\|_2 \leq C \frac{\sigma_K(x)_1}{\sqrt{K}} \quad (3.4)$$

for all x , where we recall that

$$\sigma_K(x)_p = \min_{\hat{x} \in \Sigma_K} \|x - \hat{x}\|_p. \quad (3.5)$$

This guarantees exact recovery of all possible K -sparse signals, but also ensures a degree of robustness to non-sparse signals that directly depends on how well the signals are approximated by K -sparse vectors. Such guarantees are called *instance-optimal* since they guarantee optimal performance for each instance of x [36]. This distinguishes them from guarantees that only hold for some subset of possible signals, such as sparse or compressible signals — the quality of the guarantee adapts to the particular choice of x . These are also commonly referred to as *uniform guarantees* since they hold uniformly for all x .

Our choice of norms in (3.4) is somewhat arbitrary. We could easily measure the reconstruction error using other ℓ_p norms. The choice of p , however, will limit what kinds of guarantees are possible, and will also potentially lead to alternative formulations of the NSP. See, for instance, [36]. Moreover, the form of the right-hand-side of (3.4) might seem somewhat unusual in that we measure the approximation error as $\sigma_K(x)_1/\sqrt{K}$ rather than simply something like $\sigma_K(x)_2$. However, we will see later in this course⁷ that such a guarantee is actually not possible without taking a prohibitively large number of measurements, and that (3.4) represents the best possible guarantee we can hope to obtain (see "Instance-optimal guarantees revisited" (Section 4.4)).

Later in this course, we will show that the NSP of order $2K$ is sufficient to establish a guarantee of the form (3.4) for a practical recovery algorithm (see "Noise-free signal recovery" (Section 4.2)). Moreover, the following adaptation of a theorem in [36] demonstrates that if there exists *any* recovery algorithm satisfying (3.4), then Φ must necessarily satisfy the NSP of order $2K$.

Theorem 3.2: (Theorem 3.2 of [36])

Let $\Phi : \mathbb{R}^N \rightarrow \mathbb{R}^M$ denote a sensing matrix and $\Delta : \mathbb{R}^M \rightarrow \mathbb{R}^N$ denote an arbitrary recovery algorithm. If the pair (Φ, Δ) satisfies (3.4) then Φ satisfies the NSP of order $2K$.

⁶We note that this notation will occasionally be abused to refer to the length $|\Lambda|$ vector obtained by keeping only the entries corresponding to Λ or the $M \times |\Lambda|$ matrix obtained by only keeping the columns corresponding to Λ . The usage should be clear from the context, but typically there is no substantive difference between the two.

⁷An Introduction to Compressive Sensing <<http://cnx.org/content/col11133/latest/>>

Proof:

Suppose $h \in \mathcal{N}(\Phi)$ and let Λ be the indices corresponding to the $2K$ largest entries of h . We next split Λ into Λ_0 and Λ_1 , where $|\Lambda_0| = |\Lambda_1| = K$. Set $x = h_{\Lambda_1} + h_{\Lambda^c}$ and $x' = -h_{\Lambda_0}$, so that $h = x - x'$. Since by construction $x' \in \Sigma_K$, we can apply (3.4) to obtain $x' = \Delta(\Phi x')$. Moreover, since $h \in \mathcal{N}(\Phi)$, we have

$$\Phi h = \Phi(x - x') = 0 \quad (3.6)$$

so that $\Phi x' = \Phi x$. Thus, $x' = \Delta(\Phi x)$. Finally, we have that

$$\|h_{\Lambda}\|_2 \leq \|h\|_2 = \|x - x'\|_2 = \|x - \Delta(\Phi x)\|_2 \leq C \frac{\sigma_K(x)_1}{\sqrt{K}} = \sqrt{2}C \frac{\|h_{\Lambda^c}\|_1}{\sqrt{2K}}, \quad (3.7)$$

where the last inequality follows from (3.4).

3.3 The restricted isometry property⁸

The null space property (Section 3.2) (NSP) is both necessary and sufficient for establishing guarantees of the form

$$\|\Delta(\Phi x) - x\|_2 \leq C \frac{\sigma_K(x)_1}{\sqrt{K}}, \quad (3.8)$$

but these guarantees do not account for *noise*. When the measurements are contaminated with noise or have been corrupted by some error such as quantization, it will be useful to consider somewhat stronger conditions. In [29], Candès and Tao introduced the following isometry condition on matrices Φ and established its important role in compressive sensing⁹ (CS).

Definition 3.3:

A matrix Φ satisfies the *restricted isometry property* (RIP) of order K if there exists a $\delta_K \in (0, 1)$ such that

$$(1 - \delta_K) \|x\|_2^2 \leq \|\Phi x\|_2^2 \leq (1 + \delta_K) \|x\|_2^2, \quad (3.9)$$

holds for all $x \in \Sigma_K = \{x : \|x\|_0 \leq K\}$.

If a matrix Φ satisfies the RIP of order $2K$, then we can interpret (3.9) as saying that Φ approximately preserves the distance between any pair of K -sparse vectors. This will clearly have fundamental implications concerning robustness to noise.

It is important to note that in our definition of the RIP we assume bounds that are symmetric about 1, but this is merely for notational convenience. In practice, one could instead consider arbitrary bounds

$$\alpha \|x\|_2^2 \leq \|\Phi x\|_2^2 \leq \beta \|x\|_2^2 \quad (3.10)$$

where $0 < \alpha \leq \beta < \infty$. Given any such bounds, one can always scale Φ so that it satisfies the symmetric bound about 1 in (3.9). Specifically, multiplying Φ by $\sqrt{2/(\beta + \alpha)}$ will result in an $\tilde{\Phi}$ that satisfies (3.9) with constant $\delta_K = (\beta - \alpha)/(\beta + \alpha)$. We will not explicitly show this, but one can check that all of the theorems in this course based on the assumption that Φ satisfies the RIP actually hold as long as there exists some scaling of Φ that satisfies the RIP. Thus, since we can always scale Φ to satisfy (3.9), we lose nothing by restricting our attention to this simpler bound.

Note also that if Φ satisfies the RIP of order K with constant δ_K , then for any $K' < K$ we automatically have that Φ satisfies the RIP of order K' with constant $\delta_{K'} \leq \delta_K$. Moreover, in [136] it is shown that if Φ

⁸This content is available online at <http://cnx.org/content/m37171/1.6/>.

⁹"Introduction to compressive sensing" <http://cnx.org/content/m37172/latest/>

satisfies the RIP of order K with a sufficiently small constant, then it will also automatically satisfy the RIP of order γK for certain γ , albeit with a somewhat worse constant.

Lemma 3.1: (Corollary 3.4 of [136])

Suppose that Φ satisfies the RIP of order K with constant δ_K . Let γ be a positive integer. Then Φ satisfies the RIP of order $K' = \gamma \lfloor \frac{K}{2} \rfloor$ with constant $\delta_{K'} < \gamma \cdot \delta_K$, where $\lfloor \cdot \rfloor$ denotes the floor operator.

This lemma is trivial for $\gamma = 1, 2$, but for $\gamma \geq 3$ (and $K \geq 4$) this allows us to extend from RIP of order K to higher orders. Note however, that δ_K must be sufficiently small in order for the resulting bound to be useful.

3.3.1 The RIP and stability

We will see later in this course¹⁰ that if a matrix Φ satisfies the RIP, then this is sufficient for a variety of algorithms (Section 4.6) to be able to successfully recover a sparse signal from noisy measurements. First, however, we will take a closer look at whether the RIP is actually necessary. It should be clear that the lower bound in the RIP is a necessary condition if we wish to be able to recover all sparse signals x from the measurements Φx for the same reasons that the NSP is necessary. We can say even more about the necessity of the RIP by considering the following notion of stability.

Definition 3.4:

Let $\Phi : \mathbb{R}^N \rightarrow \mathbb{R}^M$ denote a sensing matrix and $\Delta : \mathbb{R}^M \rightarrow \mathbb{R}^N$ denote a recovery algorithm. We say that the pair (Φ, Δ) is C -stable if for any $x \in \Sigma_K$ and any $e \in \mathbb{R}^M$ we have that

$$\|\Delta(\Phi x + e) - x\|_2 \leq C\|e\|. \quad (3.11)$$

This definition simply says that if we add a small amount of noise to the measurements, then the impact of this on the recovered signal should not be arbitrarily large. Theorem 3.3, p. 25 below demonstrates that the existence of any decoding algorithm (potentially impractical) that can stably recover from noisy measurements requires that Φ satisfy the lower bound of (3.9) with a constant determined by C .

Theorem 3.3:

If the pair (Φ, Δ) is C -stable, then

$$\frac{1}{C}\|x\|_2 \leq \|\Phi x\|_2 \quad (3.12)$$

for all $x \in \Sigma_{2K}$.

Proof:

Pick any $x, z \in \Sigma_K$. Define

$$e_x = \frac{\Phi(z - x)}{2} \quad \text{and} \quad e_z = \frac{\Phi(x - z)}{2}, \quad (3.13)$$

and note that

$$\Phi x + e_x = \Phi z + e_z = \frac{\Phi(x + z)}{2}. \quad (3.14)$$

¹⁰ *An Introduction to Compressive Sensing* <<http://cnx.org/content/col11133/latest/>>

Let $\hat{x} = \Delta(\Phi x + e_x) = \Delta(\Phi z + e_z)$. From the triangle inequality and the definition of C -stability, we have that

$$\begin{aligned} \|x - z\|_2 &= \|\hat{x} - \hat{x} + \hat{x} - z\|_2 \\ &\leq \|\hat{x} - x\|_2 + \|\hat{x} - z\|_2 \\ &\leq C\|e_x\| + C\|e_z\|_2 \\ &= C\|\Phi x - \Phi z\|_2. \end{aligned} \tag{3.15}$$

Since this holds for any $x, z \in \Sigma_K$, the result follows.

Note that as $C \rightarrow 1$, we have that Φ must satisfy the lower bound of (3.9) with $\delta_K = 1 - 1/C^2 \rightarrow 0$. Thus, if we desire to reduce the impact of noise in our recovered signal then we must adjust Φ so that it satisfies the lower bound of (3.9) with a tighter constant.

One might respond to this result by arguing that since the upper bound is not necessary, we can avoid redesigning Φ simply by rescaling Φ so that as long as Φ satisfies the RIP with $\delta_{2K} < 1$, the rescaled version $\alpha\Phi$ will satisfy (3.12) for any constant C . In settings where the size of the noise is independent of our choice of Φ , this is a valid point — by scaling Φ we are simply adjusting the gain on the “signal” part of our measurements, and if increasing this gain does not impact the noise, then we can achieve arbitrarily high signal-to-noise ratios, so that eventually the noise is negligible compared to the signal.

However, in practice we will typically not be able to rescale Φ to be arbitrarily large. Moreover, in many practical settings the noise is not independent of Φ . For example, suppose that the noise vector e represents quantization noise produced by a finite dynamic range quantizer with B bits. Suppose the measurements lie in the interval $[-T, T]$, and we have adjusted the quantizer to capture this range. If we rescale Φ by α , then the measurements now lie between $[-\alpha T, \alpha T]$, and we must scale the dynamic range of our quantizer by α . In this case the resulting quantization error is simply αe , and we have achieved *no reduction* in the reconstruction error.

3.3.2 Measurement bounds

We can also consider how many measurements are necessary to achieve the RIP. If we ignore the impact of δ and focus only on the dimensions of the problem (N , M , and K) then we can establish a simple lower bound. We first provide a preliminary lemma that we will need in the proof of the main theorem.

Lemma 3.2:

Let K and N satisfying $K < N/2$ be given. There exists a set $X \subset \Sigma_K$ such that for any $x \in X$ we have $\|x\|_2 \leq \sqrt{K}$ and for any $x, z \in X$ with $x \neq z$,

$$\|x - z\|_2 \geq \sqrt{K/2}, \tag{3.16}$$

and

$$\log|X| \geq \frac{K}{2} \log\left(\frac{N}{K}\right). \tag{3.17}$$

Proof:

We will begin by considering the set

$$U = \{x \in \{0, +1, -1\}^N : \|x\|_0 = K\}. \tag{3.18}$$

By construction, $\|x\|_2^2 = K$ for all $x \in U$. Thus if we construct X by picking elements from U then we automatically have $\|x\|_2 \leq \sqrt{K}$.

Next, observe that $|U| = \binom{N}{K} 2^K$. Note also that $\|x - z\|_0 \leq \|x - z\|_2^2$, and thus if $\|x - z\|_2^2 \leq K/2$ then $\|x - z\|_0 \leq K/2$. From this we observe that for any fixed $x \in U$,

$$|\{z \in U : \|x - z\|_2^2 \leq K/2\}| \leq |\{z \in U : \|x - z\|_0 \leq K/2\}| \leq \binom{N}{K/2} 3^{K/2}. \quad (3.19)$$

Thus, suppose we construct the set X by iteratively choosing points that satisfy (3.16). After adding j points to the set, there are at least

$$\binom{N}{K} 2^K - j \binom{N}{K/2} 3^{K/2} \quad (3.20)$$

points left to pick from. Thus, we can construct a set of size $|X|$ provided that

$$|X| \binom{N}{K/2} 3^{K/2} \leq \binom{N}{K} 2^K \quad (3.21)$$

Next, observe that

$$\frac{\binom{N}{K}}{\binom{N}{K/2}} = \frac{(K/2)! (N - K/2)!}{K! (N - K)!} = \prod_{i=1}^{K/2} \frac{N - K + i}{K/2 + i} \geq \left(\frac{N}{K} - \frac{1}{2}\right)^{K/2}, \quad (3.22)$$

where the inequality follows from the fact that $(n - K + i) / (K/2 + i)$ is decreasing as a function of i . Thus, if we set $|X| = (N/K)^{K/2}$ then we have

$$|X| \left(\frac{3}{4}\right)^{K/2} = \left(\frac{3N}{4K}\right)^{K/2} = \left(\frac{N}{K} - \frac{N}{4K}\right)^{K/2} \leq \left(\frac{N}{K} - \frac{1}{2}\right)^{K/2} \leq \frac{\binom{N}{K}}{\binom{N}{K/2}}. \quad (3.23)$$

Hence, (3.21) holds for $|X| = (N/K)^{K/2}$, which establishes the lemma.

Using this lemma, we can establish the following bound on the required number of measurements to satisfy the RIP.

Theorem 3.4:

Let Φ be an $M \times N$ matrix that satisfies the RIP of order $2K$ with constant $\delta \in (0, \frac{1}{2}]$. Then

$$M \geq CK \log \left(\frac{N}{K}\right) \quad (3.24)$$

where $C = 1/2 \log(\sqrt{24} + 1) \approx 0.28$.

Proof:

We first note that since Φ satisfies the RIP, then for the set of points X in Lemma 3.2, p. 26 we have,

$$\|\Phi x - \Phi z\|_2 \geq \sqrt{1 - \delta} \|x - z\|_2 \geq \sqrt{K/4} \quad (3.25)$$

for all $x, z \in X$, since $x - z \in \Sigma_{2K}$ and $\delta \leq \frac{1}{2}$. Similarly, we also have

$$\|\Phi x\|_2 \leq \sqrt{1 + \delta} \|x\|_2 \leq \sqrt{3K/2} \quad (3.26)$$

for all $x \in X$.

From the lower bound we can say that for any pair of points $x, z \in X$, if we center balls of radius $\sqrt{K/4}/2 = \sqrt{K/16}$ at Φx and Φz , then these balls will be disjoint. In turn, the upper bound tells

us that the entire set of balls is itself contained within a larger ball of radius $\sqrt{3K/2} + \sqrt{K/16}$. If we let $B^M(r) = \{x \in \mathbb{R}^M : \|x\|_2 \leq r\}$, this implies that

$$\begin{aligned} \text{Vol}\left(B^M\left(\sqrt{3K/2} + \sqrt{K/16}\right)\right) &\geq |X| \cdot \text{Vol}\left(B^M\left(\sqrt{K/16}\right)\right), \\ \left(\sqrt{3K/2} + \sqrt{K/16}\right)^M &\geq |X| \cdot \left(\sqrt{K/16}\right)^M, \\ (\sqrt{24} + 1)^M &\geq |X|, \\ M &\geq \frac{\log|X|}{\log(\sqrt{24}+1)}. \end{aligned} \quad (3.27)$$

The theorem follows by applying the bound for $|X|$ from Lemma 3.2, p. 26.

Note that the restriction to $\delta \leq \frac{1}{2}$ is arbitrary and is made merely for convenience — minor modifications to the argument establish bounds for $\delta \leq \delta_{\max}$ for any $\delta_{\max} < 1$. Moreover, although we have made no effort to optimize the constants, it is worth noting that they are already quite reasonable.

Although the proof is somewhat less direct, one can establish a similar result (in the dependence on N and K) by examining the *Gelfand width* of the ℓ_1 ball [85]. However, both this result and Theorem 3.4, p. 27 fail to capture the precise dependence of M on the desired RIP constant δ . In order to quantify this dependence, we can exploit recent results concerning the *Johnson-Lindenstrauss lemma*, which concerns embeddings of finite sets of points in low-dimensional spaces [106]. Specifically, it is shown in [104] that if we are given a point cloud with p points and wish to embed these points in $\mathbb{R}^M$ such that the squared ℓ_2 distance between any pair of points is preserved up to a factor of $1 \pm \varepsilon$, then we must have that

$$M \geq \frac{c_0 \log(p)}{\varepsilon^2}, \quad (3.28)$$

where $c_0 > 0$ is a constant.

The Johnson-Lindenstrauss lemma is closely related to the RIP. We will see in "Matrices that satisfy the RIP" (Section 3.5) that any procedure that can be used for generating a linear, distance-preserving embedding for a point cloud can also be used to construct a matrix that satisfies the RIP. Moreover, in [112] it is shown that if a matrix Φ satisfies the RIP of order $K = c_1 \log(p)$ with constant δ , then Φ can be used to construct a distance-preserving embedding for p points with $\varepsilon = \delta/4$. Combining these we obtain

$$M \geq \frac{c_0 \log(p)}{\varepsilon^2} = \frac{16c_0 K}{c_1 \delta^2}. \quad (3.29)$$

Thus, for small δ the number of measurements required to ensure that Φ satisfies the RIP of order K will be proportional to K/δ^2 , which may be significantly higher than $K \log(N/K)$. See [112] for further details.

3.4 The RIP and the NSP¹¹

Next we will show that if a matrix satisfies the restricted isometry property (Section 3.3) (RIP), then it also satisfies the null space property (Section 3.2) (NSP). Thus, the RIP is strictly stronger than the NSP.

Theorem 3.5:

Suppose that Φ satisfies the RIP of order $2K$ with $\delta_{2K} < \sqrt{2} - 1$. Then Φ satisfies the NSP of order $2K$ with constant

$$C = \frac{\sqrt{2}\delta_{2K}}{1 - (1 + \sqrt{2})\delta_{2K}}. \quad (3.30)$$

¹¹This content is available online at <http://cnx.org/content/m37176/1.5/>.

Proof:

The proof of this theorem involves two useful lemmas. The first of these follows directly from standard norm inequality by relating a K -sparse vector to a vector in $\mathbb{R}^K$. We include a simple proof for the sake of completeness.

Lemma 3.3:

Suppose $u \in \Sigma_K$. Then

$$\frac{\|u\|_1}{\sqrt{K}} \leq \|u\|_2 \leq \sqrt{K}\|u\|_\infty. \quad (3.31)$$

Proof:

For any u , $\|u\|_1 = |\langle u, \text{sgn}(u) \rangle|$. By applying the Cauchy-Schwarz inequality we obtain $\|u\|_1 \leq \|u\|_2 \|\text{sgn}(u)\|_2$. The lower bound follows since $\text{sgn}(u)$ has exactly K nonzero entries all equal to ± 1 (since $u \in \Sigma_K$) and thus $\|\text{sgn}(u)\|_2 = \sqrt{K}$. The upper bound is obtained by observing that each of the K nonzero entries of u can be upper bounded by $\|u\|_\infty$.

Below we state the second key lemma that we will need in order to prove Theorem 3.5, p. 28. This result is a general result which holds for arbitrary h , not just vectors $h \in \mathcal{N}(\Phi)$. It should be clear that when we do have $h \in \mathcal{N}(\Phi)$, the argument could be simplified considerably. However, this lemma will prove immensely useful when we turn to the problem of sparse recovery from noisy measurements (Section 4.3) later in this course¹², and thus we establish it now in its full generality. We state the lemma here, which is proven in " ℓ_1 minimization proof"¹³.

Lemma 3.4:

Suppose that Φ satisfies the RIP of order $2K$, and let $h \in \mathbb{R}^N$, $h \neq 0$ be arbitrary. Let Λ_0 be any subset of $\{1, 2, \dots, N\}$ such that $|\Lambda_0| \leq K$. Define Λ_1 as the index set corresponding to the K entries of $h_{\Lambda_0^c}$ with largest magnitude, and set $\Lambda = \Lambda_0 \cup \Lambda_1$. Then

$$\|h_\Lambda\|_2 \leq \alpha \frac{\|h_{\Lambda_0^c}\|_1}{\sqrt{K}} + \beta \frac{|\langle \Phi h_\Lambda, \Phi h \rangle|}{\|h_\Lambda\|_2}, \quad (3.32)$$

where

$$\alpha = \frac{\sqrt{2}\delta_{2K}}{1 - \delta_{2K}}, \quad \beta = \frac{1}{1 - \delta_{2K}}. \quad (3.33)$$

Again, note that Lemma 3.4, p. 29 holds for arbitrary h . In order to prove Theorem 3.5, p. 28, we merely need to apply Lemma 3.4, p. 29 to the case where $h \in \mathcal{N}(\Phi)$.

Towards this end, suppose that $h \in \mathcal{N}(\Phi)$. It is sufficient to show that

$$\|h_\Lambda\|_2 \leq C \frac{\|h_{\Lambda^c}\|_1}{\sqrt{K}} \quad (3.34)$$

holds for the case where Λ is the index set corresponding to the $2K$ largest entries of h . Thus, we can take Λ_0 to be the index set corresponding to the K largest entries of h and apply Lemma 3.4, p. 29.

The second term in Lemma 3.4, p. 29 vanishes since $\Phi h = 0$, and thus we have

$$\|h_\Lambda\|_2 \leq \alpha \frac{\|h_{\Lambda_0^c}\|_1}{\sqrt{K}}. \quad (3.35)$$

¹² An Introduction to Compressive Sensing <<http://cnx.org/content/col11133/latest/>>

¹³ " ℓ_1 minimization proof" <<http://cnx.org/content/m37187/latest/>>

Using Lemma 3.3, p. 29,

$$\|h_{\Lambda_0^c}\|_1 = \|h_{\Lambda_1}\|_1 + \|h_{\Lambda^c}\|_1 \leq \sqrt{K}\|h_{\Lambda_1}\|_2 + \|h_{\Lambda^c}\|_1 \quad (3.36)$$

resulting in

$$\|h_{\Lambda}\|_2 \leq \alpha \left(\|h_{\Lambda_1}\|_2 + \frac{\|h_{\Lambda^c}\|_1}{\sqrt{K}} \right). \quad (3.37)$$

Since $\|h_{\Lambda_1}\|_2 \leq \|h_{\Lambda}\|_2$, we have that

$$(1 - \alpha) \|h_{\Lambda}\|_2 \leq \alpha \frac{\|h_{\Lambda^c}\|_1}{\sqrt{K}}. \quad (3.38)$$

The assumption $\delta_{2K} < \sqrt{2} - 1$ ensures that $\alpha < 1$, and thus we may divide by $1 - \alpha$ without changing the direction of the inequality to establish (3.34) with constant

$$C = \frac{\alpha}{1 - \alpha} = \frac{\sqrt{2}\delta_{2K}}{1 - (1 + \sqrt{2})\delta_{2K}}, \quad (3.39)$$

as desired.

3.5 Matrices that satisfy the RIP¹⁴

We now turn to the question of how to construct matrices that satisfy the restricted isometry property (Section 3.3) (RIP). It is possible to deterministically construct matrices of size $M \times N$ that satisfy the RIP of order K , but such constructions also require M to be relatively large [52], [101]. For example, the construction in [52] requires $M = O(K^2 \log N)$ while the construction in [101] requires $M = O(KN^\alpha)$ for some constant α . In many real-world settings, these results would lead to an unacceptably large requirement on M .

Fortunately, these limitations can be overcome by randomizing the matrix construction. We will construct our random matrices as follows: given M and N , generate random matrices Φ by choosing the entries ϕ_{ij} as independent realizations from some probability distribution. We begin by observing that if all we require is that $\delta_{2K} > 0$ then we may set $M = 2K$ and draw a Φ according to a Gaussian distribution. With probability 1, any subset of $2K$ columns will be linearly independent, and hence all subsets of $2K$ columns will be bounded below by $1 - \delta_{2K}$ where $\delta_{2K} > 0$. However, suppose we wish to know the constant δ_{2K} . In order to find the value of the constant we must consider all possible $\binom{N}{K} K$ -dimensional subspaces of $\mathbb{R}^N$. From a computational perspective, this is impossible for any realistic values of N and K . Thus, we focus on the problem of achieving the RIP of order $2K$ for a specified constant δ_{2K} . Our treatment is based on the simple approach first described in [6] and later generalized to a larger class of random matrices in [129].

To ensure that the matrix will satisfy the RIP, we will impose two conditions on the random distribution. First, we require that the distribution will yield a matrix that is norm-preserving, which will require that

$$\mathbb{E}(\phi_{ij}^2) = \frac{1}{M}, \quad (3.40)$$

and hence the variance of the distribution is $1/M$. Second, we require that the distribution is a *sub-Gaussian* distribution (Section 3.7), meaning that there exists a constant $c > 0$ such that

$$\mathbb{E}(e^{\phi_{ij}t}) \leq e^{c^2 t^2 / 2} \quad (3.41)$$

¹⁴This content is available online at <<http://cnx.org/content/m37177/1.5/>>.

for all $t \in \mathbb{R}$. This says that the moment-generating function of our distribution is dominated by that of a Gaussian distribution, which is also equivalent to requiring that tails of our distribution decay *at least as fast* as the tails of a Gaussian distribution. Examples of sub-Gaussian distributions include the Gaussian distribution, the Bernoulli distribution taking values $\pm 1/\sqrt{M}$, and more generally any distribution with bounded support. See "Sub-Gaussian random variables" (Section 3.7) for more details.

For the moment, we will actually assume a bit more than sub-Gaussianity. Specifically, we will assume that the entries of Φ are *strictly* sub-Gaussian, which means that they satisfy (3.41) with $c^2 = \mathbb{E}(\phi_{ij}^2) = \frac{1}{M}$. Similar results to the following would hold for general sub-Gaussian distributions, but to simplify the constants we restrict our present attention to the strictly sub-Gaussian Φ . In this case we have the following useful result, which is proven in "Concentration of measure for sub-Gaussian random variables" (Section 3.8).

Corollary 3.1:

Suppose that Φ is an $M \times N$ matrix whose entries ϕ_{ij} are i.i.d. with ϕ_{ij} drawn according to a strictly sub-Gaussian distribution with $c^2 = 1/M$. Let $Y = \Phi x$ for $x \in \mathbb{R}^N$. Then for any $\varepsilon > 0$, and any $x \in \mathbb{R}^N$,

$$\mathbb{E}(\|Y\|_2^2) = \|x\|_2^2 \quad (3.42)$$

and

$$\mathbb{P}(|\|Y\|_2^2 - \|x\|_2^2| \geq \varepsilon \|x\|_2^2) \leq 2\exp\left(-\frac{M\varepsilon^2}{\kappa^*}\right) \quad (3.43)$$

with $\kappa^* = 2/(1 - \log(2)) \approx 6.52$.

This tells us that the norm of a sub-Gaussian random vector strongly concentrates about its mean. Using this result, in "Proof of the RIP for sub-Gaussian matrices" (Section 3.9) we provide a simple proof based on that in [6] that sub-Gaussian matrices satisfy the RIP.

Theorem 3.6:

Fix $\delta \in (0, 1)$. Let Φ be an $M \times N$ random matrix whose entries ϕ_{ij} are i.i.d. with ϕ_{ij} drawn according to a strictly sub-Gaussian distribution with $c^2 = 1/M$. If

$$M \geq \kappa_1 K \log\left(\frac{N}{K}\right), \quad (3.44)$$

then Φ satisfies the RIP of order K with the prescribed δ with probability exceeding $1 - 2e^{-\kappa_2 M}$, where κ_1 is arbitrary and $\kappa_2 = \delta^2/2\kappa^* - \log(42e/\delta)/\kappa_1$.

Note that in light of the measurement bounds in "The restricted isometry property" (Section 3.3) we see that (3.44) achieves the optimal number of measurements (up to a constant).

Using random matrices to construct Φ has a number of additional benefits. To illustrate these, we will focus on the RIP. First, one can show that for random constructions the measurements are *democratic*, meaning that it is possible to recover a signal using any sufficiently large subset of the measurements [48], [114]. Thus, by using random Φ one can be robust to the loss or corruption of a small fraction of the measurements. Second, and perhaps more significantly, in practice we are often more interested in the setting where x is sparse with respect to some basis Ψ . In this case what we actually require is that the product $\Phi\Psi$ satisfies the RIP. If we were to use a deterministic construction then we would need to explicitly take Ψ into account in our construction of Φ , but when Φ is chosen randomly we can avoid this consideration. For example, if Φ is chosen according to a Gaussian distribution and Ψ is an orthonormal basis then one can easily show that $\Phi\Psi$ will also have a Gaussian distribution, and so provided that M is sufficiently high $\Phi\Psi$ will satisfy the RIP with high probability, just as before. Although less obvious, similar results hold for sub-Gaussian distributions as well [6]. This property, sometimes referred to as *universality*, constitutes a significant advantage of using random matrices to construct Φ .

Finally, we note that since the fully random matrix approach is sometimes impractical to build in hardware, several hardware architectures have been implemented and/or proposed that enable random measurements to be acquired in practical settings. Examples include the random demodulator (Section 5.5) [173],

random filtering [175], the modulated wideband converter [132], random convolution [2], [148], and the compressive multiplexer [160]. These architectures typically use a reduced amount of randomness and are modeled via matrices Φ that have significantly more structure than a fully random matrix. Perhaps somewhat surprisingly, while it is typically not quite as easy as in the fully random case, one can prove that many of these constructions also satisfy the RIP.

3.6 Coherence¹⁵

While the spark (Section 3.2), null space property (Section 3.2) (NSP), and restricted isometry property (Section 3.3) (RIP) all provide guarantees for the recovery of sparse (Section 2.3) signals, verifying that a general matrix Φ satisfies any of these properties has a combinatorial computational complexity, since in each case one must essentially consider $\binom{N}{K}$ submatrices. In many settings it is preferable to use properties of Φ that are easily computable to provide more concrete recovery guarantees. The *coherence* of a matrix is one such property [60], [171].

Definition 3.5:

The coherence of a matrix Φ , $\mu(\Phi)$, is the largest absolute inner product between any two columns ϕ_i, ϕ_j of Φ :

$$\mu(\Phi) = \max_{1 \leq i < j \leq N} \frac{|\langle \phi_i, \phi_j \rangle|}{\|\phi_i\|_2 \|\phi_j\|_2}. \quad (3.45)$$

It is possible to show that the coherence of a matrix is always in the range $\mu(\Phi) \in \left[\sqrt{\frac{N-M}{M(N-1)}}, 1 \right]$; the lower bound is known as the Welch bound [150], [161], [189]. Note that when $N \gg M$, the lower bound is approximately $\mu(\Phi) \geq 1/\sqrt{M}$.

One can sometimes relate coherence to the spark, NSP, and RIP. For example, the coherence and spark properties of a matrix can be related by employing the Gershgorin circle theorem [87], [181].

Theorem 3.7: (Theorem 2 of [87])

The eigenvalues of an $N \times N$ matrix M with entries m_{ij} , $1 \leq i, j \leq N$, lie in the union of N discs $d_i = d_i(c_i, r_i)$, $1 \leq i \leq N$, centered at $c_i = m_{ii}$ and with radius $r_i = \sum_{j \neq i} |m_{ij}|$.

Applying this theorem on the Gram matrix $G = \Phi_\Lambda^T \Phi_\Lambda$ leads to the following straightforward result.

Lemma 3.5:

For any matrix Φ ,

$$\text{spark}(\Phi) \geq 1 + \frac{1}{\mu(\Phi)}. \quad (3.46)$$

Proof:

Since $\text{spark}(\Phi)$ does not depend on the scaling of the columns, we can assume without loss of generality that Φ has unit-norm columns. Let $\Lambda \subseteq \{1, \dots, N\}$ with $|\Lambda| = p$ determine a set of indices. We consider the restricted Gram matrix $G = \Phi_\Lambda^T \Phi_\Lambda$, which satisfies the following properties:

- $g_{ii} = 1$, $1 \leq i \leq p$;
- $|g_{ij}| \leq \mu(\Phi)$, $1 \leq i, j \leq p$, $i \neq j$.

From Theorem 3.7, (Theorem 2 of [87]), p. 32, if $\sum_{j \neq i} |g_{ij}| < |g_{ii}|$ then the matrix G is positive definite, so that the columns of Φ_Λ are linearly independent. Thus, the spark condition implies $(p-1)\mu(\Phi) < 1$ or, equivalently, $p < 1 + 1/\mu(\Phi)$ for all $p < \text{spark}(\Phi)$, yielding $\text{spark}(\Phi) \geq 1 + 1/\mu(\Phi)$.

¹⁵This content is available online at <http://cnx.org/content/m37178/1.5/>.

By merging Theorem 1 from "Null space conditions" (Section 3.2) with Lemma 3.5, p. 32, we can pose the following condition on Φ that guarantees uniqueness.

Theorem 3.8: (Theorem 12 of [60])

If

$$K < \frac{1}{2} \left(1 + \frac{1}{\mu(\Phi)} \right), \quad (3.47)$$

then for each measurement vector $y \in \mathbb{R}^M$ there exists at most one signal $x \in \Sigma_K$ such that $y = \Phi x$.

Theorem 3.8, (Theorem 12 of [60]), p. 33, together with the Welch bound, provides an upper bound on the level of sparsity K that guarantees uniqueness using coherence: $K = O(\sqrt{M})$. Another straightforward application of the Gershgorin circle theorem (Theorem 3.7, (Theorem 2 of [87]), p. 32) connects the RIP to the coherence property.

Lemma 3.6:

If Φ has unit-norm columns and coherence $\mu = \mu(\Phi)$, then Φ satisfies the RIP of order K with $\delta = (K - 1)\mu$ for all $K < 1/\mu$.

The proof of this lemma is similar to that of Lemma 3.5, p. 32.

The results given here emphasize the need for small coherence $\mu(\Phi)$ for the matrices used in CS. Coherence bounds have been studied both for deterministic and randomized matrices. For example, there are known matrices Φ of size $M \times M^2$ that achieve the coherence lower bound $\mu(\Phi) = 1/\sqrt{M}$, such as the Gabor frame generated from the Alltop sequence [100] and more general equiangular tight frames [161]. These constructions restrict the number of measurements needed to recover a K -sparse signal to be $M = O(K^2 \log N)$. Furthermore, it can be shown that when the distribution used has zero mean and finite variance, then in the asymptotic regime (as M and N grow) the coherence converges to $\mu(\Phi) = \sqrt{(2 \log N)/M}$ [21], [25], [56]. Such constructions would allow K to grow asymptotically as $M = O(K^2 \log N)$, matching the known finite-dimensional bounds.

The measurement bounds dependent on coherence are handicapped by the squared dependence on the sparsity K , but it is possible to overcome this bottleneck by shifting the types of guarantees from worst-case/deterministic behavior, to average-case/probabilistic behavior [177], [178]. In this setting, we pose a probabilistic prior on the set of K -sparse signals $x \in \Sigma_K$. It is then possible to show that if Φ has low coherence $\mu(\Phi)$ and spectral norm $\|\Phi\|_2$, and if $K = O(\mu^{-2}(\Phi) \log N)$, then the signal x can be recovered from its CS measurements $y = \Phi x$ with high probability. Note that if we replace the Welch bound, then we obtain $K = O(M \log N)$, which returns to the linear dependence of the measurement bound on the signal sparsity that appears in RIP-based results.

3.7 Sub-Gaussian random variables¹⁶

A number of distributions, notably Gaussian and Bernoulli, are known to satisfy certain concentration of measure (Section 3.8) inequalities. We will analyze this phenomenon from a more general perspective by considering the class of sub-Gaussian distributions [20].

Definition 3.6:

A random variable X is called *sub-Gaussian* if there exists a constant $c > 0$ such that

$$\mathbb{E}(\exp(Xt)) \leq \exp(c^2 t^2 / 2) \quad (3.48)$$

holds for all $t \in \mathbb{R}$. We use the notation $X \sim \text{Sub}(c^2)$ to denote that X satisfies (3.48).

The function $\mathbb{E}(\exp(Xt))$ is the *moment-generating function* of X , while the upper bound in (3.48) is the moment-generating function of a Gaussian random variable. Thus, a sub-Gaussian distribution is one

¹⁶This content is available online at <<http://cnx.org/content/m37185/1.6/>>.

whose moment-generating function is bounded by that of a Gaussian. There are a tremendous number of sub-Gaussian distributions, but there are two particularly important examples:

Example 3.1

If $X \sim \mathcal{N}(0, \sigma^2)$, i.e., X is a zero-mean Gaussian random variable with variance σ^2 , then $X \sim \text{Sub}(\sigma^2)$. Indeed, as mentioned above, the moment-generating function of a Gaussian is given by $\mathbb{E}(\exp(Xt)) = \exp(\sigma^2 t^2/2)$, and thus (3.48) is trivially satisfied.

Example 3.2

If X is a zero-mean, bounded random variable, i.e., one for which there exists a constant B such that $|X| \leq B$ with probability 1, then $X \sim \text{Sub}(B^2)$.

A common way to characterize sub-Gaussian random variables is through analyzing their moments. We consider only the mean and variance in the following elementary lemma, proven in [20].

Lemma 3.7: (Buldygin-Kozachenko [20])

If $X \sim \text{Sub}(c^2)$ then,

$$\mathbb{E}(X) = 0 \quad (3.49)$$

and

$$\mathbb{E}(X^2) \leq c^2. \quad (3.50)$$

Lemma 3.7, (Buldygin-Kozachenko [20]), p. 34 shows that if $X \sim \text{Sub}(c^2)$ then $\mathbb{E}(X^2) \leq c^2$. In some settings it will be useful to consider a more restrictive class of random variables for which this inequality becomes an equality.

Definition 3.7:

A random variable X is called *strictly sub-Gaussian* if $X \sim \text{Sub}(\sigma^2)$ where $\sigma^2 = \mathbb{E}(X^2)$, i.e., the inequality

$$\mathbb{E}(\exp(Xt)) \leq \exp(\sigma^2 t^2/2) \quad (3.51)$$

holds for all $t \in \mathbb{R}$. To denote that X is strictly sub-Gaussian with variance σ^2 , we will use the notation $X \sim \text{SSub}(\sigma^2)$.

Example 3.3

If $X \sim \mathcal{N}(0, \sigma^2)$, then $X \sim \text{SSub}(\sigma^2)$.

Example 3.4

If $X \sim U(-1, 1)$, i.e., X is uniformly distributed on the interval $[-1, 1]$, then $X \sim \text{SSub}(1/3)$.

Example 3.5

Now consider the random variable with distribution such that

$$\mathbb{P}(X = 1) = \mathbb{P}(X = -1) = \frac{1-s}{2}, \quad \mathbb{P}(X = 0) = s, \quad s \in [0, 1]. \quad (3.52)$$

For any $s \in [0, 2/3]$, $X \sim \text{SSub}(1-s)$. For $s \in (2/3, 1)$, X is not strictly sub-Gaussian.

We now provide an equivalent characterization for sub-Gaussian and strictly sub-Gaussian random variables, proven in [20], that illustrates their concentration of measure behavior.

Theorem 3.9: (Buldygin-Kozachenko [20])

A random variable $X \sim \text{Sub}(c^2)$ if and only if there exists a $t_0 \geq 0$ and a constant $a \geq 0$ such that

$$\mathbb{P}(|X| \geq t) \leq 2\exp\left(-\frac{t^2}{2a^2}\right) \quad (3.53)$$

for all $t \geq t_0$. Moreover, if $X \sim \text{SSub}(\sigma^2)$, then (3.53) holds for all $t > 0$ with $a = \sigma$.

Finally, sub-Gaussian distributions also satisfy one of the fundamental properties of a Gaussian distribution: the sum of two sub-Gaussian random variables is itself a sub-Gaussian random variable. This result is established in more generality in the following lemma.

Lemma 3.8:

Suppose that $X = [X_1, X_2, \dots, X_N]$, where each X_i is independent and identically distributed (i.i.d.) with $X_i \sim \text{Sub}(c^2)$. Then for any $\alpha \in \mathbb{R}^N$, $\langle X, \alpha \rangle \sim \text{Sub}(c^2 \|\alpha\|_2^2)$. Similarly, if each $X_i \sim \text{SSub}(\sigma^2)$, then for any $\alpha \in \mathbb{R}^N$, $\langle X, \alpha \rangle \sim \text{SSub}(\sigma^2 \|\alpha\|_2^2)$.

Proof:

Since the X_i are i.i.d., the joint distribution factors and simplifies as:

$$\begin{aligned} \mathbb{E} \left(\exp \left(t \sum_{i=1}^N \alpha_i X_i \right) \right) &= \mathbb{E} \left(\prod_{i=1}^N \exp(t \alpha_i X_i) \right) \\ &= \prod_{i=1}^N \mathbb{E} \left(\exp(t \alpha_i X_i) \right) \\ &\leq \prod_{i=1}^N \exp \left(c^2 (\alpha_i t)^2 / 2 \right) \\ &= \exp \left(\left(\sum_{i=1}^N \alpha_i^2 \right) c^2 t^2 / 2 \right). \end{aligned} \quad (3.54)$$

If the X_i are strictly sub-Gaussian, then the result follows by setting $c^2 = \sigma^2$ and observing that $\mathbb{E}(\langle X, \alpha \rangle^2) = \sigma^2 \|\alpha\|_2^2$.

3.8 Concentration of measure for sub-Gaussian random variables¹⁷

Sub-Gaussian distributions (Section 3.7) have a close relationship to the concentration of measure phenomenon [116]. To illustrate this, we note that we can combine Lemma 2 and Theorem 1 from "Sub-Gaussian random variables" (Section 3.7) to obtain deviation bounds for weighted sums of sub-Gaussian random variables. For our purposes, however, it will be more enlightening to study the *norm* of a vector of sub-Gaussian random variables. In particular, if X is a vector where each X_i is i.i.d. with $X_i \sim \text{Sub}(c)$, then we would like to know how $\|X\|_2$ deviates from its mean.

In order to establish the result, we will make use of Markov's inequality for nonnegative random variables.

Lemma 3.9: (Markov's Inequality)

For any nonnegative random variable X and $t > 0$,

$$\mathbb{P}(X \geq t) \leq \frac{\mathbb{E}(X)}{t}. \quad (3.55)$$

Proof:

Let $f(x)$ denote the probability density function for X .

$$\mathbb{E}(X) = \int_0^\infty x f(x) dx \geq \int_t^\infty x f(x) dx \geq \int_t^\infty t f(x) dx = t \mathbb{P}(X \geq t). \quad (3.56)$$

In addition, we will require the following bound on the exponential moment of a sub-Gaussian random variable.

Lemma 3.10:

Suppose $X \sim \text{Sub}(c^2)$. Then

$$\mathbb{E} \left(\exp(\lambda X^2 / 2c^2) \right) \leq \frac{1}{\sqrt{1 - \lambda}}, \quad (3.57)$$

¹⁷This content is available online at <http://cnx.org/content/m32583/1.7/>.

for any $\lambda \in [0, 1)$.

Proof:

First, observe that if $\lambda = 0$, then the lemma holds trivially. Thus, suppose that $\lambda \in (0, 1)$. Let $f(x)$ denote the probability density function for X . Since X is sub-Gaussian, we have by definition that

$$\int_{-\infty}^{\infty} \exp(tx) f(x) dx \leq \exp(c^2 t^2 / 2) \quad (3.58)$$

for any $t \in \mathbb{R}$. If we multiply by $\exp(-c^2 t^2 / 2\lambda)$, then we obtain

$$\int_{-\infty}^{\infty} \exp(tx - c^2 t^2 / 2\lambda) f(x) dx \leq \exp(c^2 t^2 (\lambda - 1) / 2\lambda). \quad (3.59)$$

Now, integrating both sides with respect to t , we obtain

$$\int_{-\infty}^{\infty} \left(\int_{-\infty}^{\infty} \exp(tx - c^2 t^2 / 2\lambda) dt \right) f(x) dx \leq \int_{-\infty}^{\infty} \exp(c^2 t^2 (\lambda - 1) / 2\lambda) dt, \quad (3.60)$$

which reduces to

$$\frac{1}{c} \sqrt{2\pi\lambda} \int_{-\infty}^{\infty} \exp(\lambda x^2 / 2c^2) f(x) dx \leq \frac{1}{c} \sqrt{\frac{2\pi\lambda}{1-\lambda}}. \quad (3.61)$$

This simplifies to prove the lemma.

We now state our main theorem, which generalizes the results of [43] and uses substantially the same proof technique.

Theorem 3.10:

Suppose that $X = [X_1, X_2, \dots, X_M]$, where each X_i is i.i.d. with $X_i \sim \text{Sub}(c^2)$ and $\mathbb{E}(X_i^2) = \sigma^2$. Then

$$\mathbb{E}(\|X\|_2^2) = M\sigma^2. \quad (3.62)$$

Moreover, for any $\alpha \in (0, 1)$ and for any $\beta \in [c^2/\sigma^2, \beta_{\max}]$, there exists a constant $\kappa^* \geq 4$ depending only on $\beta_{\max}$ and the ratio σ^2/c^2 such that

$$\mathbb{P}(\|X\|_2^2 \leq \alpha M\sigma^2) \leq \exp(-M(1-\alpha)^2/\kappa^*) \quad (3.63)$$

and

$$\mathbb{P}(\|X\|_2^2 \geq \beta M\sigma^2) \leq \exp(-M(\beta-1)^2/\kappa^*). \quad (3.64)$$

Proof:

Since the X_i are independent, we obtain

$$\mathbb{E}(\|X\|_2^2) = \sum_{i=1}^M \mathbb{E}(X_i^2) = \sum_{i=1}^M \sigma^2 = M\sigma^2 \quad (3.65)$$

and hence (3.62) holds. We now turn to (3.63) and (3.64). Let us first consider (3.64). We begin by applying Markov's inequality:

$$\begin{aligned} \mathbb{P}(\|X\|_2^2 \geq \beta M\sigma^2) &= \mathbb{P}(\exp(\lambda \|X\|_2^2) \geq \exp(\lambda \beta M\sigma^2)) \\ &\leq \frac{\mathbb{E}(\exp(\lambda \|X\|_2^2))}{\exp(\lambda \beta M\sigma^2)} \\ &= \frac{\prod_{i=1}^M \mathbb{E}(\exp(\lambda X_i^2))}{\exp(\lambda \beta M\sigma^2)}. \end{aligned} \quad (3.66)$$

Since $X_i \sim \text{Sub}(c^2)$, we have from Lemma 3.10, p. 35 that

$$\mathbb{E}(\exp(\lambda X_i^2)) = \mathbb{E}(\exp(2c^2\lambda X_i^2/2c^2)) \leq \frac{1}{\sqrt{1-2c^2\lambda}}. \quad (3.67)$$

Thus,

$$\prod_{i=1}^M \mathbb{E}(\exp(\lambda X_i^2)) \leq \left(\frac{1}{1-2c^2\lambda} \right)^{M/2} \quad (3.68)$$

and hence

$$\mathbb{P}(\|X\|_2^2 \geq \beta M \sigma^2) \leq \left(\frac{\exp(-2\lambda\beta\sigma^2)}{1-2c^2\lambda} \right)^{M/2}. \quad (3.69)$$

By setting the derivative to zero and solving for λ , one can show that the optimal λ is

$$\lambda = \frac{\beta\sigma^2 - c^2}{2c^2\sigma^2(1+\beta)}. \quad (3.70)$$

Plugging this in we obtain

$$\mathbb{P}(\|X\|_2^2 \geq \beta M \sigma^2) \leq \left(\beta \frac{\sigma^2}{c^2} \exp\left(1 - \beta \frac{\sigma^2}{c^2}\right) \right)^{M/2}. \quad (3.71)$$

Similarly,

$$\mathbb{P}(\|X\|_2^2 \leq \alpha M \sigma^2) \leq \left(\alpha \frac{\sigma^2}{c^2} \exp\left(1 - \alpha \frac{\sigma^2}{c^2}\right) \right)^{M/2}. \quad (3.72)$$

In order to combine and simplify these inequalities, note that if we define

$$\kappa^* = \max\left(4, 2 \frac{(\beta_{\max}\sigma^2/c - 1)^2}{(\beta_{\max}\sigma^2/c - 1) - \log(\beta_{\max}\sigma^2/c)}\right) \quad (3.73)$$

then we have that for any $\gamma \in [0, \beta_{\max}\sigma^2/c]$ we have the bound

$$\log(\gamma) \leq (\gamma - 1) - \frac{2(\gamma - 1)^2}{\kappa^*}, \quad (3.74)$$

and hence

$$\gamma \leq \exp\left((\gamma - 1) - \frac{2(\gamma - 1)^2}{\kappa^*}\right). \quad (3.75)$$

By setting $\gamma = \alpha\sigma^2/c^2$, (3.72) reduces to yield (3.63). Similarly, setting $\gamma = \beta\sigma^2/c^2$ establishes (3.64).

This result tells us that given a vector with entries drawn according to a sub-Gaussian distribution, we can expect the norm of the vector to concentrate around its expected value of $M\sigma^2$ with exponentially high probability as M grows. Note, however, that the range of allowable choices for β in (3.64) is limited to $\beta \geq c^2/\sigma^2 \geq 1$. Thus, for a general sub-Gaussian distribution, we may be unable to achieve an arbitrarily tight concentration. However, recall that for strictly sub-Gaussian distributions we have that $c^2 = \sigma^2$, in

which there is no such restriction. Moreover, for strictly sub-Gaussian distributions we also have the following useful corollary.¹⁸

Corollary 3.2:

Suppose that $X = [X_1, X_2, \dots, X_M]$, where each X_i is i.i.d. with $X_i \sim \text{SSub}(\sigma^2)$. Then

$$\mathbb{E}(\|X\|_2^2) = M\sigma^2 \quad (3.76)$$

and for any $\varepsilon > 0$,

$$\mathbb{P}(|\|X\|_2^2 - M\sigma^2| \geq \varepsilon M\sigma^2) \leq 2\exp\left(-\frac{M\varepsilon^2}{\kappa^*}\right) \quad (3.77)$$

with $\kappa^* = 2/(1 - \log(2)) \approx 6.52$.

Proof:

Since each $X_i \sim \text{SSub}(\sigma^2)$, we have that $X_i \sim \text{Sub}(\sigma^2)$ and $\mathbb{E}(X_i^2) = \sigma^2$, in which case we may apply Theorem 3.10, p. 36 with $\alpha = 1 - \varepsilon$ and $\beta = 1 + \varepsilon$. This allows us to simplify and combine the bounds in (3.63) and (3.64) to obtain (3.77). The value of κ^* follows from the observation that $1 + \varepsilon \leq 2$ so that we can set $\beta_{\max} = 2$.

Finally, from Corollary 3.2, p. 38 we also have the following additional useful corollary. This result generalizes the main results of [1] to the broader family of general strictly sub-Gaussian distributions via a much simpler proof.

Corollary 3.3:

Suppose that Φ is an $M \times N$ matrix whose entries ϕ_{ij} are i.i.d. with $\phi_{ij} \sim \text{SSub}(1/M)$. Let $Y = \Phi x$ for $x \in \mathbb{R}^N$. Then for any $\varepsilon > 0$, and any $x \in \mathbb{R}^N$,

$$\mathbb{E}(\|Y\|_2^2) = \|x\|_2^2 \quad (3.78)$$

and

$$\mathbb{P}(|\|Y\|_2^2 - \|x\|_2^2| \geq \varepsilon \|x\|_2^2) \leq 2\exp\left(-\frac{M\varepsilon^2}{\kappa^*}\right) \quad (3.79)$$

with $\kappa^* = 2/(1 - \log(2)) \approx 6.52$.

Proof:

Let ϕ_i denote the i^{th} row of Φ . Observe that if Y_i denotes the first element of Y , then $Y_i = \langle \phi_i, x \rangle$, and thus by Lemma 2 from "Sub-Gaussian random variables" (Section 3.7), $Y_i \sim \text{SSub}(\|x\|_2^2/M)$. Applying Corollary 3.2, p. 38 to the M -dimensional random vector Y , we obtain (3.79).

3.9 Proof of the RIP for sub-Gaussian matrices¹⁹

We now show how to exploit the concentration of measure (Section 3.8) properties of sub-Gaussian distributions (Section 3.7) to provide a simple proof that sub-Gaussian matrices satisfy the restricted isometry

¹⁸Corollary 3.2, p. 38 exploits the strictness in the strictly sub-Gaussian distribution twice — first to ensure that $\beta \in (1, 2]$ is an admissible range for β and then to simplify the computation of κ^* . One could easily establish a different version of this corollary for non-strictly sub-Gaussian vectors but for which we consider a more restricted range of ε provided that $c^2/\sigma^2 < 2$. However, since most of the distributions of interest in this thesis are indeed strictly sub-Gaussian, we do not pursue this route. Note also that if one is interested in very small ε , then there is considerable room for improvement in the constant C^* .

¹⁹This content is available online at <http://cnx.org/content/m37186/1.4/>.

property (Section 3.3) (RIP). Specifically, we wish to show that by constructing an $M \times N$ matrix Φ at random with M sufficiently large, then with high probability there exists a $\delta_K \in (0, 1)$ such that

$$(1 - \delta_K) \|x\|_2^2 \leq \|\Phi x\|_2^2 \leq (1 + \delta_K) \|x\|_2^2 \quad (3.80)$$

holds for all $x \in \Sigma_K$ (where Σ_K denotes the set of all signals x with at most K nonzeros).

We begin by observing that if all we require is that $\delta_{2K} > 0$, then we may set $M = 2K$ and draw a Φ according to a Gaussian distribution, or indeed any continuous univariate distribution. In this case, with probability 1, any subset of $2K$ columns will be linearly independent, and hence all subsets of $2K$ columns will be bounded below by $1 - \delta_{2K}$ where $\delta_{2K} > 0$. However, suppose we wish to know the constant δ_{2K} . In order to find the value of the constant we must consider all possible $\binom{N}{K} K$ -dimensional subspaces of $\mathbb{R}^N$. From a computational perspective, this is impossible for any realistic values of N and K . Moreover, in light of the lower bounds we described earlier in this course, the actual value of δ_{2K} in this case is likely to be very close to 1. Thus, we focus instead on the problem of achieving the RIP of order $2K$ for a specified constant δ_{2K} .

To ensure that the matrix will satisfy the RIP, we will impose two conditions on the random distribution. First, we require that the distribution is sub-Gaussian. In order to simplify our argument, we will use the simpler results stated in Corollary 2 from "Concentration of measure for sub-Gaussian random variables" (Section 3.8), which we briefly recall.

Corollary 3.4:

Suppose that Φ is an $M \times N$ matrix whose entries ϕ_{ij} are i.i.d. with $\phi_{ij} \sim \text{SSub}(1/M)$. Let $Y = \Phi x$ for $x \in \mathbb{R}^N$. Then for any $\varepsilon > 0$, and any $x \in \mathbb{R}^N$,

$$\mathbb{E}(\|Y\|_2^2) = \|x\|_2^2 \quad (3.81)$$

and

$$\mathbb{P}(|\|Y\|_2^2 - \|x\|_2^2| \geq \varepsilon \|x\|_2^2) \leq 2 \exp\left(-\frac{M\varepsilon^2}{\kappa^*}\right) \quad (3.82)$$

with $\kappa^* = 2/(1 - \log(2)) \approx 6.52$.

By exploiting this theorem, we assume that the distribution used to construct Φ is *strictly* sub-Gaussian. This is done simply to yield more concrete constants. The argument could easily be modified to establish a similar result for general sub-Gaussian distributions by instead using Theorem 2 from "Concentration of measure for sub-Gaussian random variables" (Section 3.8).

Our second condition is that we require that the distribution yield a matrix that is approximately norm-preserving, which will require that

$$\mathbb{E}(\phi_{ij}^2) = \frac{1}{M}, \quad (3.83)$$

and hence the variance is $1/M$.

We shall now show how the concentration of measure inequality in Corollary 3.4, p. 39 can be used together with covering arguments to prove the RIP for sub-Gaussian random matrices. Our general approach will be to construct nets of points in each K -dimensional subspace, apply (3.82) to all of these points through a union bound, and then extend the result from our finite set of points to all possible K -dimensional signals. Thus, in order to prove the result, we will require the following upper bound on the number of points required to construct the nets of points. (For an overview of results similar to Lemma 3.11, p. 39 and of various related concentration of measure results, we refer the reader to the excellent introduction of [5].)

Lemma 3.11:

Let $\varepsilon \in (0, 1)$ be given. There exists a set of points Q such that $|Q| \leq (3/\varepsilon)^K$ and for any $x \in \mathbb{R}^K$ with $\|x\|_2 \leq 1$ there is a point $q \in Q$ satisfying $\|x - q\|_2 \leq \varepsilon$.

Proof:

We construct Q in a greedy fashion. We first select an arbitrary point $q_1 \in \mathbb{R}^K$ with $\|q_1\|_2 \leq 1$. We then continue adding points to Q so that at step i we add a point $q_i \in \mathbb{R}^K$ with $\|q_i\|_2 \leq 1$ which satisfies $\|q_i - q_j\|_2 > \varepsilon$ for all $j < i$. This continues until we can add no more points (and hence for any $x \in \mathbb{R}^K$ with $\|x\|_2 \leq 1$ there is a point $q \in Q$ satisfying $\|x - q\|_2 \leq \varepsilon$.) Now we wish to bound $|Q|$. Observe that if we center balls of radius $\varepsilon/2$ at each point in Q , then these balls are disjoint and lie within a ball of radius $1 + \varepsilon/2$. Thus, if $B^K(r)$ denotes a ball of radius r in $\mathbb{R}^K$, then

$$|Q| \cdot \text{Vol}(B^K(\varepsilon/2)) \leq \text{Vol}(B^K(1 + \varepsilon/2)) \quad (3.84)$$

and hence

$$\begin{aligned} |Q| &\leq \frac{\text{Vol}(B^K(1 + \varepsilon/2))}{\text{Vol}(B^K(\varepsilon/2))} \\ &= \frac{(1 + \varepsilon/2)^K}{(\varepsilon/2)^K} \\ &\leq (3/\varepsilon)^K. \end{aligned} \quad (3.85)$$

We now turn to our main theorem, which is based on the proof given in [7].

Theorem 3.11:

Fix $\delta \in (0, 1)$. Let Φ be an $M \times N$ random matrix whose entries ϕ_{ij} are i.i.d. with $\phi_{ij} \sim \text{SSub}(1/M)$. If

$$M \geq \kappa_1 K \log\left(\frac{N}{K}\right), \quad (3.86)$$

then Φ satisfies the RIP of order K with the prescribed δ with probability exceeding $1 - 2e^{-\kappa_2 M}$, where κ_1 is arbitrary and $\kappa_2 = \delta^2/2\kappa^* - \log(42e/\delta)/\kappa_1$.

Proof:

First note that it is enough to prove (3.80) in the case $\|x\|_2 = 1$, since Φ is linear. Next, fix an index set $T \subset \{1, 2, \dots, N\}$ with $|T| = K$, and let X_T denote the K -dimensional subspace spanned by the columns of Φ_T . We choose a finite set of points Q_T such that $Q_T \subseteq X_T$, $\|q\|_2 \leq 1$ for all $q \in Q_T$, and for all $x \in X_T$ with $\|x\|_2 \leq 1$ we have

$$\min_{q \in Q_T} \|x - q\|_2 \leq \delta/14. \quad (3.87)$$

From Lemma 3.11, p. 39, we can choose such a set Q_T with $|Q_T| \leq (42/\delta)^K$. We then repeat this process for each possible index set T , and collect all the sets Q_T together:

$$Q = \bigcup_{T: |T|=K} Q_T. \quad (3.88)$$

There are $\binom{N}{K}$ possible index sets T . We can bound this number by

$$\binom{N}{K} = \frac{N(N-1)(N-2)\cdots(N-K+1)}{K!} \leq \frac{N^K}{K!} \leq \left(\frac{eN}{K}\right)^K, \quad (3.89)$$

where the last inequality follows since from Sterling's approximation we have $K! \geq (K/e)^K$. Hence $|Q| \leq (42eN/\delta K)^K$. Since the entries of Φ are drawn according to a strictly sub-Gaussian

distribution, from Corollary 3.4, p. 39 we have (3.82). We next use the union bound to apply (3.82) to this set of points with $\varepsilon = \delta/\sqrt{2}$, with the result that, with probability exceeding

$$1 - 2(42eN/\delta K)^K e^{-M\delta^2/2\kappa^*}, \quad (3.90)$$

we have

$$(1 - \delta/\sqrt{2}) \|q\|_2^2 \leq \|\Phi q\|_2^2 \leq (1 + \delta/\sqrt{2}) \|q\|_2^2, \quad \text{for all } q \in Q. \quad (3.91)$$

We observe that if M satisfies (3.86) then

$$\log\left(\frac{42eN}{\delta K}\right)^K \leq K \log\left(\frac{N}{K}\right) \log\left(\frac{42e}{\delta}\right) \leq \frac{M \log(42e/\delta)}{\kappa_1} \quad (3.92)$$

and thus (3.90) exceeds $1 - 2e^{-\kappa_2 M}$ as desired.

We now define A as the smallest number such that

$$\|\Phi x\|_2 \leq \sqrt{1+A} \|x\|_2, \quad \text{for all } x \in \Sigma_K, \|x\|_2 \leq 1. \quad (3.93)$$

Our goal is to show that $A \leq \delta$. For this, we recall that for any $x \in \Sigma_K$ with $\|x\|_2 \leq 1$, we can pick a $q \in Q$ such that $\|x - q\|_2 \leq \delta/14$ and such that $x - q \in \Sigma_K$ (since if $x \in X_T$, we can pick $q \in Q_T \subset X_T$ satisfying $\|x - q\|_2 \leq \delta/14$). In this case we have

$$\|\Phi x\|_2 \leq \|\Phi q\|_2 + \|\Phi(x - q)\|_2 \leq \sqrt{1 + \delta/\sqrt{2}} + \sqrt{1+A} \cdot \delta/14. \quad (3.94)$$

Since by definition A is the smallest number for which (3.93) holds, we obtain $\sqrt{1+A} \leq \sqrt{1 + \delta/\sqrt{2}} + \sqrt{1+A} \cdot \delta/14$. Therefore

$$\sqrt{1+A} \leq \frac{\sqrt{1 + \delta/\sqrt{2}}}{1 - \delta/14} \leq \sqrt{1 + \delta}, \quad (3.95)$$

as desired. We have proved the upper inequality in (3.80). The lower inequality follows from this since

$$\|\Phi x\|_2 \geq \|\Phi q\|_2 - \|\Phi(x - q)\|_2 \geq \sqrt{1 - \delta/\sqrt{2}} - \sqrt{1 + \delta} \cdot \delta/14 \geq \sqrt{1 - \delta}, \quad (3.96)$$

which completes the proof.

Above we prove above that the RIP holds with high probability when the matrix Φ is drawn according to a strictly sub-Gaussian distribution. However, we are often interested in signals that are sparse or compressible in some orthonormal basis $\Psi \neq I$, in which case we would like the matrix $\Phi\Psi$ to satisfy the RIP. In this setting it is easy to see that by choosing our net of points in the K -dimensional subspaces spanned by sets of K columns of Ψ , Theorem 3.11, p. 40 will establish the RIP for $\Phi\Psi$ for Φ again drawn from a sub-Gaussian distribution. This *universality* of Φ with respect to the sparsity-inducing basis is an attractive property that was initially observed for the Gaussian distribution (based on symmetry arguments), but we can now see is a property of more general sub-Gaussian distributions. Indeed, it follows that with high probability such a Φ will simultaneously satisfy the RIP with respect to an exponential number of fixed bases.

Chapter 4

ℓ_1 -norm minimization

4.1 Signal recovery via ℓ_1 -norm minimization¹

As we will see later in this course², there now exist a wide variety of approaches to recover a sparse (Section 2.3) signal x from a small number of linear measurements. We begin by considering a natural first approach to the problem of sparse recovery.

Given measurements $y = \Phi x$ and the knowledge that our original signal x is sparse or compressible (Section 2.4), it is natural to attempt to recover x by solving an optimization problem of the form

$$\hat{x} = \underset{z}{\operatorname{argmin}} \|z\|_0 \quad \text{subject to} \quad z \in \mathcal{B}(y), \quad (4.1)$$

where $\mathcal{B}(y)$ ensures that $\hat{x}$ is consistent with the measurements y . Recall that $\|z\|_0 = |\operatorname{supp}(z)|$ simply counts the number of nonzero entries in z , so (4.1) simply seeks out the sparsest signal consistent with the observed measurements. For example, if our measurements are exact and noise-free, then we can set $\mathcal{B}(y) = \{z : \Phi z = y\}$. When the measurements have been contaminated with a small amount of bounded noise, we could instead set $\mathcal{B}(y) = \{z : \|\Phi z - y\|_2 \leq \varepsilon\}$. In both cases, (4.1) finds the sparsest x that is consistent with the measurements y .

Note that in (4.1) we are inherently assuming that x itself is sparse. In the more common setting where $x = \Psi\alpha$, we can easily modify the approach and instead consider

$$\hat{\alpha} = \underset{z}{\operatorname{argmin}} \|z\|_0 \quad \text{subject to} \quad z \in \mathcal{B}(y) \quad (4.2)$$

where $\mathcal{B}(y) = \{z : \Phi\Psi z = y\}$ or $\mathcal{B}(y) = \{z : \|\Phi\Psi z - y\|_2 \leq \varepsilon\}$. By setting $\tilde{\Phi} = \Phi\Psi$ we see that (4.1) and (4.2) are essentially identical. Moreover, as noted in "Matrices that satisfy the RIP" (Section 3.5), in many cases the introduction of Ψ does not significantly complicate the construction of matrices Φ such that $\tilde{\Phi}$ will satisfy the desired properties. Thus, for most of the remainder of this course we will restrict our attention to the case where $\Psi = I$. It is important to note, however, that this restriction does impose certain limits in our analysis when Ψ is a general dictionary and not an orthonormal basis. For example, in this case $\|\hat{x} - x\|_2 = \|\Psi\hat{c} - \Psi c\|_2 \neq \|\hat{\alpha} - \alpha\|_2$, and thus a bound on $\|\hat{c} - c\|_2$ cannot directly be translated into a bound on $\|\hat{x} - x\|_2$, which is often the metric of interest.

Although it is possible to analyze the performance of (4.1) under the appropriate assumptions on Φ , we do not pursue this strategy since the objective function $\|\cdot\|_0$ is nonconvex, and hence (4.1) is potentially very difficult to solve. In fact, one can show that for a general matrix Φ , even finding a solution that approximates

¹This content is available online at <http://cnx.org/content/m37179/1.5/>.

²An Introduction to Compressive Sensing <http://cnx.org/content/col11133/latest/>

the true minimum is NP-hard. One avenue for translating this problem into something more tractable is to replace $\|\cdot\|_0$ with its convex approximation $\|\cdot\|_1$. Specifically, we consider

$$\hat{x} = \underset{z}{\operatorname{argmin}} \|z\|_1 \quad \text{subject to} \quad z \in \mathcal{B}(y). \quad (4.3)$$

Provided that $\mathcal{B}(y)$ is convex, (4.3) is computationally feasible. In fact, when $\mathcal{B}(y) = \{z : \Phi z = y\}$, the resulting problem can be posed as a linear program [33].

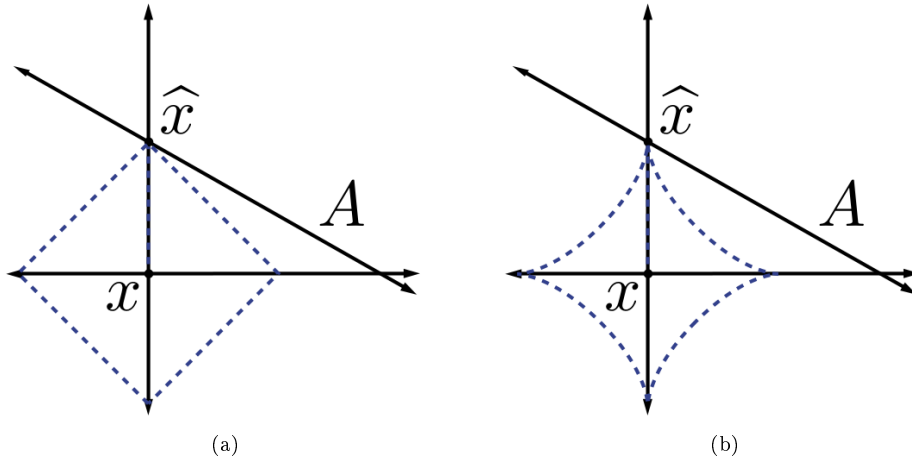


Figure 4.1: Best approximation of a point in $\mathbb{R}^2$ by a one-dimensional subspace using the ℓ_1 norm and the ℓ_p quasinorm with $p = \frac{1}{2}$. (a) Approximation in ℓ_1 norm (b) Approximation in ℓ_p quasinorm

It is clear that replacing (4.1) with (4.3) transforms a computationally intractable problem into a tractable one, but it may not be immediately obvious that the solution to (4.3) will be at all similar to the solution to (4.1). However, there are certainly intuitive reasons to expect that the use of ℓ_1 minimization will indeed promote sparsity. As an example, recall the example we discussed earlier shown in Figure 4.1. In this case the solutions to the ℓ_1 minimization problem coincided exactly with the solution to the ℓ_p minimization problem for any $p < 1$, and notably, is sparse. Moreover, the use of ℓ_1 minimization to promote or exploit sparsity has a long history, dating back at least to the work of Beurling on Fourier transform extrapolation from partial observations [14].

Additionally, in a somewhat different context, in 1965 Logan [118] showed that a bandlimited signal can be perfectly recovered in the presence of *arbitrary* corruptions on a small interval. Again, the recovery method consists of searching for the bandlimited signal that is closest to the observed signal in the ℓ_1 norm. This can be viewed as further validation of the intuition gained from Figure 4.1 — the ℓ_1 norm is well-suited to sparse errors.

Historically, the use of ℓ_1 minimization on large problems finally became practical with the explosion of computing power in the late 1970's and early 1980's. In one of its first applications, it was demonstrated that geophysical signals consisting of spike trains could be recovered from only the high-frequency components of these signals by exploiting ℓ_1 minimization [117], [165], [186]. Finally, in the 1990's there was renewed interest in these approaches within the signal processing community for the purpose of finding sparse approximations (Section 2.4) to signals and images when represented in overcomplete dictionaries or unions of bases [33], [125]. Separately, ℓ_1 minimization received significant attention in the statistics literature as a method for variable selection in linear regression (Section 5.1), known as the Lasso [168].

Thus, there are a variety of reasons to suspect that ℓ_1 minimization will provide an accurate method for sparse signal recovery. More importantly, this also provides a computationally tractable approach to the sparse signal recovery problem. We now provide an overview of ℓ_1 minimization in both the noise-free (Section 4.2) and noisy (Section 4.3) settings from a theoretical perspective. We will then further discuss algorithms for performing ℓ_1 minimization (Section 4.7) later in this course³.

4.2 Noise-free signal recovery⁴

We now begin our analysis of

$$\hat{x} = \underset{z}{\operatorname{argmin}} \|z\|_1 \quad \text{subject to} \quad z \in \mathcal{B}(y). \quad (4.4)$$

for various specific choices of $\mathcal{B}(y)$. In order to do so, we require the following general result which builds on Lemma 4 from " ℓ_1 minimization proof"⁵. The key ideas in this proof follow from [22].

Lemma 4.1:

Suppose that Φ satisfies the restricted isometry property (Section 3.3) (RIP) of order $2K$ with $\delta_{2K} < \sqrt{2} - 1$. Let $x, \hat{x} \in \mathbb{R}^N$ be given, and define $h = \hat{x} - x$. Let Λ_0 denote the index set corresponding to the K entries of x with largest magnitude and Λ_1 the index set corresponding to the K entries of $h_{\Lambda_0^c}$ with largest magnitude. Set $\Lambda = \Lambda_0 \cup \Lambda_1$. If $\|\hat{x}\|_1 \leq \|x\|_1$, then

$$\|h\|_2 \leq C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_1 \frac{|\langle \Phi h_{\Lambda}, \Phi h \rangle|}{\|h_{\Lambda}\|_2}. \quad (4.5)$$

where

$$C_0 = 2 \frac{1 - (1 - \sqrt{2}) \delta_{2K}}{1 - (1 + \sqrt{2}) \delta_{2K}}, \quad C_1 = \frac{2}{1 - (1 + \sqrt{2}) \delta_{2K}}. \quad (4.6)$$

Proof:

We begin by observing that $h = h_{\Lambda} + h_{\Lambda^c}$, so that from the triangle inequality

$$\|h\|_2 \leq \|h_{\Lambda}\|_2 + \|h_{\Lambda^c}\|_2. \quad (4.7)$$

We first aim to bound $\|h_{\Lambda^c}\|_2$. From Lemma 3 from " ℓ_1 minimization proof"⁶ we have

$$\|h_{\Lambda^c}\|_2 = \left\| \sum_{j \geq 2} h_{\Lambda_j} \right\|_2 \leq \sum_{j \geq 2} \|h_{\Lambda_j}\|_2 \leq \frac{\|h_{\Lambda_0^c}\|_1}{\sqrt{K}}, \quad (4.8)$$

where the Λ_j are defined as before, i.e., Λ_1 is the index set corresponding to the K largest entries of $h_{\Lambda_0^c}$ (in absolute value), Λ_2 as the index set corresponding to the next K largest entries, and so on.

We now wish to bound $\|h_{\Lambda_0^c}\|_1$. Since $\|x\|_1 \geq \|\hat{x}\|_1$, by applying the triangle inequality we obtain

$$\begin{aligned} \|x\|_1 &\geq \|x + h\|_1 = \|x_{\Lambda_0} + h_{\Lambda_0}\|_1 + \|x_{\Lambda_0^c} + h_{\Lambda_0^c}\|_1 \\ &\geq \|x_{\Lambda_0}\|_1 - \|h_{\Lambda_0}\|_1 + \|h_{\Lambda_0^c}\|_1 - \|x_{\Lambda_0^c}\|_1. \end{aligned} \quad (4.9)$$

³ *An Introduction to Compressive Sensing* <<http://cnx.org/content/col11133/latest/>>

⁴ This content is available online at <<http://cnx.org/content/m37181/1.6/>>.

⁵ " ℓ_1 minimization proof" <<http://cnx.org/content/m37187/latest/>>

⁶ " ℓ_1 minimization proof" <<http://cnx.org/content/m37187/latest/>>

Rearranging and again applying the triangle inequality,

$$\begin{aligned}\|h_{\Lambda_0^c}\|_1 &\leq \|x\|_1 - \|x_{\Lambda_0}\|_1 + \|h_{\Lambda_0}\|_1 + \|x_{\Lambda_0^c}\|_1 \\ &\leq \|x - x_{\Lambda_0}\|_1 + \|h_{\Lambda_0}\|_1 + \|x_{\Lambda_0^c}\|_1.\end{aligned}\quad (4.10)$$

Recalling that $\sigma_K(x)_1 = \|x_{\Lambda_0^c}\|_1 = \|x - x_{\Lambda_0}\|_1$,

$$\|h_{\Lambda_0^c}\|_1 \leq \|h_{\Lambda_0}\|_1 + 2\sigma_K(x)_1. \quad (4.11)$$

Combining this with (4.8) we obtain

$$\|h_{\Lambda^c}\|_2 \leq \frac{\|h_{\Lambda_0}\|_1 + 2\sigma_K(x)_1}{\sqrt{K}} \leq \|h_{\Lambda_0}\|_2 + 2\frac{\sigma_K(x)_1}{\sqrt{K}} \quad (4.12)$$

where the last inequality follows from standard bounds on ℓ_p norms (Lemma 1 from "The RIP and the NSP" (Section 3.4)). By observing that $\|h_{\Lambda_0}\|_2 \leq \|h_{\Lambda}\|_2$ this combines with (4.7) to yield

$$\|h\|_2 \leq 2\|h_{\Lambda}\|_2 + 2\frac{\sigma_K(x)_1}{\sqrt{K}}. \quad (4.13)$$

We now turn to establishing a bound for $\|h_{\Lambda}\|_2$. Combining Lemma 4 from " ℓ_1 minimization proof"⁷ with (4.11) and again applying standard bounds on ℓ_p norms we obtain

$$\begin{aligned}\|h_{\Lambda}\|_2 &\leq \alpha \frac{\|h_{\Lambda_0^c}\|_1}{\sqrt{K}} + \beta \frac{|\langle \Phi h_{\Lambda}, \Phi h \rangle|}{\|h_{\Lambda}\|_2} \\ &\leq \alpha \frac{\|h_{\Lambda_0}\|_1 + 2\sigma_K(x)_1}{\sqrt{K}} + \beta \frac{|\langle \Phi h_{\Lambda}, \Phi h \rangle|}{\|h_{\Lambda}\|_2} \\ &\leq \alpha \|h_{\Lambda_0}\|_2 + 2\alpha \frac{\sigma_K(x)_1}{\sqrt{K}} + \beta \frac{|\langle \Phi h_{\Lambda}, \Phi h \rangle|}{\|h_{\Lambda}\|_2}.\end{aligned}\quad (4.14)$$

Since $\|h_{\Lambda_0}\|_2 \leq \|h_{\Lambda}\|_2$,

$$(1 - \alpha) \|h_{\Lambda}\|_2 \leq 2\alpha \frac{\sigma_K(x)_1}{\sqrt{K}} + \beta \frac{|\langle \Phi h_{\Lambda}, \Phi h \rangle|}{\|h_{\Lambda}\|_2}. \quad (4.15)$$

The assumption that $\delta_{2K} < \sqrt{2} - 1$ ensures that $\alpha < 1$. Dividing by $(1 - \alpha)$ and combining with (4.13) results in

$$\|h\|_2 \leq \left(\frac{4\alpha}{1 - \alpha} + 2 \right) \frac{\sigma_K(x)_1}{\sqrt{K}} + \frac{2\beta}{1 - \alpha} \frac{|\langle \Phi h_{\Lambda}, \Phi h \rangle|}{\|h_{\Lambda}\|_2}. \quad (4.16)$$

Plugging in for α and β yields the desired constants.

Lemma 4.1, p. 45 establishes an error bound for the class of ℓ_1 minimization algorithms described by (4.4) when combined with a measurement matrix Φ satisfying the RIP. In order to obtain specific bounds for concrete examples of $\mathcal{B}(y)$, we must examine how requiring $\hat{x} \in \mathcal{B}(y)$ affects $|\langle \Phi h_{\Lambda}, \Phi h \rangle|$. As an example, in the case of noise-free measurements we obtain the following theorem.

Theorem 4.1: (Theorem 1.1 of [22])

Suppose that Φ satisfies the RIP of order $2K$ with $\delta_{2K} < \sqrt{2} - 1$ and we obtain measurements of the form $y = \Phi x$. Then when $\mathcal{B}(y) = \{z : \Phi z = y\}$, the solution $\hat{x}$ to (4.4) obeys

$$\|\hat{x} - x\|_2 \leq C_0 \frac{\sigma_K(x)_1}{\sqrt{K}}. \quad (4.17)$$

⁷" ℓ_1 minimization proof" <<http://cnx.org/content/m37187/latest/>>

Proof:

Since $x \in \mathcal{B}(y)$ we can apply Lemma 4.1, p. 45 to obtain that for $h = \hat{x} - x$,

$$\|h\|_2 \leq C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_1 \frac{|\langle \Phi h_\Lambda, \Phi h \rangle|}{\|h_\Lambda\|_2}. \quad (4.18)$$

Furthermore, since $\hat{x} \in \mathcal{B}(y)$ we also have that $y = \Phi \hat{x} = \Phi x$ and hence $\Phi h = 0$. Therefore the second term vanishes, and we obtain the desired result.

Theorem 4.1, (Theorem 1.1 of [22]), p. 46 is rather remarkable. By considering the case where $x \in \Sigma_K = \{x : \|x\|_0 \leq K\}$ we can see that provided Φ satisfies the RIP — which as shown earlier allows for as few as $O(K \log(N/K))$ measurements — we can recover any K -sparse x exactly. This result seems improbable on its own, and so one might expect that the procedure would be highly sensitive to noise, but we will see next that Lemma 4.1, p. 45 can also be used to demonstrate that this approach is actually stable.

Note that Theorem 4.1, (Theorem 1.1 of [22]), p. 46 assumes that Φ satisfies the RIP. One could easily modify the argument to replace this with the assumption that Φ satisfies the null space property (Section 3.2) (NSP) instead. Specifically, if we are only interested in the noiseless setting, in which case h lies in the null space of Φ , then Lemma 4.1, p. 45 simplifies and its proof could be broken into two steps: (i) show that if Φ satisfies the RIP then it satisfies the NSP (as shown in "The RIP and the NSP" (Section 3.4)), and (ii) the NSP implies the simplified version of Lemma 4.1, p. 45. This proof directly mirrors that of Lemma 4.1, p. 45. Thus, by the same argument as in the proof of Theorem 4.1, (Theorem 1.1 of [22]), p. 46, it is straightforward to show that if Φ satisfies the NSP then it will obey the same error bound.

4.3 Signal recovery in noise⁸

The ability to perfectly reconstruct a sparse (Section 2.3) signal from noise-free (Section 4.2) measurements represents a promising result. However, in most real-world systems the measurements are likely to be contaminated by some form of noise. For instance, in order to process data in a computer we must be able to represent it using a finite number of bits, and hence the measurements will typically be subject to quantization error. Moreover, systems which are implemented in physical hardware will be subject to a variety of different types of noise depending on the setting.

Perhaps somewhat surprisingly, one can show that it is possible to modify

$$\hat{x} = \underset{z}{\operatorname{argmin}} \|z\|_1 \quad \text{subject to} \quad z \in \mathcal{B}(y). \quad (4.19)$$

to stably recover sparse signals under a variety of common noise models [27], [31], [98]. As might be expected, the restricted isometry property (Section 3.3) (RIP) is extremely useful in establishing performance guarantees in noise.

In our analysis we will make repeated use of Lemma 1 from "Noise-free signal recovery" (Section 4.2), so we repeat it here for convenience.

Lemma 4.2:

Suppose that Φ satisfies the RIP of order $2K$ with $\delta_{2K} < \sqrt{2} - 1$. Let $x, \hat{x} \in \mathbb{R}^N$ be given, and define $h = \hat{x} - x$. Let Λ_0 denote the index set corresponding to the K entries of x with largest magnitude and Λ_1 the index set corresponding to the K entries of $h_{\Lambda_0^c}$ with largest magnitude. Set $\Lambda = \Lambda_0 \cup \Lambda_1$. If $\|\hat{x}\|_1 \leq \|x\|_1$, then

$$\|h\|_2 \leq C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_1 \frac{|\langle \Phi h_\Lambda, \Phi h \rangle|}{\|h_\Lambda\|_2}. \quad (4.20)$$

⁸This content is available online at <http://cnx.org/content/m37182/1.5/>.

where

$$C_0 = 2 \frac{1 - (1 - \sqrt{2}) \delta_{2K}}{1 - (1 + \sqrt{2}) \delta_{2K}}, \quad C_1 = \frac{2}{1 - (1 + \sqrt{2}) \delta_{2K}}. \quad (4.21)$$

4.3.1 Bounded noise

We first provide a bound on the worst-case performance for uniformly bounded noise, as first investigated in [27].

Theorem 4.2: (Theorem 1.2 of [23])

Suppose that Φ satisfies the RIP of order $2K$ with $\delta_{2K} < \sqrt{2} - 1$ and let $y = \Phi x + e$ where $\|e\|_2 \leq \varepsilon$.

Then when $\mathcal{B}(y) = \{z : \|\Phi z - y\|_2 \leq \varepsilon\}$, the solution $\hat{x}$ to (4.19) obeys

$$\|\hat{x} - x\|_2 \leq C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_2 \varepsilon, \quad (4.22)$$

where

$$C_0 = 2 \frac{1 - (1 - \sqrt{2}) \delta_{2K}}{1 - (1 + \sqrt{2}) \delta_{2K}}, \quad C_2 = 4 \frac{\sqrt{1 + \delta_{2K}}}{1 - (1 + \sqrt{2}) \delta_{2K}}. \quad (4.23)$$

Proof:

We are interested in bounding $\|h\|_2 = \|\hat{x} - x\|_2$. Since $\|e\|_2 \leq \varepsilon$, $x \in \mathcal{B}(y)$, and therefore we know that $\|\hat{x}\|_1 \leq \|x\|_1$. Thus we may apply Lemma 4.2, p. 47, and it remains to bound $|\langle \Phi h_\Lambda, \Phi h \rangle|$. To do this, we observe that

$$\|\Phi h\|_2 = \left\| \Phi \begin{pmatrix} \hat{x} \\ -x \end{pmatrix} \right\|_2 = \|\Phi \hat{x} - y + y - \Phi x\|_2 \leq \|\Phi \hat{x} - y\|_2 + \|y - \Phi x\|_2 \leq 2\varepsilon \quad (4.24)$$

where the last inequality follows since $\hat{x} \in \mathcal{B}(y)$. Combining this with the RIP and the Cauchy-Schwarz inequality we obtain

$$|\langle \Phi h_\Lambda, \Phi h \rangle| \leq \|\Phi h_\Lambda\|_2 \|\Phi h\|_2 \leq 2\varepsilon \sqrt{1 + \delta_{2K}} \|\Phi h_\Lambda\|_2. \quad (4.25)$$

Thus,

$$\|h\|_2 \leq C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_1 2\varepsilon \sqrt{1 + \delta_{2K}} = C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_2 \varepsilon, \quad (4.26)$$

completing the proof.

In order to place this result in context, consider how we would recover a sparse vector x if we happened to already know the K locations of the nonzero coefficients, which we denote by Λ_0 . This is referred to as the *oracle estimator*. In this case a natural approach is to reconstruct the signal using a simple pseudoinverse:

$$\begin{aligned} \hat{x}_{\Lambda_0} &= \Phi_{\Lambda_0}^\dagger y = (\Phi_{\Lambda_0}^T \Phi_{\Lambda_0})^{-1} \Phi_{\Lambda_0}^T y \\ \hat{x}_{\Lambda_0^c} &= 0. \end{aligned} \quad (4.27)$$

The implicit assumption in (4.27) is that Φ_{Λ_0} has full column-rank (and hence we are considering the case where Φ_{Λ_0} is the $M \times K$ matrix with the columns indexed by Λ_0^c removed) so that there is a unique solution to the equation $y = \Phi_{\Lambda_0} x_{\Lambda_0}$. With this choice, the recovery error is given by

$$\|\hat{x} - x\|_2 = \|(\Phi_{\Lambda_0}^T \Phi_{\Lambda_0})^{-1} \Phi_{\Lambda_0}^T (\Phi x + e) - x\|_2 = \|(\Phi_{\Lambda_0}^T \Phi_{\Lambda_0})^{-1} \Phi_{\Lambda_0}^T e\|_2. \quad (4.28)$$

We now consider the worst-case bound for this error. Using standard properties of the singular value decomposition, it is straightforward to show that if Φ satisfies the RIP of order $2K$ (with constant δ_{2K}), then the largest singular value of $\Phi_{\Lambda_0}^T$ lies in the range $[1/\sqrt{1+\delta_{2K}}, 1/\sqrt{1-\delta_{2K}}]$. Thus, if we consider the worst-case recovery error over all e such that $\|e\|_2 \leq \varepsilon$, then the recovery error can be bounded by

$$\frac{\varepsilon}{\sqrt{1+\delta_{2K}}} \leq \|\hat{x} - x\|_2 \leq \frac{\varepsilon}{\sqrt{1-\delta_{2K}}}. \quad (4.29)$$

Therefore, if x is exactly K -sparse, then the guarantee for the pseudoinverse recovery method, which is given *perfect knowledge of the true support of x* , cannot improve upon the bound in Theorem 4.2, (Theorem 1.2 of [23]), p. 48 by more than a constant value.

We now examine a slightly different noise model. Whereas Theorem 4.2, (Theorem 1.2 of [23]), p. 48 assumed that the noise norm $\|e\|_2$ was small, the theorem below analyzes a different recovery algorithm known as the *Dantzig selector* in the case where $\|\Phi^T e\|_\infty$ is small [31]. We will see below that this will lead to a simple analysis of the performance of this algorithm in Gaussian noise.

Theorem 4.3:

Suppose that Φ satisfies the RIP of order $2K$ with $\delta_{2K} < \sqrt{2} - 1$ and we obtain measurements of the form $y = \Phi x + e$ where $\|\Phi^T e\|_\infty \leq \lambda$. Then when $\mathcal{B}(y) = \{z : \|\Phi^T (\Phi z - y)\|_\infty \leq \lambda\}$, the solution $\hat{x}$ to (4.19) obeys

$$\|\hat{x} - x\|_2 \leq C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_3 \sqrt{K} \lambda, \quad (4.30)$$

where

$$C_0 = 2 \frac{1 - (1 - \sqrt{2}) \delta_{2K}}{1 - (1 + \sqrt{2}) \delta_{2K}}, \quad C_3 = \frac{4\sqrt{2}}{1 - (1 + \sqrt{2}) \delta_{2K}}. \quad (4.31)$$

Proof:

The proof mirrors that of Theorem 4.2, (Theorem 1.2 of [23]), p. 48. Since $\|\Phi^T e\|_\infty \leq \lambda$, we again have that $x \in \mathcal{B}(y)$, so $\|\hat{x}\|_1 \leq \|x\|_1$ and thus Lemma 4.2, p. 47 applies. We follow a similar approach as in Theorem 4.2, (Theorem 1.2 of [23]), p. 48 to bound $|\langle \Phi h_\Lambda, \Phi h \rangle|$. We first note that

$$\|\Phi^T \Phi h\|_\infty \leq \|\Phi^T (\Phi \hat{x} - y)\|_\infty + \|\Phi^T (y - \Phi x)\|_\infty \leq 2\lambda \quad (4.32)$$

where the last inequality again follows since $x, \hat{x} \in \mathcal{B}(y)$. Next, note that $\Phi h_\Lambda = \Phi_\Lambda h_\Lambda$. Using this we can apply the Cauchy-Schwarz inequality to obtain

$$|\langle \Phi h_\Lambda, \Phi h \rangle| = |\langle h_\Lambda, \Phi_\Lambda^T \Phi h \rangle| \leq \|h_\Lambda\|_2 \|\Phi_\Lambda^T \Phi h\|_2. \quad (4.33)$$

Finally, since $\|\Phi^T \Phi h\|_\infty \leq 2\lambda$, we have that every coefficient of $\Phi^T \Phi h$ is at most 2λ , and thus $\|\Phi_\Lambda^T \Phi h\|_2 \leq \sqrt{2K} (2\lambda)$. Thus,

$$\|h\|_2 \leq C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_1 2\sqrt{2K} \lambda = C_0 \frac{\sigma_K(x)_1}{\sqrt{K}} + C_3 \sqrt{K} \lambda, \quad (4.34)$$

as desired.

4.3.2 Gaussian noise

Finally, we also examine the performance of these approaches in the presence of Gaussian noise. The case of Gaussian noise was first considered in [98], which examined the performance of ℓ_0 minimization with noisy measurements. We now see that Theorem 4.2, (Theorem 1.2 of [23]), p. 48 and Theorem 4.3, p. 49 can be leveraged to provide similar guarantees for ℓ_1 minimization. To simplify our discussion we will restrict our attention to the case where $x \in \Sigma_K = \{x : \|x\|_0 \leq K\}$, so that $\sigma_K(x)_1 = 0$ and the error bounds in Theorem 4.2, (Theorem 1.2 of [23]), p. 48 and Theorem 4.3, p. 49 depend only on the noise e .

To begin, suppose that the coefficients of $e \in \mathbb{R}^M$ are i.i.d. according to a Gaussian distribution with mean zero and variance σ^2 . Since the Gaussian distribution is itself sub-Gaussian, we can apply results such as Corollary 1 from "Concentration of measure for sub-Gaussian random variables" (Section 3.8) to show that there exists a constant $c_0 > 0$ such that for any $\varepsilon > 0$,

$$\mathbb{P}(\|e\|_2 \geq (1 + \varepsilon) \sqrt{M}\sigma) \leq \exp(-c_0 \varepsilon^2 M). \quad (4.35)$$

Applying this result to Theorem 4.2, (Theorem 1.2 of [23]), p. 48 with $\varepsilon = 1$, we obtain the following result for the special case of Gaussian noise.

Corollary 4.1:

Suppose that Φ satisfies the RIP of order $2K$ with $\delta_{2K} < \sqrt{2} - 1$. Furthermore, suppose that $x \in \Sigma_K$ and that we obtain measurements of the form $y = \Phi x + e$ where the entries of e are i.i.d. $\mathcal{N}(0, \sigma^2)$. Then when $\mathcal{B}(y) = \{z : \|\Phi z - y\|_2 \leq 2\sqrt{M}\sigma\}$, the solution $\hat{x}$ to (4.19) obeys

$$\|\hat{x} - x\|_2 \leq 8 \frac{\sqrt{1 + \delta_{2K}}}{1 - (1 + \sqrt{2})\delta_{2K}} \sqrt{M}\sigma \quad (4.36)$$

with probability at least $1 - \exp(-c_0 M)$.

We can similarly consider Theorem 4.3, p. 49 in the context of Gaussian noise. If we assume that the columns of Φ have unit norm, then each coefficient of $\Phi^T e$ is a Gaussian random variable with mean zero and variance σ^2 . Using standard tail bounds for the Gaussian distribution (see Theorem 1 from "Sub-Gaussian random variables" (Section 3.7)), we have that

$$\mathbb{P}(|[\Phi^T e]_i| \geq t\sigma) \leq \exp(-t^2/2) \quad (4.37)$$

for $i = 1, 2, \dots, n$. Thus, using the union bound over the bounds for different i , we obtain

$$\mathbb{P}(\|\Phi^T e\|_\infty \geq 2\sqrt{\log N}\sigma) \leq N \exp(-2\log N) = \frac{1}{N}. \quad (4.38)$$

Applying this to Theorem 4.3, p. 49, we obtain the following result, which is a simplified version of Theorem 1.1 of [31].

Corollary 4.2:

Suppose that Φ has unit-norm columns and satisfies the RIP of order $2K$ with $\delta_{2K} < \sqrt{2} - 1$. Furthermore, suppose that $x \in \Sigma_K$ and that we obtain measurements of the form $y = \Phi x + e$ where the entries of e are i.i.d. $\mathcal{N}(0, \sigma^2)$. Then when $\mathcal{B}(y) = \{z : \|\Phi^T(\Phi z - y)\|_\infty \leq 2\sqrt{\log N}\sigma\}$, the solution $\hat{x}$ to (4.19) obeys

$$\|\hat{x} - x\|_2 \leq 4\sqrt{2} \frac{\sqrt{1 + \delta_{2K}}}{1 - (1 + \sqrt{2})\delta_{2K}} \sqrt{K \log N} \sigma \quad (4.39)$$

with probability at least $1 - \frac{1}{N}$.

Ignoring the precise constants and the probabilities with which the bounds hold (which we have made no effort to optimize), we observe that if $M = O(K \log N)$ then these results appear to be essentially the same. However, there is a subtle difference. Specifically, if M and N are fixed and we consider the effect of varying K , we can see that Corollary 4.2, p. 50 yields a bound that is adaptive to this change, providing a stronger guarantee when K is small, whereas the bound in Corollary 4.1, p. 50 does not improve as K is reduced. Thus, while they provide very similar guarantees, there are certain circumstances where the Dantzig selector is preferable. See [31] for further discussion of the comparative advantages of these approaches.

4.4 Instance-optimal guarantees revisited⁹

We now briefly return to the noise-free (Section 4.2) setting to take a closer look at instance-optimal guarantees for recovering non-sparse signals. To begin, recall that in Theorem 1 from "Noise-free signal recovery" (Section 4.2) we bounded the ℓ_2 -norm of the reconstruction error of

$$\hat{x} = \underset{z}{\operatorname{argmin}} \|z\|_1 \quad \text{subject to} \quad z \in \mathcal{B}(y). \quad (4.40)$$

as

$$\|\hat{x} - x\|_2 \leq C_0 \sigma_K(x)_1 / \sqrt{K} \quad (4.41)$$

when $\mathcal{B}(y) = \{z : \Phi z = y\}$. One can generalize this result to measure the reconstruction error using the ℓ_p -norm for any $p \in [1, 2]$. For example, by a slight modification of these arguments, one can also show that $\|\hat{x} - x\|_1 \leq C_0 \sigma_K(x)_1$ (see [24]). This leads us to ask whether we might replace the bound for the ℓ_2 error with a result of the form $\|\hat{x} - x\|_2 \leq C \sigma_K(x)_2$. Unfortunately, obtaining such a result requires an unreasonably large number of measurements, as quantified by the following theorem of [35].

Theorem 4.4: (Theorem 5.1 of [35])

Suppose that Φ is an $M \times N$ matrix and that $\Delta : \mathbb{R}^M \rightarrow \mathbb{R}^N$ is a recovery algorithm that satisfies

$$\|x - \Delta(\Phi x)\|_2 \leq C \sigma_K(x)_2 \quad (4.42)$$

for some $K \geq 1$, then $M > \left(1 - \sqrt{1 - 1/C^2}\right) N$.

Proof:

We begin by letting $h \in \mathbb{R}^N$ denote any vector in $\mathcal{N}(\Phi)$. We write $h = h_\Lambda + h_{\Lambda^c}$ where Λ is an arbitrary set of indices satisfying $|\Lambda| \leq K$. Set $x = h_{\Lambda^c}$, and note that $\Phi x = \Phi h_{\Lambda^c} = \Phi h - \Phi h_\Lambda = -\Phi h_\Lambda$ since $h \in \mathcal{N}(\Phi)$. Since $h_\Lambda \in \Sigma_K$, (4.42) implies that $\Delta(\Phi x) = \Delta(-\Phi h_\Lambda) = -h_\Lambda$. Hence, $\|x - \Delta(\Phi x)\|_2 = \|h_{\Lambda^c} - (-h_\Lambda)\|_2 = \|h\|_2$. Furthermore, we observe that $\sigma_K(x)_2 \leq \|x\|_2$, since by definition $\sigma_K(x)_2 \leq \|x - \tilde{x}\|_2$ for all $\tilde{x} \in \Sigma_K$, including $\tilde{x} = 0$. Thus $\|h\|_2 \leq C \|h_{\Lambda^c}\|_2$. Since $\|h\|_2^2 = \|h_\Lambda\|_2^2 + \|h_{\Lambda^c}\|_2^2$, this yields

$$\|h_\Lambda\|_2^2 = \|h\|_2^2 - \|h_{\Lambda^c}\|_2^2 \leq \|h\|_2^2 - \frac{1}{C^2} \|h\|_2^2 = \left(1 - \frac{1}{C^2}\right) \|h\|_2^2. \quad (4.43)$$

This must hold for any vector $h \in \mathcal{N}(\Phi)$ and for any set of indices Λ such that $|\Lambda| \leq K$. In particular, let $\{v_i\}_{i=1}^{N-M}$ be an orthonormal basis for $\mathcal{N}(\Phi)$, and define the vectors $\{h_i\}_{i=1}^N$ as follows:

$$h_j = \sum_{i=1}^{N-M} v_i(j) v_i. \quad (4.44)$$

⁹This content is available online at <http://cnx.org/content/m37183/1.6/>.

We note that $h_j = \sum_{i=1}^{N-M} < e_j, v_i > v_i$ where e_j denotes the vector of all zeros except for a 1 in the j -th entry. Thus we see that $h_j = P_{\mathcal{N}} e_j$ where $P_{\mathcal{N}}$ denotes an orthogonal projection onto $\mathcal{N}(\Phi)$. Since $\|P_{\mathcal{N}} e_j\|_2^2 + \|P_{\mathcal{N}}^\perp e_j\|_2^2 = \|e_j\|_2^2 = 1$, we have that $\|h_j\|_2 \leq 1$. Thus, by setting $\Lambda = \{j\}$ for h_j we observe that

$$\left| \sum_{i=1}^{N-M} |v_i(j)|^2 \right|^2 = |h_j(j)|^2 \leq \left(1 - \frac{1}{C^2}\right) \|h_j\|_2^2 \leq 1 - \frac{1}{C^2}. \quad (4.45)$$

Summing over $j = 1, 2, \dots, N$, we obtain

$$N\sqrt{1 - 1/C^2} \geq \sum_{j=1}^N \sum_{i=1}^{N-M} |v_i(j)|^2 = \sum_{i=1}^{N-M} \sum_{j=1}^N |v_i(j)|^2 = \sum_{i=1}^{N-M} \|v_i\|_2^2 = N - M, \quad (4.46)$$

and thus $M \geq \left(1 - \sqrt{1 - 1/C^2}\right) N$ as desired.

Thus, if we want a bound of the form (4.42) that holds for *all* signals x with a constant $C \approx 1$, then regardless of what recovery algorithm we use we will need to take $M \approx N$ measurements. However, in a sense this result is overly pessimistic, and we will now see that the results we just established for signal recovery in noise can actually allow us to overcome this limitation by essentially treating the approximation error as noise.

Towards this end, notice that all the results concerning ℓ_1 minimization stated thus far are deterministic instance-optimal guarantees that apply simultaneously to all x given any matrix that satisfies the restricted isometry property (Section 3.5) (RIP). This is an important theoretical property, but as noted in "Matrices that satisfy the RIP" (Section 3.5), in practice it is very difficult to obtain a deterministic guarantee that the matrix Φ satisfies the RIP. In particular, constructions that rely on randomness are only known to satisfy the RIP with high probability. As an example, recall Theorem 1 from "Matrices that satisfy the RIP" (Section 3.5), which opens the door to slightly weaker results that hold only with high probability.

Theorem 4.5:

Fix $\delta \in (0, 1)$. Let Φ be an $M \times N$ random matrix whose entries ϕ_{ij} are i.i.d. with ϕ_{ij} drawn according to a strictly sub-Gaussian distribution (Section 3.7) with $c^2 = 1/M$. If

$$M \geq \kappa_1 K \log \left(\frac{N}{K} \right), \quad (4.47)$$

then Φ satisfies the RIP of order K with the prescribed δ with probability exceeding $1 - 2e^{-\kappa_2 M}$, where κ_1 is arbitrary and $\kappa_2 = \delta^2/2\kappa^* - \log(42e/\delta)/\kappa_1$.

Even within the class of probabilistic results, there are two distinct flavors. The typical approach is to combine a probabilistic construction of a matrix that will satisfy the RIP with high probability with the previous results in this chapter. This yields a procedure that, with high probability, will satisfy a deterministic guarantee applying to all possible signals x . A weaker kind of result is one that states that given a signal x , we can draw a random matrix Φ and with high probability expect certain performance *for that signal* x . This type of guarantee is sometimes called *instance-optimal in probability*. The distinction is essentially whether or not we need to draw a new random Φ for each signal x . This may be an important distinction in practice, but if we assume for the moment that it is permissible to draw a new matrix Φ for each x , then we can see that Theorem 4.4, (Theorem 5.1 of [35]), p. 51 may be somewhat pessimistic. In order to establish our main result we will rely on the fact, previously used in "Matrices that satisfy the RIP" (Section 3.5), that sub-Gaussian matrices preserve the norm of an arbitrary vector with high probability. Specifically, a slight modification of Corollary 1 from "Matrices that satisfy the RIP" (Section 3.5) shows that for any $x \in \mathbb{R}^N$, if we choose Φ according to the procedure in Theorem 4.5, p. 52, then we also have that

$$\mathbb{P}(\|\Phi x\|_2^2 \geq 2\|x\|_2^2) \leq \exp(-\kappa_3 M) \quad (4.48)$$

with $\kappa_3 = 4/\kappa^*$. Using this we obtain the following result.

Theorem 4.6:

Let $x \in \mathbb{R}^N$ be fixed. Set $\delta_{2K} < \sqrt{2} - 1$. Suppose that Φ is an $M \times N$ sub-Gaussian random matrix with $M \geq \kappa_1 K \log(N/K)$. Suppose we obtain measurements of the form $y = \Phi x$. Set $\varepsilon = 2\sigma_K(x)_2$. Then with probability exceeding $1 - 2\exp(-\kappa_2 M) - \exp(-\kappa_3 M)$, when $\mathcal{B}(y) = \{z : \|\Phi z - y\|_2 \leq \varepsilon\}$, the solution $\hat{x}$ to (4.40) obeys

$$\|\hat{x} - x\|_2 \leq \frac{8\sqrt{1 + \delta_{2K}} - (1 + \sqrt{2})\delta_{2K}}{1 - (1 + \sqrt{2})\delta_{2K}} \sigma_K(x)_2. \quad (4.49)$$

Proof:

First we recall that, as noted above, from Theorem 4.5, p. 52 we have that Φ will satisfy the RIP of order $2K$ with probability at least $1 - 2\exp(-\kappa_2 M)$. Next, let Λ denote the index set corresponding to the K entries of x with largest magnitude and write $x = x_\Lambda + x_{\Lambda^c}$. Since $x_\Lambda \in \Sigma_K$, we can write $\Phi x = \Phi x_\Lambda + \Phi x_{\Lambda^c} = \Phi x_\Lambda + e$. If Φ is sub-Gaussian then from Lemma 2 from "Sub-Gaussian random variables" (Section 3.7) we have that Φx_{Λ^c} is also sub-Gaussian, and one can apply (4.48) to obtain that with probability at least $1 - \exp(-\kappa_3 M)$, $\|\Phi x_{\Lambda^c}\|_2 \leq 2\|x_{\Lambda^c}\|_2 = 2\sigma_K(x)_2$. Thus, applying the union bound we have that with probability exceeding $1 - 2\exp(-\kappa_2 M) - \exp(-\kappa_3 M)$, we satisfy the necessary conditions to apply Theorem 1 from "Signal recovery in noise" (Section 4.3) to x_Λ , in which case $\sigma_K(x_\Lambda)_1 = 0$ and hence

$$\|\hat{x} - x_\Lambda\|_2 \leq 2C_2 \sigma_K(x)_2. \quad (4.50)$$

From the triangle inequality we thus obtain

$$\|\hat{x} - x\|_2 = \|\hat{x} - x_\Lambda + x_\Lambda - x\|_2 \leq \|\hat{x} - x_\Lambda\|_2 + \|x_\Lambda - x\|_2 \leq (2C_2 + 1) \sigma_K(x)_2 \quad (4.51)$$

which establishes the theorem.

Thus, although it is not possible to achieve a deterministic guarantee of the form in (4.42) without taking a prohibitively large number of measurements, it is possible to show that such performance guarantees can hold with high probability while simultaneously taking far fewer measurements than would be suggested by Theorem 4.4, (Theorem 5.1 of [35]), p. 51. Note that the above result applies only to the case where the parameter is selected correctly, which requires some limited knowledge of x , namely $\sigma_K(x)_2$. In practice this limitation can easily be overcome through a parameter selection technique such as cross-validation [188], but there also exist more intricate analyses of ℓ_1 minimization that show it is possible to obtain similar performance without requiring an oracle for parameter selection [190]. Note that Theorem 4.6, p. 53 can also be generalized to handle other measurement matrices and to the case where x is compressible rather than sparse. Moreover, this proof technique is applicable to a variety of the greedy algorithms described later in this course that do not require knowledge of the noise level to establish similar results [34], [137].

4.5 The cross-polytope and phase transitions¹⁰

The analysis of ℓ_1 minimization based on the restricted isometry property (Section 3.3) (RIP) described in "Signal recovery in noise" (Section 4.3) allows us to establish a variety of guarantees under different noise settings, but one drawback is that the analysis of how many measurements are actually required for a matrix to satisfy the RIP is relatively loose. An alternative approach to analyzing ℓ_1 minimization algorithms is to

¹⁰This content is available online at <<http://cnx.org/content/m37184/1.5/>>.

examine them from a more geometric perspective. Towards this end, we define the closed ℓ_1 ball, also known as the *cross-polytope*:

$$C^N = \{x \in \mathbb{R}^N : \|x\|_1 \leq 1\}. \quad (4.52)$$

Note that C^N is the convex hull of $2N$ points $\{p_i\}_{i=1}^{2N}$. Let $\Phi C^N \subseteq \mathbb{R}^M$ denote the convex polytope defined as either the convex hull of $\{\Phi p_i\}_{i=1}^{2N}$ or equivalently as

$$\Phi C^N = \{y \in \mathbb{R}^M : y = \Phi x, x \in C^N\}. \quad (4.53)$$

For any $x \in \Sigma_K = \{x : \|x\|_0 \leq K\}$, we can associate a K -face of C^N with the support and sign pattern of x . One can show that the number of K -faces of ΦC^N is precisely the number of index sets of size K for which signals supported on them can be recovered by

$$\hat{x} = \underset{z}{\operatorname{argmin}} \|z\|_1 \quad \text{subject to} \quad z \in \mathcal{B}(y). \quad (4.54)$$

with $\mathcal{B}(y) = \{z : \Phi z = y\}$. Thus, ℓ_1 minimization yields the same solution as ℓ_0 minimization for all $x \in \Sigma_K$ if and only if the number of K -faces of ΦC^N is identical to the number of K -faces of C^N . Moreover, by counting the number of K -faces of ΦC^N , we can quantify exactly what fraction of sparse vectors can be recovered using ℓ_1 minimization with Φ as our sensing matrix. See [55], [57], [63], [64], [65] for more details and [66] for an overview of the implications of this body of work. Note also that by replacing the cross-polytope with certain other polytopes (the simplex and the hypercube), one can apply the same technique to obtain results concerning the recovery of more limited signal classes, such as sparse signals with nonnegative or bounded entries [66].

Given this result, one can then study random matrix constructions from this perspective to obtain probabilistic bounds on the number of K -faces of ΦC^N with Φ is generated at random, such as from a Gaussian distribution. Under the assumption that $K = \rho M$ and $M = \gamma N$, one can obtain asymptotic results as $N \rightarrow \infty$. This analysis leads to the *phase transition* phenomenon, where for large problem sizes there are sharp thresholds dictating that the fraction of K -faces preserved will tend to either one or zero with high probability, depending on ρ and γ [66].

These results provide sharp bounds on the minimum number of measurements required in the noiseless setting. In general, these bounds are significantly stronger than the corresponding measurement bounds obtained within the RIP-based framework given in "Noise-free signal recovery" (Section 4.2), which tend to be extremely loose in terms of the constants involved. However, these sharper bounds also require somewhat more intricate analysis and typically more restrictive assumptions on Φ (such as it being Gaussian). Thus, one of the main strengths of the RIP-based analysis presented in "Noise-free signal recovery" (Section 4.2) and "Signal recovery in noise" (Section 4.3) is that it gives results for a broad class of matrices that can also be extended to noisy settings.

4.6 Sparse recovery algorithms¹¹

Given noisy compressive measurements $y = \Phi x + e$ of a signal x , a core problem in compressive sensing¹² (CS) is to recover a sparse (Section 2.3) signal x from a set of measurements (Section 3.1) y . Considerable efforts have been directed towards developing algorithms that perform fast, accurate, and stable reconstruction of x from y . As we have seen in previous chapters (Section 3.1), a "good" CS matrix Φ typically satisfies certain geometric conditions, such as the restricted isometry property (Section 3.3) (RIP). Practical algorithms exploit this fact in various ways in order to drive down the number of measurements, enable faster reconstruction, and ensure robustness to both numerical and stochastic errors.

The design of sparse recovery algorithms are guided by various criteria. Some important ones are listed as follows.

¹¹This content is available online at <http://cnx.org/content/m37292/1.3/>.

¹²"Introduction to compressive sensing" <http://cnx.org/content/m37172/latest/>

- **Minimal number of measurements.** Sparse recovery algorithms must require approximately the same number of measurements (up to a small constant) required for the stable embedding of K -sparse signals.
- **Robustness to measurement noise and model mismatch** Sparse recovery algorithms must be stable with regards to perturbations of the input signal, as well as noise added to the measurements; both types of errors arise naturally in practical systems.
- **Speed.** Sparse recovery algorithms must strive towards expending minimal computational resources, Keeping in mind that a lot of applications in CS deal with very high-dimensional signals.
- **Performance guarantees.** In previous chapters (Section 4.1), we have already seen a range of performance guarantees that hold for sparse signal recovery using ℓ_1 minimization. In evaluating other algorithms, we will have the same considerations. For example, we can choose to design algorithms that possess instance-optimal or probabilistic guarantees (Section 4.4). We can also choose to focus on algorithm performance for the recovery of exactly K -sparse signals x , or consider performance for the recovery of general signals xs . Alternately, we can also consider algorithms that are accompanied by performance guarantees in either the noise-free (Section 4.2) or noisy (Section 4.3) settings.

A multitude of algorithms satisfying some (or even all) of the above have been proposed in the literature. While it is impossible to describe all of them in this chapter, we refer the interested reader to the DSP resources webpage¹³ for a more complete list of recovery algorithms. Broadly speaking, recovery methods tend to fall under three categories: convex optimization-based approaches (Section 4.7), greedy methods (Section 4.8), and combinatorial techniques (Section 4.9). The rest of the chapter discusses several properties and example algorithms of each flavor of CS reconstruction.

4.7 Convex optimization-based methods¹⁴

An important class of sparse recovery algorithms (Section 4.6) fall under the purview of *convex optimization*. Algorithms in this category seek to optimize a convex function $f(\cdot)$ of the unknown variable x over a (possibly unbounded) convex subset of $\mathbb{R}^N$.

4.7.1 Setup

Let $J(x)$ be a convex *sparsity-promoting* cost function (i.e., $J(x)$ is small for sparse x .) To recover a sparse signal representation $\hat{x}$ from measurements $y = \Phi x$, $\Phi \in \mathbb{R}^{M \times N}$, we may either solve

$$\min_x \{J(x) : y = \Phi x\}, \quad (4.55)$$

when there is no noise, or solve

$$\min_x \{J(x) : H(\Phi x, y) \leq \varepsilon\} \quad (4.56)$$

when there is noise in the measurements. Here, H is a cost function that penalizes the distance between the vectors Φx and y . For an appropriate penalty parameter μ , (4.56) is equivalent to the *unconstrained* formulation:

$$\min_x J(x) + \mu H(\Phi x, y) \quad (4.57)$$

for some $\mu > 0$. The parameter μ may be chosen by trial-and-error, or by statistical techniques such as cross-validation [16].

For convex programming algorithms, the most common choices of J and H are usually chosen as follows: $J(x) = \|x\|_1$, the ℓ_1 -norm of x , and $H(\Phi x, y) = \frac{1}{2} \|\Phi x - y\|_2^2$, the ℓ_2 -norm of the error between the

¹³<http://dsp.rice.edu/cs>

¹⁴This content is available online at <<http://cnx.org/content/m37293/1.5/>>.

observed measurements and the linear projections of the target vector x . In statistics, minimizing this H subject to $\|x\|_1 \leq \delta$ is known as the *Lasso* problem. More generally, $J(\cdot)$ acts as a regularization term and can be replaced by other, more complex, functions; for example, the desired signal may be piecewise constant, and simultaneously have a sparse representation under a known basis transform Ψ . In this case, we may use a mixed regularization term:

$$J(x) = TV(x) + \lambda \|x\|_1 \quad (4.58)$$

It might be tempting to use conventional convex optimization packages for the above formulations ((4.55), (4.56), and (4.57)). Nevertheless, the above problems pose two key challenges which are specific to practical problems encountered in CS¹⁵: (i) real-world applications are invariably large-scale (an image of a resolution of 1024×1024 pixels leads to optimization over a million variables, well beyond the reach of any standard optimization software package); (ii) the objective function is nonsmooth, and standard smoothing techniques do not yield very good results. Hence, for these problems, conventional algorithms (typically involving matrix factorizations) are not effective or even applicable. These unique challenges encountered in the context of CS have led to considerable interest in developing improved sparse recovery algorithms in the optimization community.

4.7.2 Linear programming

In the noiseless (Section 4.2) case, the ℓ_1 -minimization (Section 4.1) problem (obtained by substituting $J(x) = \|x\|_1$ in (4.55)) can be recast as a linear program (LP) with equality constraints. These can be solved in polynomial time ($O(N^3)$) using standard interior-point methods [17]. This was the first feasible reconstruction algorithm used for CS recovery and has strong theoretical guarantees, as shown earlier in this course (Section 4.1). In the noisy case, the problem can be recast as a second-order cone program (SOCP) with quadratic constraints. Solving LPs and SOCPs is a principal thrust in optimization research; nevertheless, their application in practical CS problems is limited due to the fact that both the signal dimension N , and the number of constraints M , can be very large in many scenarios. Note that both LPs and SOCPs correspond to the constrained formulations in (4.55) and (4.56) and are solved using *first order* interior-point methods.

A newer algorithm called “l1_ls” [110] is based on an interior-point algorithm that uses a preconditioned conjugate gradient (PCG) method to approximately solve linear systems in a truncated-Newton framework. The algorithm exploits the structure of the Hessian to construct their preconditioner; thus, this is a second order method. Computational results show that about a hundred PCG steps are sufficient for obtaining accurate reconstruction. This method has been typically shown to be slower than first-order methods, but could be faster in cases where the true target signal is highly sparse.

4.7.3 Fixed-point continuation

As opposed to solving the constrained formulation, an alternate approach is to solve the unconstrained formulation in (4.57). A widely used method for solving ℓ_1 -minimization problems of the form

$$\min_x \mu \|x\|_1 + H(x), \quad (4.59)$$

for a convex and differentiable H , is an iterative procedure based on *shrinkage* (also called soft thresholding; see (4.60) below). In the context of solving (4.59) with a quadratic H , this method was independently proposed and analyzed in [11], [83], [133], [141], and then further studied or extended in [38], [44], [74], [76], [96], [191]. Shrinkage is a classic method used in wavelet-based image denoising. The shrinkage operator on

¹⁵“Introduction to compressive sensing” <<http://cnx.org/content/m37172/latest/>>

any scalar component can be defined as follows:

$$\text{shrink}(t, \alpha) = \begin{cases} t - \alpha & \text{if } t > \alpha, \\ 0 & \text{if } -\alpha \leq t \leq \alpha, \text{ and} \\ t + \alpha & \text{if } t < -\alpha. \end{cases} \quad (4.60)$$

This concept can be used effectively to solve (4.59). In particular, the basic algorithm can be written as following the fixed-point iteration: for $i = 1, \dots, N$, the i^{th} coefficient of x at the $(k+1)^{\text{th}}$ time step is given by

$$x_i^{k+1} = \text{shrink}((x^k - \tau [\text{U+25BD}] H(x^k))_i, \mu\tau) \quad (4.61)$$

where $\tau > 0$ serves as a step-length for gradient descent (which may vary with k) and μ is as specified by the user. It is easy to see that the larger μ is, the larger the allowable distance between x^{k+1} and x^k . For a quadratic penalty term $H(\cdot)$, the gradient $[\text{U+25BD}] H$ can be easily computed as a linear function of x^k ; thus each iteration of (4.61) essentially boils down to a small number of matrix-vector multiplications.

The simplicity of the iterative approach is quite appealing, both from a computational, as well as a code-design standpoint. Various modifications, enhancements, and generalizations to this approach have been proposed, both to improve the efficiency of the basic iteration in (4.61), and to extend its applicability to various kinds of J [77], [84], [191]. In principle, the basic iteration in (4.61) would not be practically effective without a continuation (or path-following) strategy [96], [191] in which we choose a gradually decreasing sequence of values for the parameter μ to guide the intermediate iterates towards the final optimal solution.

This procedure is known as *continuation*; in [96], the performance of an algorithm known as Fixed-Point Continuation (FPC) has been compared favorably with another similar method known as Gradient Projection for Sparse Reconstruction (GPSR) [84] and “l1_ls” [110]. A key aspect to solving the unconstrained optimization problem is the choice of the parameter μ . As discussed above, for CS recovery, μ may be chosen by trial and error; for the noiseless constrained formulation, we may solve the corresponding unconstrained minimization by choosing a large value for μ .

In the case of recovery from noisy compressive measurements, a commonly used choice for the convex cost function $H(x)$ is the square of the norm of the *residual*. Thus we have:

$$\begin{aligned} H(x) &= \|y - \Phi x\|_2^2 \\ [\text{U+25BD}] H(x) &= 2\Phi^\top (y - \Phi x). \end{aligned} \quad (4.62)$$

For this particular choice of penalty function, (4.61) reduces to the following iteration:

$$x_i^{k+1} = \text{shrink}((x^k - \tau [\text{U+25BD}] H(y - \Phi x^k))_i, \mu\tau) \quad (4.63)$$

which is run until convergence to a fixed point. The algorithm is detailed in pseudocode form below.

```

Inputs: CS matrix  $\Phi$ , signal measurements  $y$ , parameter sequence  $\mu\_n$ 
Outputs: Signal estimate  $\hat{x}$ 
initialize:  $\hat{x}_0 = 0$ ,  $r = y$ ,  $k = 0$ .
while altting criterion false do
  1.  $k \leftarrow k + 1$ 
  2.  $\hat{x} \leftarrow \hat{x} - \tau \Phi^\top r$  {take a gradient step}
  3.  $\hat{x} \leftarrow \text{shrink}(\hat{x}, \mu\_k \tau)$  {perform soft thresholding}
  4.  $r \leftarrow y - \Phi \hat{x}$  {update measurement residual}

```

```

end while
return  $\hat{x} \leftarrow \hat{x}$ 

```

4.7.4 Bregman iteration methods

It turns out that an efficient method to obtain the solution to the constrained optimization problem in (4.55) can be devised by solving a small number of the unconstrained problems in the form of (4.57). These subproblems are commonly referred to as *Bregman iterations*. A simple version can be written as follows:

$$\begin{aligned} y^{k+1} &= y^k + y - \Phi x^k \\ x^{k+1} &= \operatorname{argmin} J(x) + \frac{\mu}{2} \|\Phi x - y^{k+1}\|^2. \end{aligned} \quad (4.64)$$

The problem in the second step can be solved by the algorithms reviewed above. Bregman iterations were introduced in [143] for constrained total variation minimization problems, and was proved to converge for closed, convex functions $J(x)$. In [192], it is applied to (4.55) for $J(x) = \|x\|_1$ and shown to converge in a finite number of steps for any $\mu > 0$. For moderate μ , the number of iterations needed is typically lesser than 5. Compared to the alternate approach that solves (4.55) through directly solving the unconstrained problem in (4.57) with a very large μ , Bregman iterations are often more stable and sometimes much faster.

4.7.5 Discussion

All the methods discussed in this section optimize a convex function (usually the ℓ_1 -norm) over a convex (possibly unbounded) set. This implies *guaranteed* convergence to the global optimum. In other words, given that the sampling matrix Φ satisfies the conditions specified in "Signal recovery via ℓ_1 minimization" (Section 4.1), convex optimization methods will recover the underlying signal x . In addition, convex relaxation methods also guarantee *stable* recovery by reformulating the recovery problem as the SOCP, or the unconstrained formulation.

4.8 Greedy algorithms¹⁶

4.8.1 Setup

As opposed to solving a (possibly computationally expensive) convex optimization (Section 4.7) program, an alternate flavor to sparse recovery (Section 4.6) is to apply methods of *sparse approximation*. Recall that the goal of sparse recovery is to recover the *sparsest* vector x which explains the linear measurements y . In other words, we aim to solve the (nonconvex) problem:

$$\min_{\mathcal{I}} \{|\mathcal{I}| : y = \sum_{i \in \mathcal{I}} \phi_i x_i\}, \quad (4.65)$$

where $\mathcal{I}$ denotes a particular subset of the indices $i = 1, \dots, N$, and ϕ_i denotes the i^{th} column of Φ . It is well known that searching over the power set formed by the columns of Φ for the optimal subset $\mathcal{I}^*$ with smallest cardinality is NP-hard. Instead, classical sparse approximation methods tackle this problem by *greedily* selecting columns of Φ and forming successively better approximations to y .

4.8.2 Matching Pursuit

Matching Pursuit (MP), named and introduced to the signal processing community by Mallat and Zhang [126], [127], is an iterative greedy algorithm that decomposes a signal into a linear combination

¹⁶This content is available online at <<http://cnx.org/content/m37294/1.4/>>.

of elements from a dictionary. In sparse recovery, this dictionary is merely the sampling matrix $\Phi \in \mathbb{R}^{M \times N}$; we seek a sparse representation (x) of our “signal” y .

MP is conceptually very simple. A key quantity in MP is the *residual* $r \in \mathbb{R}^M$; the residual represents the as-yet “unexplained” portion of the measurements. At each iteration of the algorithm, we select a vector from the dictionary that is maximally correlated with the residual r :

$$\lambda_k = \underset{\lambda}{\operatorname{argmax}} \frac{\langle r_k, \phi_\lambda \rangle \phi_\lambda}{\|\phi_\lambda\|^2}. \quad (4.66)$$

Once this column is selected, we possess a “better” representation of the signal, since a new coefficient indexed by λ_k has been added to our signal approximation. Thus, we update both the residual and the approximation as follows:

$$\begin{aligned} r_k &= r_{k-1} - \frac{\langle r_{k-1}, \phi_{\lambda_k} \rangle \phi_{\lambda_k}}{\|\phi_{\lambda_k}\|^2}, \\ \hat{x}_{\lambda_k} &= \hat{x}_{\lambda_k} + \langle r_{k-1}, \phi_{\lambda_k} \rangle. \end{aligned} \quad (4.67)$$

and repeat the iteration. A suitable stopping criterion is when the norm of r becomes smaller than some quantity. MP is described in pseudocode form below.

```

Inputs: Measurement matrix  $\Phi$  , signal measurements  $y$ 
Outputs: Sparse signal  $\hat{x}$ 
initialize:  $\hat{x}_0 = 0$ ,  $r = y$ ,  $i = 0$ .
while  $\text{altting criterion false}$  do
    1.  $i \leftarrow i + 1$ 
    2.  $b \leftarrow \Phi^T r$  {form residual signal estimate}
    3.  $\hat{x}_i \leftarrow \hat{x}_{i-1} + \mathbf{T}(b)$  {update largest magnitude coefficient}
    4.  $r \leftarrow r - \Phi \hat{x}_i$  {update measurement residual}
end while
return  $\hat{x} \leftarrow \hat{x}_i$ 

```

Although MP is intuitive and can find an accurate approximation of the signal, it possesses two major drawbacks: (i) it offers no guarantees in terms of recovery error; indeed, it does not exploit the special structure present in the dictionary Φ ; (ii) the required number of iterations required can be quite large. The complexity of MP is $O(MNT)$ [72] , where T is the number of MP iterations

4.8.3 Orthogonal Matching Pursuit (OMP)

Matching Pursuit (MP) can prove to be computationally infeasible for many problems, since the complexity of MP grows linearly in the number of iterations T . By employing a simple modification of MP, the maximum number of MP iterations can be upper bounded as follows. At any iteration k , Instead of subtracting the contribution of the dictionary element with which the residual r is maximally correlated, we compute the projection of r onto the *orthogonal subspace* to the linear span of the currently selected dictionary elements. This quantity thus better represents the “unexplained” portion of the residual, and is subtracted from r to form a new residual, and the process is repeated. If Φ_Ω is the submatrix formed by the columns of Φ selected at time step t , the following operations are performed:

$$\begin{aligned} x_k &= \underset{x}{\operatorname{argmin}} \|y - \Phi_\Omega x\|_2, \\ \hat{\alpha}_t &= \Phi_\Omega^T r_t, \\ r_t &= y - \Phi_\Omega \hat{\alpha}_t. \end{aligned} \quad (4.68)$$

These steps are repeated until convergence. This is known as Orthogonal Matching Pursuit (OMP) [144]. Tropp and Gilbert [172] proved that OMP can be used to recover a sparse signal with high probability using compressive measurements. The algorithm converges in at most K iterations, where K is the sparsity, but requires the added computational cost of orthogonalization at each iteration. Indeed, the total complexity of OMP can be shown to be $O(MNK)$.

While OMP is provably fast and can be shown to lead to exact recovery, the guarantees accompanying OMP for sparse recovery are weaker than those associated with optimization techniques (Section 4.1). In particular, the reconstruction guarantees are *not uniform*, i.e., it cannot be shown that a single measurement matrix with $M = CK \log N$ rows can be used to recover every possible K -sparse signal with $M = CK \log N$ measurements. (Although it is possible to obtain such uniform guarantees when it is acceptable to take more measurements. For example, see [49].) Another issue with OMP is robustness to noise; it is unknown whether the solution obtained by OMP will only be perturbed slightly by the addition of a small amount of noise in the measurements. Nevertheless, OMP is an efficient method for CS recovery, especially when the signal sparsity K is low. A pseudocode representation of OMP is shown below.

```

Inputs: Measurement matrix  $\Phi$ , signal measurements  $y$ 
Outputs: Sparse representation  $\hat{x}$ 
Initialize:  $\hat{\theta}_0 = 0$ ,  $r = y$ ,  $\Omega = \emptyset$ ,  $i = 0$ .
while altting criterion false do
    1.  $i \leftarrow i + 1$ 
    2.  $b \leftarrow \Phi^T r$  {form residual signal estimate}
    3.  $\Omega \leftarrow \Omega \cup \text{supp}(\mathbf{T}(b, 1))$  {add index of residual's largest magnitude entry to signal support}
    4.  $\hat{x}_{i|\Omega} \leftarrow \Phi_{\Omega}^+ r$ ,  $\hat{x}_{i|\Omega}^C \leftarrow 0$  {form signal estimate}
    5.  $r \leftarrow y - \Phi \hat{x}_i$  {update measurement residual}
end while
return  $\hat{x} \leftarrow \hat{x}_i$ 

```

4.8.4 Stagewise Orthogonal Matching Pursuit (StOMP)

Orthogonal Matching Pursuit is ineffective when the signal is not very sparse as the computational cost increases quadratically with the number of nonzeros K . In this setting, Stagewise Orthogonal Matching Pursuit (StOMP) [58] is a better choice for approximately sparse signals in a large-scale setting.

StOMP offers considerable computational advantages over ℓ_1 minimization (Section 4.7) and Orthogonal Matching Pursuit for large scale problems with sparse solutions. The algorithm starts with an initial residual $r_0 = y$ and calculates the set of all projections $\Phi^T r_{k-1}$ at the k^{th} stage (as in OMP). However, instead of picking a single dictionary element, it uses a threshold parameter τ to determine the next best *set of columns* of Φ whose correlations with the current residual exceed τ . The new residual is calculated using a least squares estimate of the signal using this expanded set of columns, just as before.

Unlike OMP, the number of iterations in StOMP is fixed and chosen before hand; $S = 10$ is recommended in [58]. In general, the complexity of StOMP is $O(KN \log N)$, a significant improvement over OMP. However, StOMP does not bring in its wake any reconstruction guarantees. StOMP also has moderate memory requirements compared to OMP where the orthogonalization requires the maintenance of a Cholesky factorization of the dictionary elements.

4.8.5 Compressive Sampling Matching Pursuit (CoSaMP)

Greedy pursuit algorithms (such as MP and OMP) alleviate the issue of computational complexity encountered in optimization-based sparse recovery, but lose the associated strong guarantees for uniform signal recovery, given a requisite number of measurements of the signal. In addition, it is unknown whether these greedy algorithms are robust to signal and/or measurement noise.

There have been some recent attempts to develop greedy algorithms (Regularized OMP [139], [140], Compressive Sampling Matching Pursuit (CoSaMP) [138] and Subspace Pursuit [42]) that bridge this gap between uniformity and complexity. Intriguingly, the restricted isometry property (Section 3.3) (RIP), developed in the context of analyzing ℓ_1 minimization (Section 4.1), plays a central role in such algorithms. Indeed, if the matrix Φ satisfies the RIP of order K , this implies that every subset of K columns of the matrix is approximately orthonormal. This property is used to prove strong convergence results of these greedy-like methods.

One variant of such an approach is employed by the CoSaMP algorithm. An interesting feature of CoSaMP is that unlike MP, OMP and StOMP, new indices in a signal estimate can be added *as well as deleted* from the current set of chosen indices. In contrast, greedy pursuit algorithms suffer from the fact that a chosen index (or equivalently, a chosen atom from the dictionary Φ remains in the signal representation until the end. A brief description of CoSaMP is as follows: at the start of a given iteration i , suppose the signal estimate is $\hat{x}_{i-1}$.

- Form signal residual estimate: $e \leftarrow \Phi^T r$
- Find the biggest $2K$ coefficients of the signal residual e ; call this set of indices Ω .
- Merge supports: $T \leftarrow \Omega \cup \text{supp}(\hat{x}_{i-1})$.
- Form signal estimate b by subspace projection: $b|_T \leftarrow \Phi_T^\dagger y$, $b|_{T^c} \leftarrow 0$.
- Prune b by retaining its K largest coefficients. Call this new estimate $\hat{x}_i$.
- Update measurement residual: $r \leftarrow y - \Phi \hat{x}_i$.

This procedure is summarized in pseudocode form below.

```

Inputs: Measurement matrix  $\Phi$  , measurements  $y$ , signal sparsity  $K$ 
Output:  $K$ -sparse approximation  $\hat{x}$  to true signal representation  $x$ 
Initialize:  $\hat{x}_0 = 0$  ,  $r = y$ ;  $i = 0$ 
while alting criterion false do
  1.
    i ← i + 1
  2.  $e \leftarrow \Phi^T r$  {form signal residual estimate}
  3.  $\Omega \leftarrow \text{supp}(\mathbf{T}(e, 2K))$  {prune signal residual estimate}
  4.  $T \leftarrow \Omega \cup \text{supp}(\hat{x}_{i-1})$  {merge supports}
  5.  $b|_T \leftarrow \Phi_T^\dagger y$ ,  $b|_{T^c} \leftarrow 0$  {form signal estimate}
  6.  $\hat{x}_i \leftarrow \mathbf{T}(b, K)$  {prune signal estimate}
  7.  $r \leftarrow y - \Phi \hat{x}_i$  {update measurement residual}
end while
return  $\hat{x} \leftarrow \hat{x}_i$ 

```

As discussed in [138], the key computational issues for CoSaMP are the formation of the signal residual, and the method used for subspace projection in the signal estimation step. Under certain general assumptions, the computational cost of CoSaMP can be shown to be $O(MN)$, which is *independent* of the sparsity of the original signal. This represents an improvement over both greedy algorithms as well as convex methods.

While CoSaMP arguably represents the state of the art in sparse recovery algorithm performance, it possesses one drawback: the algorithm requires prior knowledge of the sparsity K of the target signal. An incorrect choice of input sparsity may lead to a worse guarantee than the actual error incurred by a weaker algorithm such as OMP. The stability bounds accompanying CoSaMP ensure that the error due to an incorrect parameter choice is bounded, but it is not yet known how these bounds translate into practice.

4.8.6 Iterative Hard Thresholding

Iterative Hard Thresholding (IHT) is a well-known algorithm for solving nonlinear inverse problems. The structure of IHT is simple: starting with an initial estimate $\hat{x}_0$, iterative hard thresholding (IHT) obtains a sequence of estimates using the iteration:

$$\hat{x}_{i+1} = \mathbf{T} \left(\hat{x}_i + \Phi^T \left(y - \Phi \hat{x}_i \right), K \right). \quad (4.69)$$

In [15], Blumensath and Davies proved that this sequence of iterations converges to a fixed point $\hat{x}$; further, if the matrix Φ possesses the RIP, they showed that the recovered signal $\hat{x}$ satisfies an instance-optimality guarantee of the type described earlier (Section 4.1). The guarantees (as well as the proof technique) are reminiscent of the ones that are derived in the development of other algorithms such as ROMP and CoSaMP.

4.8.7 Discussion

While convex optimization techniques are powerful methods for computing sparse representations, there are also a variety of greedy/iterative methods for solving such problems. Greedy algorithms rely on iterative approximation of the signal coefficients and support, either by iteratively identifying the support of the signal until a convergence criterion is met, or alternatively by obtaining an improved estimate of the sparse signal at each iteration by accounting for the mismatch to the measured data. Some greedy methods can actually be shown to have performance guarantees that match those obtained for convex optimization approaches. In fact, some of the more sophisticated greedy algorithms are remarkably similar to those used for ℓ_1 minimization described previously (Section 4.7). However, the techniques required to prove performance guarantees are substantially different. There also exist iterative techniques for sparse recovery based on message passing schemes for sparse graphical models. In fact, some greedy algorithms (such as those in [13], [102]) can be directly interpreted as message passing methods [62].

4.9 Combinatorial algorithms¹⁷

In addition to convex optimization (Section 4.7) and greedy pursuit (Section 4.8) approaches, there is another important class of sparse recovery algorithms that we will refer to as *combinatorial algorithms*. These algorithms, mostly developed by the theoretical computer science community, in many cases pre-date the compressive sensing¹⁸ literature but are highly relevant to the sparse signal recovery problem (Section 4.6).

4.9.1 Setup

The oldest combinatorial algorithms were developed in the context of *group testing* [78], [107], [158]. In the group testing problem, we suppose that there are N total items, of which an unknown subset of K elements are anomalous and need to be identified. For example, we might wish to identify defective products in an industrial setting, or identify a subset of diseased tissue samples in a medical context. In both of these cases the vector x indicates which elements are anomalous, i.e., $x_i \neq 0$ for the K anomalous elements and $x_i = 0$ otherwise. Our goal is to design a collection of tests that allow us to identify the support (and possibly the values of the nonzeros) of x while also minimizing the number of tests performed. In the simplest practical setting these tests are represented by a binary matrix Φ whose entries ϕ_{ij} are equal to 1 if and only if the j^{th} item is used in the i^{th} test. If the output of the test is linear with respect to the inputs, then the problem of recovering the vector x is essentially the same as the standard sparse recovery problem.

Another application area in which combinatorial algorithms have proven useful is computation on *data streams* [39], [134]. Suppose that x_i represents the number of packets passing through a network router

¹⁷This content is available online at <<http://cnx.org/content/m37295/1.3/>>.

¹⁸"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

with destination i . Simply storing the vector x is typically infeasible since the total number of possible destinations (represented by a 32-bit IP address) is $N = 2^{32}$. Thus, instead of attempting to store x directly, one can store $y = \Phi x$ where Φ is an $M \times N$ matrix with $M \ll N$. In this context the vector y is often called a *sketch*. Note that in this problem y is computed in a different manner than in the compressive sensing context. Specifically, in the network traffic example we do not ever observe x_i directly; rather, we observe increments to x_i (when a packet with destination i passes through the router). Thus we construct y iteratively by adding the i^{th} column to y each time we observe an increment to x_i , which we can do since $y = \Phi x$ is linear. When the network traffic is dominated by traffic to a small number of destinations, the vector x is compressible, and thus the problem of recovering x from the sketch Φx is again essentially the same as the sparse recovery problem.

Several combinatorial algorithms for sparse recovery have been developed in the literature. A non-exhaustive list includes Random Fourier Sampling [93], HHS Pursuit [93], and Sparse Sequential Matching Pursuit [12]. We do not provide a full discussion of each of these algorithms; instead, we describe two simple methods that highlight the flavors of combinatorial sparse recovery — *count-min* and *count-median*.

4.9.2 The count-min sketch

Define H as the set of all discrete-valued functions $h : \{1, \dots, N\} \rightarrow \{1, \dots, m\}$. Note that H is a finite set of size m^N . Each function $h \in H$ can be specified by a binary *characteristic matrix* $\phi(h)$ of size $m \times N$, with each column being a binary vector with exactly one 1 at the location j , where $j = h(i)$. To construct the overall *sampling matrix* Φ , we choose d functions $h_1, \dots, h_d$ independently from the uniform distribution defined on H , and vertically concatenate their characteristic matrices. Thus, if $M = md$, Φ is a binary matrix of size $M \times N$ with each column containing exactly d ones.

Now given any signal x , we acquire linear measurements $y = \Phi x$. It is easy to visualize the measurements via the following two properties. First, the coefficients of the measurement vector y are naturally grouped according to the “mother” binary functions $\{h_1, \dots, h_d\}$. Second, consider the i^{th} coefficient of the measurement vector y , which corresponds to the mother binary function h . Then, the expression for y_i is simply given by:

$$y_i = \sum_{j:h(j)=i} x_j. \quad (4.70)$$

In other words, for a fixed signal coefficient index j , each measurement y_i as expressed above consists of an observation of x_j corrupted by other signal coefficients mapped to the same i by the function h . Signal recovery essentially consists of estimating the signal values from these “corrupted” observations.

The *count-min* algorithm is useful in the special case where the entries of the original signal are positive. Given measurements y using the sampling matrix Φ as constructed above, the estimate of the j^{th} signal entry is given by:

$$\hat{x}_j = \min_l y_i : h_l(j) = i. \quad (4.71)$$

Intuitively, this means that the estimate of x_j is formed by simply looking at all measurements that comprise of x_j corrupted by other signal values, and picking the one with the lowest magnitude. Despite the simplicity of this algorithm, it is accompanied by an arguably powerful instance-optimality guarantee: if $d = C \log N$ and $m = 4/\alpha K$, then with high probability, the recovered signal $\hat{x}$ satisfies:

$$\|x - \hat{x}\|_{\infty} \leq \alpha/K \cdot \|x - x^*\|_1, \quad (4.72)$$

where x^* represents the best K -term approximation of x in the ℓ_1 sense.

4.9.3 The count-median sketch

For the general setting when the coefficients of the original signal could be either positive or negative, a similar algorithm known as *count-median* can be used. Instead of picking the minimum of the measurements, we compute the median of all those measurements that are comprised of a corrupted version of x_j and declare it as the signal coefficient estimate, i.e.,

$$\hat{x}_j = \underset{i}{\text{median}} \ y_i : h_i(j) = i. \quad (4.73)$$

The recovery guarantees for count-median are similar to that for count-min, with a different value of the failure probability constant. An important feature of both count-min and count-median is that they require that the measurements be *perfectly noiseless*, in contrast to optimization/greedy algorithms which can tolerate small amounts of measurement noise.

4.9.4 Summary

Although we ultimately wish to recover a sparse signal from a small number of linear measurements in both of these settings, there are some important differences between such settings and the compressive sensing setting studied in this course¹⁹. First, in these settings it is natural to assume that the designer of the reconstruction algorithm also has full control over Φ , and is thus free to choose Φ in a manner that reduces the amount of computation required to perform recovery. For example, it is often useful to design Φ so that it has few nonzeros, i.e., the sensing matrix itself is also sparse [10], [89], [103]. In general, most methods involve careful construction of the sensing matrix (Section 3.1) Φ , which is in contrast with the optimization and greedy methods that work with any matrix satisfying a generic condition such as the restricted isometry property (Section 3.3). This additional degree of freedom can lead to significantly faster algorithms [32], [41], [90], [92].

Second, note that the computational complexity of all the convex methods and greedy algorithms described above is always at least linear in N , since in order to recover x we must at least incur the computational cost of reading out all N entries of x . This may be acceptable in many typical compressive sensing applications, but this becomes impractical when N is extremely large, as in the network monitoring example. In this context, one may seek to develop algorithms whose complexity is linear only in the *length of the representation* of the signal, i.e., its sparsity K . In this case the algorithm does not return a complete reconstruction of x but instead returns only its K largest elements (and their indices). As surprising as it may seem, such algorithms are indeed possible. See [90], [92] for examples.

4.10 Bayesian methods²⁰

4.10.1 Setup

Throughout this course²¹, we have almost exclusively worked within a deterministic signal framework. In other words, our signal x is fixed and belongs to a known set of signals. In this section, we depart from this framework and assume that the sparse (Section 2.3) (or compressible (Section 2.4)) signal of interest arises from a known *probability distribution*, i.e., we assume sparsity promoting *priors* on the elements of x , and recover from the stochastic measurements $y = \Phi x$ a probability distribution on each nonzero element of x . Such an approach falls under the purview of *Bayesian* methods for sparse recovery (Section 4.6).

The algorithms discussed in this section demonstrate a digression from the conventional sparse recovery techniques typically used in compressive sensing²² (CS). We note that none of these algorithms are accompanied by guarantees on the number of measurements required, or the fidelity of signal reconstruction; indeed,

¹⁹An Introduction to Compressive Sensing <<http://cnx.org/content/col11133/latest/>>

²⁰This content is available online at <<http://cnx.org/content/m37359/1.4/>>.

²¹An Introduction to Compressive Sensing <<http://cnx.org/content/col11133/latest/>>

²²"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

in a Bayesian signal modeling framework, there is no well-defined notion of “reconstruction error”. However, such methods do provide insight into developing recovery algorithms for rich classes of signals, and may be of considerable practical interest.

4.10.2 Sparse recovery via belief propagation

As we will see later in this course, there are significant parallels to be drawn between error correcting codes and sparse recovery [152]. In particular, sparse codes such as LDPC codes have had grand success. The advantage that sparse coding matrices may have in efficient encoding of signals and their low complexity decoding algorithms, is transferable to CS encoding and decoding with the use of *sparse* sensing matrices (Section 3.1) Φ . The sparsity in the Φ matrix is equivalent to the sparsity in LDPC coding graphs.

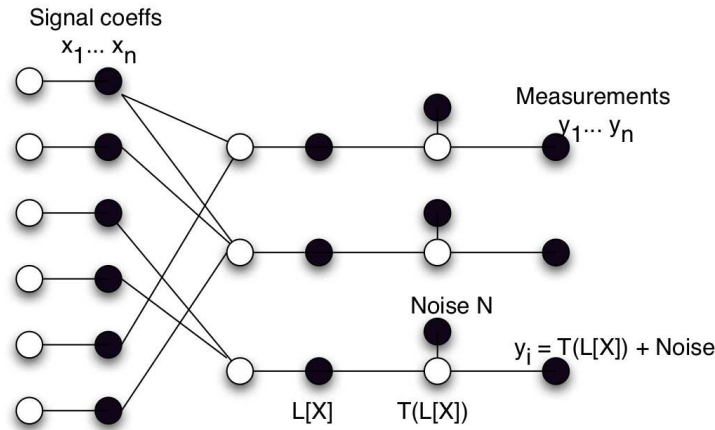


Figure 4.2: Factor graph depicting the relationship between the variables involved in CS decoding using BP. Variable nodes are black and the constraint nodes are white.

A sensing matrix Φ that defines the relation between the signal x and measurements y can be represented as a bipartite graph of signal coefficient nodes $x(i)$ and measurement nodes $y(i)$ [152], [156]. The factor graph in Figure 4.2 represents the relationship between the signal coefficients and measurements in the CS decoding problem.

The choice of signal probability density is of practical interest. In many applications, the signals of interest need to be modeled as being compressible (as opposed to being strictly sparse). This behavior is modeled by a two-state Gaussian mixture distribution, with each signal coefficient taking either a “large” or “small” coefficient value state. Assuming that the elements of x are i.i.d., it can be shown that small coefficients occur more frequently than the large coefficients. Other distributions besides the two-state Gaussian may also be used to model the coefficients, for e.g., the i.i.d. Laplace prior on the coefficients of x .

The ultimate goal is to estimate (i.e., decode) x , given y and Φ . The decoding problem takes the form of a Bayesian inference problem in which we want to approximate the marginal distributions of each of the $x(i)$ coefficients conditioned on the observed measurements $y(i)$. We can then estimate the Maximum Likelihood Estimate (MLE), or the Maximum a Posteriori (MAP) estimates of the coefficients from their distributions. This sort of inference can be solved using a variety of methods; for example, the popular belief propagation method (BP) [152] can be applied to solve for the coefficients approximately. Although exact inference in

arbitrary graphical models is an NP hard problem, inference using BP can be employed when Φ is sparse enough, i.e., when most of the entries in the matrix are equal to zero.

4.10.3 Sparse Bayesian learning

Another probabilistic approach used to estimate the components of x is by using Relevance Vector Machines (RVMs). An RVM is essentially a Bayesian learning method that produces sparse classification by linearly weighting a small number of fixed basis functions from a large dictionary of potential candidates (for more details the interested reader may refer to [170], [169]). From the CS perspective, we may view this as a method to determine the elements of a sparse x which linearly weight the basis functions comprising the columns of Φ .

The RVM setup employs a hierarchy of priors; first, a Gaussian prior is assigned to each of the N elements of x ; subsequently, a Gamma prior is assigned to the inverse-variance α_i of the i^{th} Gaussian prior. Therefore each α_i controls the strength of the prior on its associated weight in x_i . If x is the sparse vector to be reconstructed, its associated Gaussian prior is given by:

$$p(x|\alpha) = \prod_{i=1}^N \mathcal{N}(x_i|0, \alpha_i^{-1}) \quad (4.74)$$

and the Gamma prior on α is written as:

$$p(\alpha|a, b) = \prod_{i=1}^N \Gamma(\alpha_i|a, b) \quad (4.75)$$

The overall prior on x can be analytically evaluated to be the Student-t distribution, which can be designed to peak at $x_i = 0$ with appropriate choice of a and b . This enables the desired solution x to be sparse. The RVM approach can be visualized using a graphical model similar to the one in "Sparse recovery via belief propagation" (Section 4.10.2: Sparse recovery via belief propagation). Using the observed measurements y , the posterior density on each x_i is estimated by an iterative algorithm (e.g., Markov Chain Monte Carlo (MCMC) methods). For a detailed analysis of the RVM with a measurement noise prior, refer to [105], [169].

Alternatively, we can eliminate the need to set the hyperparameters a and b as follows. Assuming Gaussian measurement noise with mean 0 and variance σ^2 , we can directly find the marginal log likelihood for α and maximize it by the EM algorithm (or directly differentiate) to find estimates for α .

$$\mathcal{L}(\alpha) = \log p(y|\alpha, \sigma^2) = \log \int p(y|x, \sigma^2) p(y|\alpha) dx. \quad (4.76)$$

4.10.4 Bayesian compressive sensing

Unfortunately, evaluation of the log-likelihood in the original RVM setup involves taking the inverse of an $N \times N$ matrix, rendering the algorithm's complexity to be $O(N^3)$. A fast alternative algorithm for the RVM is available which monotonically maximizes the marginal likelihoods of the priors by a gradient ascent, resulting in an algorithm with complexity $O(NM^2)$. Here, basis functions are sequentially added and deleted, thus building the model up constructively, and the true sparsity of the signal x is exploited to minimize model complexity. This is known as Fast Marginal Likelihood Maximization, and is employed by the Bayesian Compressive Sensing (BCS) algorithm [105] to efficiently evaluate the posterior densities of x_i .

A key advantage of the BCS algorithm is that it enables evaluation of "error bars" on each estimated coefficient of x ; these give us an idea of the (in)accuracies of these estimates. These error bars could be used to *adaptively select* the linear projections (i.e., the rows of the matrix Φ) to reduce uncertainty in the signal. This provides an intriguing connection between CS and machine learning techniques such as experimental design and active learning [80], [123].

Chapter 5

Applications of Compressive Sensing

5.1 Linear regression and model selection¹

Many of the sparse recovery algorithms (Section 4.6) we have described so far in this course² were originally developed to address the problem of sparse linear regression and model selection in statistics. In this setting we are given some data consisting of a set of input variables and response variables. We will suppose that there are a total of N input variables, and we observe a total of M input and response pairs. We can represent the set of input variable observations as an $M \times N$ matrix Φ , and the set of response variable observations as an $M \times 1$ vector y .

In linear regression, it is assumed that y can be approximated as a linear function of the input variables, i.e., there exists an x such that $y \approx \Phi x$. However, when the number of input variables is large compared to the number of observations, i.e., $M \ll N$, this becomes extremely challenging because we wish to estimate N parameters from far fewer than N observations. In general this would be impossible to overcome, but in practice it is common that only a few input variables are actually necessary to predict the response variable. In this case the x that we wish to estimate is sparse, and we can apply all of the techniques that we have learned so far for sparse recovery to estimate x . In this setting, not only does sparsity aid us in our goal of obtaining a regression, but it also performs *model selection* by identifying the most relevant variables in predicting the response.

5.2 Sparse error correction³

In communications, error correction refers to mechanisms that can detect and correct errors in the data that appear due to distortion in the transmission channel. Standard approaches for error correction rely on repetition schemes, redundancy checks, or nearest neighbor code search. We consider the particular case in which a signal x with M entries is coded by taking length- N linearly independent codewords $\{\phi_1, \dots, \phi_M\}$, with $N > M$ and summing them using the entries of x as coefficients. The received message is a length- N code $y = \sum_{m=1}^M \phi_m x_m = \Phi f$, where Φ is a matrix that has the different codewords for columns. We assume that the transmission channel corrupts the entries of y in an additive way, so that the received data is $y = \Phi x + e$, where e is an error vector.

The techniques developed for sparse recovery (Section 4.6) in the context of compressive sensing⁴ (CS) provide a number of methods to estimate the error vector e — therefore making it possible to correct it and obtain the signal x — when e is sufficiently sparse (Section 2.3) [28]. To estimate the error, we build a matrix Θ that is a basis for the orthogonal subspace to the span of the matrix Φ , i.e., an $(N - M) \times N$ matrix Θ

¹This content is available online at <http://cnx.org/content/m37360/1.3/>.

²*An Introduction to Compressive Sensing* <http://cnx.org/content/col11133/latest/>

³This content is available online at <http://cnx.org/content/m37361/1.3/>.

⁴"Introduction to compressive sensing" <http://cnx.org/content/m37172/latest/>

that holds $\Theta\Phi = 0$. When such a matrix is obtained, we can modify the measurements by multiplying them with the matrix to obtain $\tilde{y} = \Theta y = \Theta\Phi x + \Theta e = \Theta e$. If the matrix Θ is well-suited for CS (i.e., it satisfies a condition such as the restricted isometry property (Section 3.3)) and e is sufficiently sparse, then the error vector e can be estimated accurately using CS. Once the estimate $\hat{e}$ is obtained, the error-free measurements can be estimated as $\hat{y} = y - \hat{e}$, and the signal can be recovered as $\hat{x} = \Phi^\dagger \hat{y} = \Phi^\dagger y - \Phi^\dagger \hat{e}$. As an example, when the codewords ϕ_m have random independent and identically distributed sub-Gaussian (Section 3.7) entries, then a K -sparse error can be corrected if $M < N - CK \log N / K$ for a fixed constant C (see "Matrices that satisfy the RIP" (Section 3.5)).

5.3 Group testing and data stream algorithms⁵

Another scenario where compressive sensing⁶ and sparse recovery algorithms (Section 4.6) can be potentially useful is the context of *group testing* and the related problem of *computation on data streams*.

5.3.1 Group testing

Among the historically oldest of all sparse recovery algorithms were developed in the context of *combinatorial group testing* [79], [108], [159]. In this problem we suppose that there are N total items and K anomalous elements that we wish to find. For example, we might wish to identify defective products in an industrial setting, or identify a subset of diseased tissue samples in a medical context. In both of these cases the vector x indicates which elements are anomalous, i.e., $x_i \neq 0$ for the K anomalous elements and $x_i = 0$ otherwise. Our goal is to design a collection of tests that allow us to identify the support (and possibly the values of the nonzeros) of x while also minimizing the number of tests performed. In the simplest practical setting these tests are represented by a binary matrix Φ whose entries ϕ_{ij} are equal to 1 if and only if the j^{th} item is used in the i^{th} test. If the output of the test is linear with respect to the inputs, then the problem of recovering the vector x is essentially the same as the standard sparse recovery problem in compressive sensing.

5.3.2 Computation on data streams

Another application area in which ideas related to compressive sensing have proven useful is computation on *data streams* [40], [135]. As an example of a typical data streaming problem, suppose that x_i represents the number of packets passing through a network router with destination i . Simply storing the vector x is typically infeasible since the total number of possible destinations (represented by a 32-bit IP address) is $N = 2^{32}$. Thus, instead of attempting to store x directly, one can store $y = \Phi x$ where Φ is an $M \times N$ matrix with $M \ll N$. In this context the vector y is often called a *sketch*. Note that in this problem y is computed in a different manner than in the compressive sensing context. Specifically, in the network traffic example we do not ever observe x_i directly, rather we observe increments to x_i (when a packet with destination i passes through the router). Thus we construct y iteratively by adding the i^{th} column to y each time we observe an increment to x_i , which we can do since $y = \Phi x$ is linear. When the network traffic is dominated by traffic to a small number of destinations, the vector x is compressible, and thus the problem of recovering x from the sketch Φx is again essentially the same as the sparse recovery problem in compressive sensing.

⁵This content is available online at <http://cnx.org/content/m37362/1.4/>.

⁶"Introduction to compressive sensing" <http://cnx.org/content/m37172/latest/>

5.4 Compressive medical imaging⁷

5.4.1 MR image reconstruction

Magnetic Resonance Imaging (MRI) is a medical imaging technique based on the core principle that protons in water molecules in the human body align themselves in a magnetic field. MRI machines repeatedly pulse magnetic fields to cause water molecules in the human body to disorient and then reorient themselves, which causes a release of detectable radiofrequencies. We assume that the object to be imaged as a collection of voxels. The MRI's magnetic pulses are sent incrementally along a gradient leading to a different phase and frequency encoding for each column and row of voxels respectively. Abstracting away from the technicalities of the physical process, the magnetic field measured in MRI acquisition corresponds to a Fourier coefficient of the imaged object; the object can then be recovered by an inverse Fourier transform. , we can view the MRI as measuring Fourier samples.

A major limitation of the MRI process is the linear relation between the number of measured data samples and scan times. Long-duration MRI scans are more susceptible to physiological motion artifacts, add discomfort to the patient, and are expensive [119]. Therefore, minimizing scan time without compromising image quality is of direct benefit to the medical community.

The theory of compressive sensing⁸ (CS) can be applied to MR image reconstruction by exploiting the transform-domain sparsity of MR images [120], [121], [122], [179]. In standard MRI reconstruction, undersampling in the Fourier domain results in aliasing artifacts when the image is reconstructed. However, when a known transform renders the object image sparse (Section 2.3) or compressible (Section 2.4), the image can be reconstructed using sparse recovery (Section 4.6) methods. While the discrete cosine and wavelet transforms are commonly used in CS to reconstruct these images, the use of total variation norm minimization also provides high-quality reconstruction.

5.4.2 Electroencephalography

Electroencephalography (EEG) and Magnetoencephalography (MEG) are two popular noninvasive methods to characterize brain function by measuring scalp electric potential distributions and magnetic fields due to neuronal firing. EEG and MEG provide temporal resolution on the millisecond timescale characteristic of neural population activity and can also help to estimate the current sources inside the brain by solving an inverse problem [94].

Models for neuromagnetic sources suggest that the underlying activity is often limited in spatial extent. Based on this idea, algorithms like FOCUSS (Focal Underdetermined System Solution) are used to identify highly localized sources by assuming a sparse model to solve an underdetermined problem [95].

FOCUSS is a recursive linear estimation procedure, based on a weighted pseudo-inverse solution. The algorithm assigns a current (with nonlinear current location parameters) to each element within a region so that the unknown current values can be related linearly to the measurements. The weights at each step are derived from the solution of the previous iterative step. The algorithm converges to a source distribution in which the number of parameters required to describe source currents does not exceed the number of measurements. The initialization determines which of the localized solutions the algorithm converges to.

5.5 Analog-to-information conversion⁹

We now consider the application of compressive sensing¹⁰ (CS) to the problem of designing a system that can acquire a continuous-time signal $x(t)$. Specifically, we would like to build an *analog-to-digital converter* (ADC) that avoids having to sample $x(t)$ at its Nyquist rate when $x(t)$ is sparse. In this context, we will assume that $x(t)$ has some kind of sparse (Section 2.3) structure in the Fourier domain, meaning that it

⁷This content is available online at <<http://cnx.org/content/m37363/1.4/>>.

⁸"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

⁹This content is available online at <<http://cnx.org/content/m37375/1.4/>>.

¹⁰"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

is still bandlimited but that much of the spectrum is empty. We will discuss the different possible signal models for mathematically capturing this structure in greater detail below. For now, the challenge is that our measurement system (Section 3.1) must be built using analog hardware. This imposes severe restrictions on the kinds of operations we can perform.

5.5.1 Analog measurement model

To be more concrete, since we are dealing with a continuous-time signal $x(t)$, we must also consider continuous-time test functions $\{\phi_j(t)\}_{j=1}^M$. We then consider a finite window of time, say $t \in [0, T]$, and would like to collect M measurements of the form

$$y[j] = \int_0^T x(t) \phi_j(t) dt. \quad (5.1)$$

Building an analog system to collect such measurements will require three main components:

1. hardware for generating the test signals $\phi_j(t)$;
2. M correlators that multiply the signal $x(t)$ with each respective $\phi_j(t)$;
3. M integrators with a zero-valued initial state.

We could then sample and quantize the output of each of the integrators to collect the measurements $y[j]$. Of course, even in this somewhat idealized setting, it should be clear that what we can build in hardware will constrain our choice of $\phi_j(t)$ since we cannot reliably and accurately produce (and reproduce) arbitrarily complex $\phi_j(t)$ in analog hardware. Moreover, the architecture described above requires M correlator/integrator pairs operating in parallel, which will be potentially prohibitively expensive both in dollar cost as well as costs such as size, weight, and power (SWAP).

As a result, there have been a number of efforts to design simpler architectures, chiefly by carefully designing structured $\phi_j(t)$. The simplest to describe and historically earliest idea is to choose $\phi_j(t) = \delta(t - t_j)$, where $\{t_j\}_{j=1}^M$ denotes a sequence of M locations in time at which we would like to sample the signal $x(t)$. Typically, if the number of measurements we are acquiring is lower than the Nyquist-rate, then these locations cannot simply be uniformly spaced in the interval $[0, T]$, but must be carefully chosen. Note that this approach simply requires a single traditional ADC with the ability to sample on a non-uniform grid, avoiding the requirement for M parallel correlator/integrator pairs. Such non-uniform sampling systems have been studied in other contexts outside of the CS framework. For example, there exist specialized fast algorithms for the recovery of extremely large Fourier-sparse signals. The algorithm uses samples at a non-uniform sequence of locations that are highly structured, but where the initial location is chosen using a (pseudo)random seed. This literature provides guarantees similar to those available from standard CS [88], [91]. Additionally, there exist frameworks for the sampling and recovery of multi-band signals, whose Fourier transforms are mostly zero except for a few frequency bands. These schemes again use non-uniform sampling patterns based on coset sampling [19], [18], [82], [81], [131], [182]. Unfortunately, these approaches are often highly sensitive to *jitter*, or error in the timing of when the samples are taken.

We will consider a rather different approach, which we call the *random demodulator* [111], [115], [174].¹¹ The architecture of the random demodulator is depicted in Figure 5.1. The analog input $x(t)$ is correlated with a pseudorandom square pulse of ± 1 's, called the *chipping sequence* $p_c(t)$, which alternates between values at a rate of N_a Hz, where N_a Hz is at least as fast as the Nyquist rate of $x(t)$. The mixed signal is integrated over a time period $1/M_a$ and sampled by a traditional integrate-and-dump back-end ADC at M_a Hz $\ll N_a$ Hz. In this case our measurements are given by

$$y[j] = \int_{(j-1)/M_a}^{j/M_a} p_c(t) x(t) dt. \quad (5.2)$$

¹¹A correlator is also known as a “demodulator” due to its most common application: demodulating radio signals.

In practice, data is processed in time blocks of period T , and we define $N = N_a T$ as the number of elements in the chipping sequence, and $M = M_a T$ as the number of measurements. We will discuss the discretization of this model below, but the key observation is that the correlator and chipping sequence operate at a fast rate, while the back-end ADC operates at a low rate. In hardware it is easier to build a high-rate modulator/chipping sequence combination than a high-rate ADC [115]. In fact, many systems already use components of this front end for binary phase shift keying demodulation, as well as for other conventional communication schemes such as CDMA.

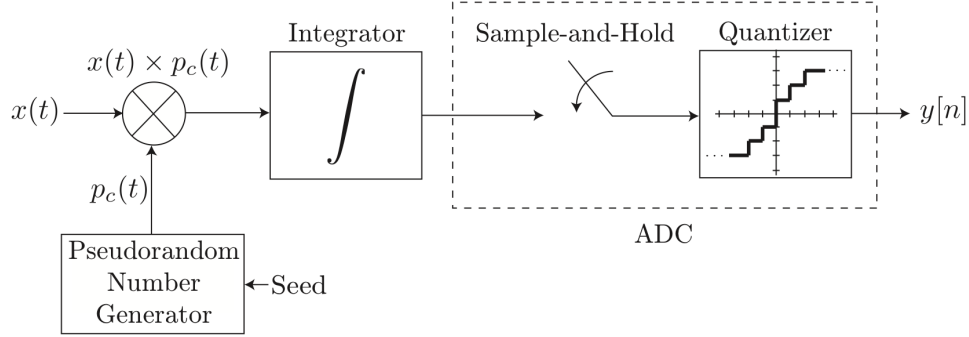


Figure 5.1: Random demodulator block diagram.

5.5.2 Discrete formulation

Although the random demodulator directly acquires compressive measurements without first sampling $x(t)$, it is equivalent to a system which first samples $x(t)$ at its Nyquist-rate to yield a discrete-time vector x , and then applies a matrix Φ to obtain the measurements $y = \Phi x$. To see this we let $p_c[n]$ denote the sequence of ± 1 used to generate the signal $p_c(t)$, i.e., $p_c(t) = p_c[n]$ for $t \in [(n-1)/N_a, n/N_a]$. As an example, consider the first measurement, or the case of $j = 1$. In this case, $t \in [0, 1/M_a]$, so that $p_c(t)$ is determined by $p_c[n]$ for $n = 1, 2, \dots, N_a/M_a$. Thus, from (5.2) we obtain

$$\begin{aligned} y[1] &= \int_0^{1/M_a} p_c(t) x(t) dt \\ &= \sum_{n=1}^{N_a/M_a} p_c[n] \int_{(n-1)/N_a}^{n/N_a} x(t) dt. \end{aligned} \quad (5.3)$$

But since N_a is the Nyquist-rate of $x(t)$, $\int_{(n-1)/N_a}^{n/N_a} x(t) dt$ simply calculates the average value of $x(t)$ on the n^{th} interval, yielding a sample denoted $x[n]$. Thus, we obtain

$$y[1] = \sum_{n=1}^{N_a/M_a} p_c[n] x[n]. \quad (5.4)$$

In general, our measurement process is equivalent to multiplying the signal x with the random sequence of ± 1 's in $p_c[n]$ and then summing every sequential block of N_a/M_a coefficients. We can represent this as a banded matrix Φ containing N_a/M_a pseudorandom ± 1 s per row. For example, with $N = 12$, $M = 4$, and

$T = 1$, such a Φ is expressed as

$$\Phi = \begin{bmatrix} -1 & +1 & +1 & & & & \\ & & & -1 & +1 & -1 & \\ & & & & & & +1 & +1 & -1 \\ & & & & & & & & & +1 & -1 & -1 \end{bmatrix}. \quad (5.5)$$

In general, Φ will have M rows and each row will contain N/M nonzeros. Note that matrices satisfying this structure are extremely efficient to apply, requiring only $O(N)$ computations compared to $O(MN)$ in the general case. This is extremely useful during recovery.

A detailed analysis of the random demodulator in [174] studied the properties of these matrices applied to a particular signal model. Specifically, it is shown that if Ψ represents the $N \times N$ normalized discrete Fourier transform (DFT) matrix, then the matrix $\Phi\Psi$ will satisfy the restricted isometry property (Section 3.3) (RIP) with high probability, provided that

$$M = O(K \log^2(N/K)), \quad (5.6)$$

where the probability is taken with respect to the random choice of $p_c[n]$. This means that if $x(t)$ is a periodic (or finite-length) signal such that once it is sampled it is sparse or compressible in the basis Ψ , then it should be possible to recover $x(t)$ from the measurements provided by the random demodulator. Moreover, it is empirically demonstrated that combining ℓ_1 minimization (Section 4.1) with the random demodulator can recover K -sparse (in Ψ) signals with

$$M \geq CK \log(N/K + 1) \quad (5.7)$$

measurements where $C \approx 1.7$ [174].

Note that the signal model considered in [174] is somewhat restrictive, since even a pure tone will not yield a sparse DFT unless the frequency happens to be equal to k/N_a for some integer k . Perhaps a more realistic signal model is the multi-band signal model of [19], [18], [82], [81], [131], [182], where the signal is assumed to be bandlimited outside of K bands each of bandwidth B , where KB is much less than the total possible bandwidth. It remains unknown whether the random demodulator can be exploited to recover such signals. Moreover, there also exist other CS-inspired architectures that we have not explored in this [3], [149], [176], and this remains an active area of research. We have simply provided an overview of one of the more promising approaches in order to illustrate the potential applicability of the ideas of this course¹² to the problem of analog-to-digital conversion.

5.6 Single-pixel camera¹³

5.6.1 Architecture

Several hardware architectures have been proposed that apply the theory of compressive sensing¹⁴ (CS) in an imaging setting [69], [128], [147]. We will focus on the so-called *single-pixel camera* [69], [162], [163], [184], [185]. The single-pixel camera is an optical computer that sequentially measures the inner products $y[j] = \langle x, \phi_j \rangle$ between an N -pixel sampled version of the incident light-field from the scene under view (denoted by x) and a set of N -pixel test functions $\{\phi_j\}_{j=1}^M$. The architecture is illustrated in Figure 5.2, and an aerial view of the camera in the lab is shown in Figure 5.3. As shown in these figures, the light-field is focused by a lens (Lens 1 in Figure 5.3) not onto a CCD or CMOS sampling array but rather onto a spatial light modulator (SLM). An SLM modulates the intensity of a light beam according to a control

¹²An Introduction to Compressive Sensing <<http://cnx.org/content/col11133/latest/>>

¹³This content is available online at <<http://cnx.org/content/m37369/1.4/>>.

¹⁴"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

signal. A simple example of a transmissive SLM that either passes or blocks parts of the beam is an overhead transparency. Another example is a liquid crystal display (LCD) projector.

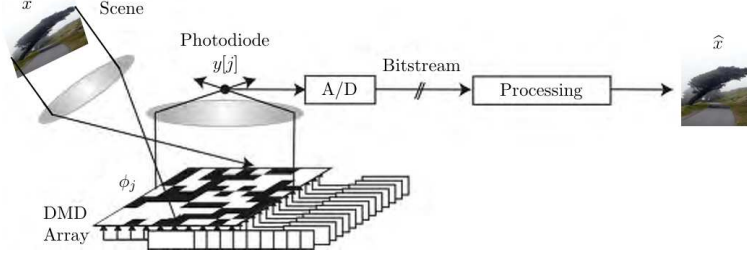


Figure 5.2: Single-pixel camera block diagram. Incident light-field (corresponding to the desired image x) is reflected off a digital micromirror device (DMD) array whose mirror orientations are modulated according to the pseudorandom pattern ϕ_j supplied by a random number generator. Each different mirror pattern produces a voltage at the single photodiode that corresponds to one measurement $y[j]$.

The Texas Instruments (TI) digital micromirror device (DMD) is a reflective SLM that selectively redirects parts of the light beam. The DMD consists of an array of bacterium-sized, electrostatically actuated micro-mirrors, where each mirror in the array is suspended above an individual static random access memory (SRAM) cell. Each mirror rotates about a hinge and can be positioned in one of two states (± 10 degrees from horizontal) according to which bit is loaded into the SRAM cell; thus light falling on the DMD can be reflected in two directions depending on the orientation of the mirrors.

Each element of the SLM corresponds to a particular element of ϕ_j (and its corresponding pixel in x). For a given ϕ_j , we can orient the corresponding element of the SLM either towards (corresponding to a 1 at that element of ϕ_j) or away from (corresponding to a 0 at that element of ϕ_j) a second lens (Lens 2 in Figure 5.3). This second lens collects the reflected light and focuses it onto a single photon detector (the single pixel) that integrates the product of x and ϕ_j to compute the measurement $y[j] = \langle x, \phi_j \rangle$ as its output voltage. This voltage is then digitized by an A/D converter. Values of ϕ_j between 0 and 1 can be obtained by dithering the mirrors back and forth during the photodiode integration time. By reshaping x into a column vector and the ϕ_j into row vectors, we can thus model this system as computing the product $y = \Phi x$, where each row of Φ corresponds to a ϕ_j . To compute randomized measurements, we set the mirror orientations ϕ_j randomly using a pseudorandom number generator, measure $y[j]$, and then repeat the process M times to obtain the measurement vector y .

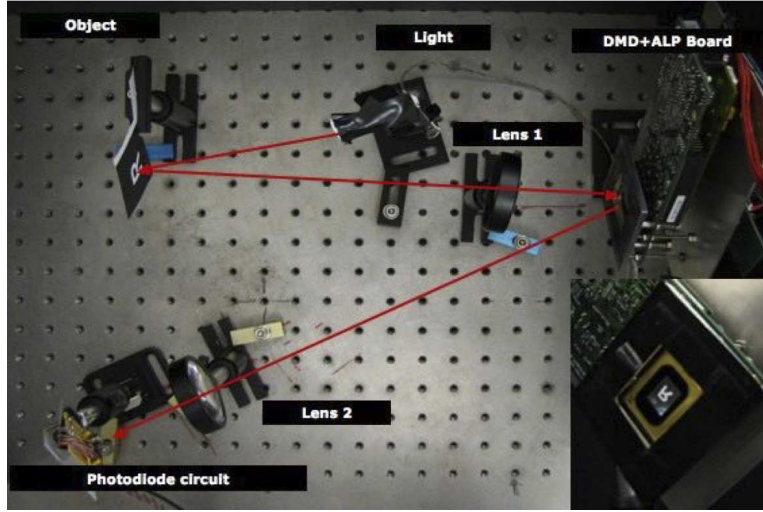


Figure 5.3: Aerial view of the single-pixel camera in the lab.

The single-pixel design reduces the required size, complexity, and cost of the photon detector array down to a single unit, which enables the use of exotic detectors that would be impossible in a conventional digital camera. Example detectors include a photomultiplier tube or an avalanche photodiode for low-light (photon-limited) imaging, a sandwich of several photodiodes sensitive to different light wavelengths for multimodal sensing, a spectrometer for hyperspectral imaging, and so on.

In addition to sensing flexibility, the practical advantages of the single-pixel design include the facts that the quantum efficiency of a photodiode is higher than that of the pixel sensors in a typical CCD or CMOS array and that the fill factor of a DMD can reach 90% whereas that of a CCD/CMOS array is only about 50%. An important advantage to highlight is that each CS measurement receives about $N/2$ times more photons than an average pixel sensor, which significantly reduces image distortion from dark noise and read-out noise.

The single-pixel design falls into the class of multiplex cameras. The baseline standard for multiplexing is classical raster scanning, where the test functions $\{\phi_j\}$ are a sequence of delta functions $\delta[n - j]$ that turn on each mirror in turn. There are substantial advantages to operating in a CS rather than raster scan mode, including fewer total measurements (M for CS rather than N for raster scan) and significantly reduced dark noise. See [69] for a more detailed discussion of these issues.

Figure 5.4 (a) and (b) illustrates a target object (a black-and-white printout of an “R”) x and reconstructed image $\hat{x}$ taken by the single-pixel camera prototype in Figure 5.3 using $N = 256 \times 256$ and $M = N/50$ [69]. Figure 5.4(c) illustrates an $N = 256 \times 256$ color single-pixel photograph of a printout of the Mandrill test image taken under low-light conditions using RGB color filters and a photomultiplier tube with $M = N/10$. In both cases, the images were reconstructed using total variation minimization, which is closely related to wavelet coefficient ℓ_1 minimization (Section 4.1) [30].

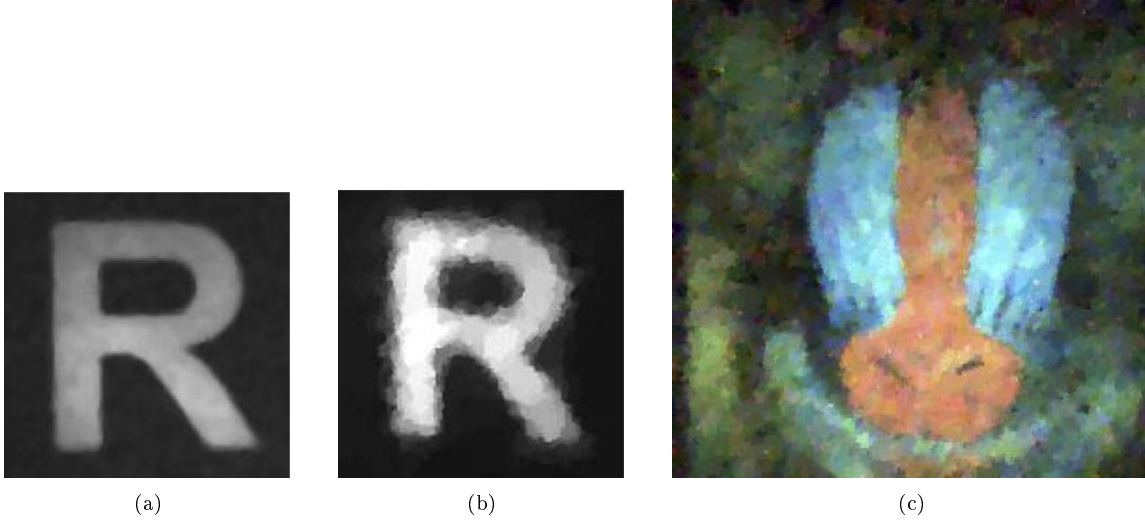


Figure 5.4: Sample image reconstructions from single-pixel camera. (a) 256×256 conventional image of a black-and-white “R”. (b) Image reconstructed from $M = 1300$ single-pixel camera measurements ($50 \times$ sub-Nyquist). (c) 256×256 pixel color reconstruction of a printout of the Mandrill test image imaged in a low-light setting using a single photomultiplier tube sensor, RGB color filters, and $M = 6500$ random measurements.

5.6.2 Discrete formulation

Since the DMD array is programmable, we can employ arbitrary test functions ϕ_j . However, even when we restrict the ϕ_j to be $\{0, 1\}$ -valued, storing these patterns for large values of N is impractical. Furthermore, as noted above, even pseudorandom Φ can be computationally problematic during recovery. Thus, rather than purely random Φ , we can also consider Φ that admit a fast transform-based implementation by taking random submatrices of a Walsh, Hadamard, or noiselet transform [37], [26]. We will describe the Walsh transform for the purpose of illustration.

We will suppose that N is a power of 2 and let $W_{\log_2 N}$ denote the $N \times N$ Walsh transform matrix. We begin by setting $W_0 = 1$, and we now define W_j recursively as

$$W_j = \frac{1}{\sqrt{2}} \begin{bmatrix} W_{j-1} & W_{j-1} \\ W_{j-1} & -W_{j-1} \end{bmatrix}. \quad (5.8)$$

This construction produces an orthonormal matrix with entries of $\pm 1/\sqrt{N}$ that admits a fast implementation requiring $O(N \log N)$ computations to apply. As an example, note that

$$W_1 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \quad (5.9)$$

and

$$W_2 = \frac{1}{2} \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 \end{bmatrix}. \quad (5.10)$$

We can exploit these constructions as follows. Suppose that $N = 2^B$ and generate W_B . Let I_Γ denote a $M \times N$ submatrix of the identity I obtained by picking a random set of M rows, so that $I_\Gamma W_B$ is the submatrix of W_B consisting of the rows of W_B indexed by Γ . Furthermore, let D denote a random $N \times N$ permutation matrix. We can generate Φ as

$$\Phi = \left(\frac{1}{2} \sqrt{N} I_\Gamma W_B + \frac{1}{2} \right) D. \quad (5.11)$$

Note that $\frac{1}{2} \sqrt{N} I_\Gamma W_B + \frac{1}{2}$ merely rescales and shifts $I_\Gamma W_B$ to have $\{0, 1\}$ -valued entries, and recall that each row of Φ will be reshaped into a 2-D matrix of numbers that is then displayed on the DMD array. Furthermore, D can be thought of as either permuting the pixels or permuting the columns of W_B . This step adds some additional randomness since some of the rows of the Walsh matrix are highly correlated with coarse scale wavelet basis functions — but permuting the pixels eliminates this structure. Note that at this point we do not have any strict guarantees that such Φ combined with a wavelet basis Ψ will yield a product $\Phi\Psi$ satisfying the restricted isometry property (Section 3.3), but this approach seems to work well in practice.

5.7 Hyperspectral imaging¹⁵

Standard digital color images of a scene of interest consist of three components — red, green and blue — which contain the intensity level for each of the pixels in three different groups of wavelengths. This concept has been extended in the *hyperspectral* and *multispectral* imaging sensing modalities, where the data to be acquired consists of a three-dimensional *datacube* that has two spatial dimensions x and y and one spectral dimension λ .

In simple terms, a datacube is a 3-D function $f(x, y, \lambda)$ that can be represented as a stacking of intensities of the scene at different wavelengths. An example datacube is shown in Figure 5.5. Each of its entries is called a voxel. We also define a pixel's *spectral signature* as the stacking of its voxels in the spectral dimension $f(x, y) = \{f(x, y, \lambda)\}_\lambda$. The spectral signature of a pixel can give a wealth of information about the corresponding point in the scene that is not captured by its color. For example, using spectral signatures, it is possible to identify the type of material observed (for example, vegetation vs. ground vs. water), or its chemical composition.

Datacubes are high-dimensional, since the standard number of pixels present in a digitized image is multiplied by the number of spectral bands desired. However, considerable structure is present in the observed data. The spatial structure common in natural images is also observed in hyperspectral imaging, while each pixel's spectral signature is usually smooth.

¹⁵This content is available online at <http://cnx.org/content/m37370/1.4/>.

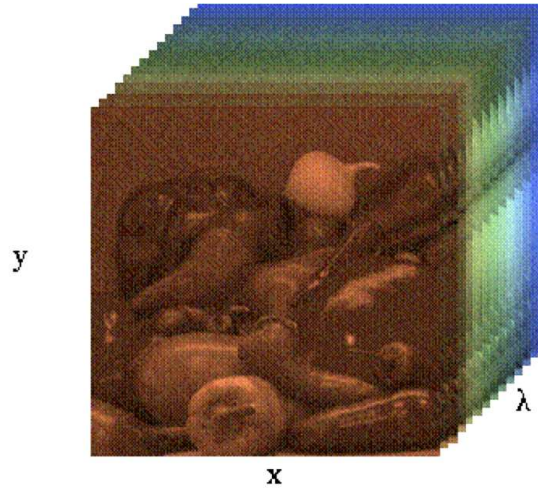


Figure 5.5: Example hyperspectral datacube, with labeled dimensions.

Compressive sensing¹⁶ (CS) architectures for hyperspectral imaging perform lower-dimensional projections that multiplex in the spatial domain, the spectral domain, or both. Below, we detail three example architectures, as well as three possible models to sparsify hyperspectral datacubes.

5.7.1 Compressive hyperspectral imaging architectures

5.7.1.1 Single pixel hyperspectral camera

The single pixel camera (Section 5.6) uses a single photodetector to record random projections of the light emanated from the image, with the different random projections being captured in sequence. A single pixel hyperspectral camera requires a light modulating element that is reflective across the wavelengths of interest, as well as a sensor that can record the desired spectral bands separately [109]. A block diagram is shown in Figure 5.6.

The single sensor consists of a single spectrometer that spans the necessary wavelength range, which replaces the photodiode. The spectrometer records the intensity of the light reflected by the modulator in each wavelength. The same digital micromirror device (DMD) provides reflectivity for wavelengths from near infrared to near ultraviolet. Thus, by converting the datacube into a vector sorted by spectral band, the matrix that operates on the data to obtain the CS measurements is represented as

$$\Phi = \begin{bmatrix} \Phi_{x,y} & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \Phi_{x,y} & \cdots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \cdots & \Phi_{x,y} \end{bmatrix}. \quad (5.12)$$

This architecture performs multiplexing only in the spatial domain, i.e. dimensions x and y , since there is no mixing of the different spectral bands along the dimension λ .

¹⁶"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

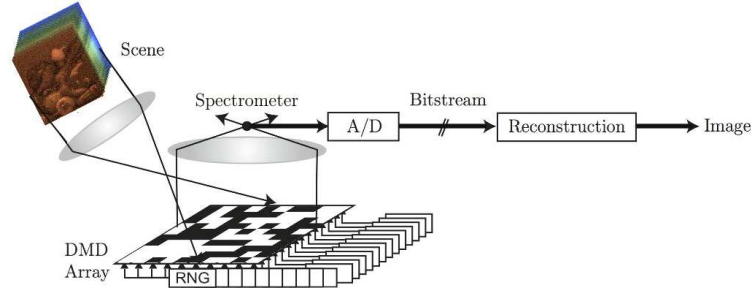


Figure 5.6: Block diagram for a single pixel hyperspectral camera. The photodiode is replaced by a spectrometer that captures the modulated light intensity for all spectral bands, for each of the CS measurements.

5.7.1.2 Dual disperser coded aperture snapshot spectral imager

The dual disperser coded aperture snapshot spectral imager (DD-CASSI), shown in Figure 5.7, is an architecture that combines separate multiplexing in the spatial and spectral domain, which is then sensed by a wide-wavelength sensor/pixel array, thus flattening the spectral dimension [86].

First, a dispersive element separates the different spectral bands, which still overlap in the spatial domain. In simple terms, this element shears the datacube, with each spectral slice being displaced from the previous by a constant amount in the same spatial dimension. The resulting datacube is then masked using the coded aperture, whose effect is to "punch holes" in the sheared datacube by blocking certain pixels of light. Subsequently, a second dispersive element acts on the masked, sheared datacube; however, this element shears in the opposite direction, effectively inverting the shearing of the first dispersive element. The resulting datacube is upright, but features "sheared" holes of datacube voxels that have been masked out.

The resulting modified datacube is then received by a sensor array, which flattens the spectral dimension by measuring the sum of all the wavelengths received; the received light field resembles the target image, allowing for optical adjustments such as focusing. In this way, the measurements consist of full sampling in the spatial x and y dimensions, with an aggregation effect in the spectral λ dimension.

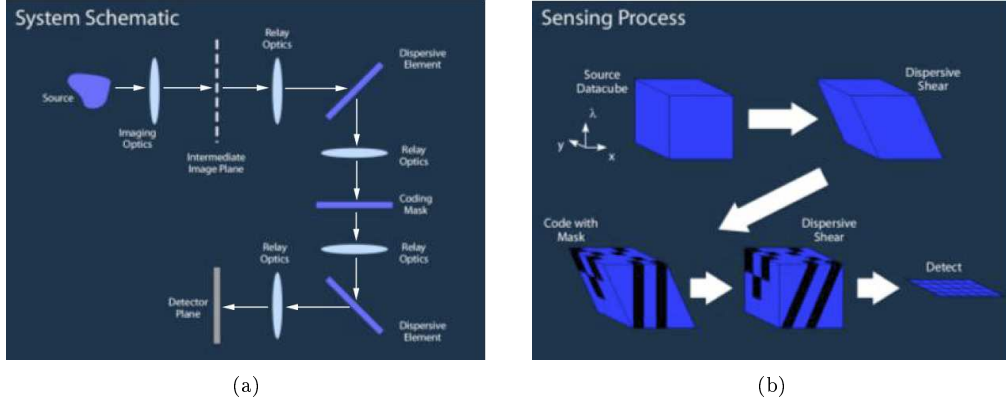


Figure 5.7: Dual disperser coded aperture snapshot spectral imager (DD-CASSI). (a) Schematic of the DD-CASSI components. (b) Illustration of the datacube processing performed by the components.

5.7.1.3 Single disperser coded aperture snapshot spectral imager

The single disperser coded aperture snapshot spectral imager (SD-CASSI), shown in Figure 5.8, is a simplification of the DD-CASSI architecture in which the first dispersive element is removed [183]. Thus, the light field received at the sensors does not resemble the target image. Furthermore, since the shearing is not reversed, the area occupied by the sheared datacube is larger than that of the original datacube, requiring a slightly larger number of pixels for the capture.

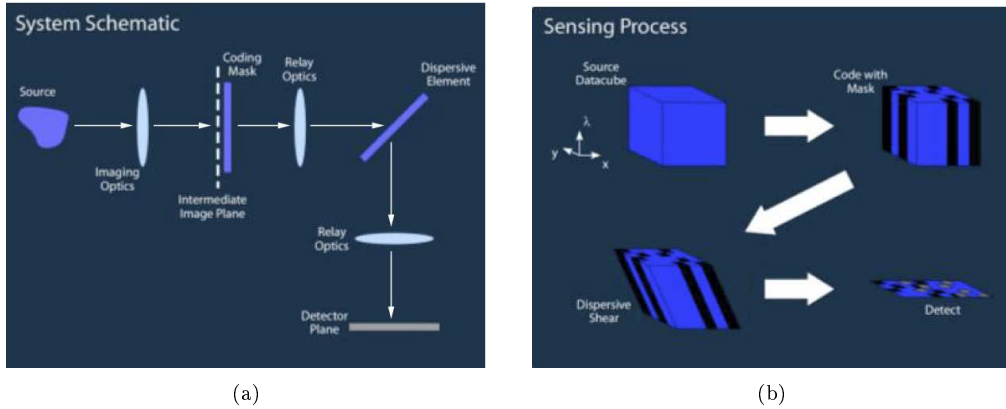


Figure 5.8: Single disperser coded aperture snapshot spectral imager (SD-CASSI). (a) Schematic of the SD-CASSI components. (b) Illustration of the datacube processing performed by the components.

5.7.2 Sparsity structures for hyperspectral datacubes

5.7.2.1 Dyadic Multiscale Partitioning

This sparsity (Section 2.3) structure assumes that the spectral signature for all pixels in a neighborhood is close to constant; that is, that the datacube is piecewise constant with smooth borders in the spatial dimensions. The complexity of an image is then given by the number of spatial dyadic squares with constant spectral signature necessary to accurately approximate the datacube; see Figure 5.9. A reconstruction algorithm then searches for the signal of lowest complexity (i.e., with the fewest dyadic squares) that generates compressive measurements close to those observed [86].

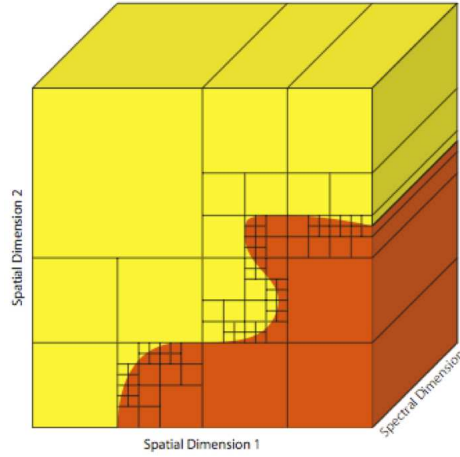


Figure 5.9: Example dyadic square partition for piecewise spatially constant datacube.

5.7.2.2 Spatial-only sparsity

This sparsity structure operates on each spectral band separately and assumes the same type of sparsity structure for each band [183]. The sparsity basis is drawn from those commonly used in images, such as wavelets, curvelets, or the discrete cosine basis. Since each basis operates only on a band, the resulting sparsity basis for the datacube can be represented as a block-diagonal matrix:

$$\Psi = \begin{bmatrix} \Psi_{x,y} & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \Psi_{x,y} & \cdots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \cdots & \Psi_{x,y} \end{bmatrix}. \quad (5.13)$$

5.7.2.3 Kronecker product sparsity

This sparsity structure employs separate sparsity bases for the spatial dimensions and the spectral dimension, and builds a sparsity basis for the datacube using the Kronecker product of these two [68]:

$$\Psi = \Psi_{\lambda} \otimes \Psi_{x,y} = \begin{bmatrix} \Psi_{\lambda} [1, 1] \Psi_{x,y} & \Psi_{\lambda} [1, 2] \Psi_{x,y} & \cdots \\ \Psi_{\lambda} [2, 1] \Psi_{x,y} & \Psi_{\lambda} [2, 2] \Psi_{x,y} & \cdots \\ \vdots & \vdots & \ddots \end{bmatrix}. \quad (5.14)$$

In this manner, the datacube sparsity bases simultaneously enforces both spatial and spectral structure, potentially achieving a sparsity level lower than the sums of the spatial sparsities for the separate spectral slices, depending on the level of structure between them and how well can this structure be captured through sparsity.

5.7.3 Summary

Compressive sensing¹⁷ will make the largest impact in applications with very large, high dimensional datasets that exhibit considerable amounts of structure. Hyperspectral imaging is a leading example of such applications; the sensor architectures and data structure models surveyed in this module show initial promising work in this new direction, enabling new ways of simultaneously sensing and compressing such data. For standard sensing architectures, the data structures surveyed also enable new transform coding-based compression schemes.

5.8 Compressive processing of manifold-modeled data¹⁸

A powerful data model for many applications is the geometric notion of a low-dimensional *manifold*. Data that possesses merely K intrinsic degrees of freedom” can be assumed to lie on a K -dimensional manifold in the high-dimensional ambient space. Once the manifold model is identified, any point on it can be represented using essentially K pieces of information. For instance, suppose a stationary camera of resolution N observes a truck moving down along a straight line on a highway. Then, the set of images captured by the camera forms a 1-dimensional manifold in the image space $\mathbb{R}^N$. Another example is the set of images captured by a static camera observing a cube that rotates in 3 dimensions. (Figure 5.10).

¹⁷“Introduction to compressive sensing” <<http://cnx.org/content/m37172/latest/>>

¹⁸This content is available online at <<http://cnx.org/content/m37371/1.6/>>.

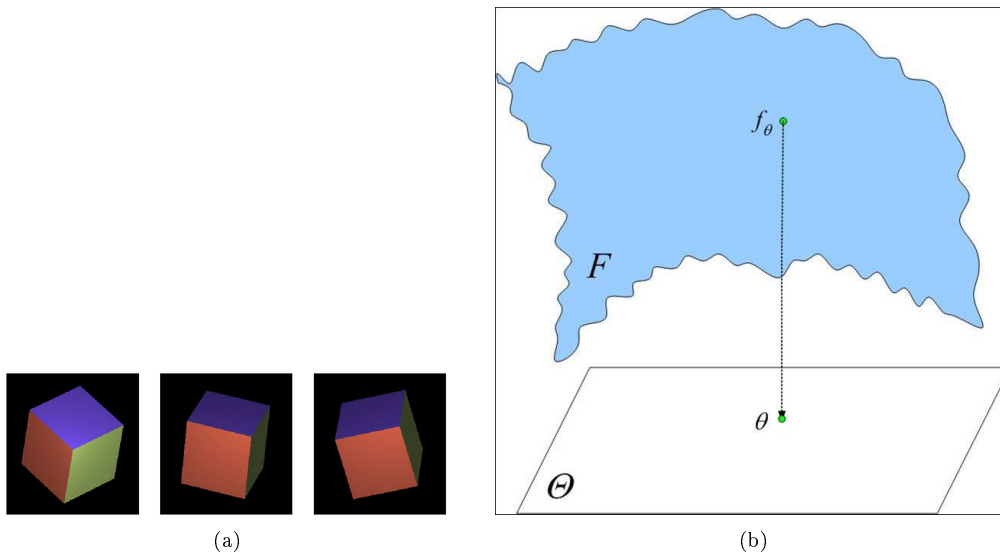


Figure 5.10: (a) A rotating cube has 3 degrees of freedom, thus giving rise to a 3-dimensional manifold in image space. (b) Illustration of a manifold F parametrized by a K -dimensional vector θ .

In many applications, it is beneficial to explicitly characterize the structure (alternately, identify the parameters) of the manifold formed by a set of observed signals. This is known as *manifold learning* and has been the subject of considerable study over the last several years; well-known manifold learning algorithms include Isomap [166], LLE [151], and Hessian eigenmaps [61]. An informal example is as follows: if a 2-dimensional manifold were to be imagined as the surface of a twisted sheet of rubber, manifold learning can be described as the process of “unraveling” the sheet and stretching it out on a 2D flat surface. Figure 5.11 indicates the performance of Isomap on a simple 2-dimensional dataset comprising of images of a translating disk.

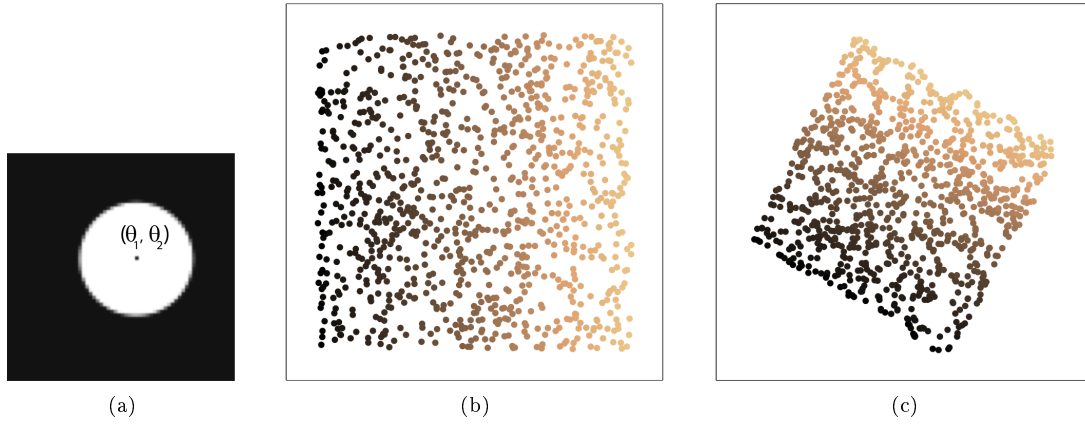


Figure 5.11: (a) Input data consisting of 1000 images of a disk shifted in $K = 2$ dimensions, parametrized by an articulation vector (θ_1, θ_2) . (b) True θ_1 and θ_2 values of the sampled data. (c) Isomap embedding learned from original data in $\mathbb{R}^N$.

A *linear, nonadaptive* manifold dimensionality reduction technique has recently been introduced that employs the technique of random projections [8]. Consider a K -dimensional manifold $\mathcal{M}$ in the ambient space $\mathbb{R}^N$ and its projection onto a random subspace of dimension $M = CK \log(N)$; note that $K < M < N$. The result of [8] is that the pairwise metric structure of sample points from $\mathcal{M}$ is preserved with high accuracy under projection from $\mathbb{R}^N$ to $\mathbb{R}^M$. This is analogous to the result for compressive sensing¹⁹ of sparse (Section 2.3) signals (see "The restricted isometry property" (Section 3.3); however, the difference is that the number of projections required to preserve the ensemble structure does *not* depend on the sparsity of the individual images, but rather on the dimension of the underlying manifold.

This result has far reaching implications; it suggests that a wide variety of signal processing tasks can be performed *directly on the random projections* acquired by these devices, thus saving valuable sensing, storage and processing costs. In particular, this enables provably efficient manifold learning in the projected domain [99]. Figure 5.12 illustrates the performance of Isomap on the translating disk dataset under varying numbers of random projections.

¹⁹"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

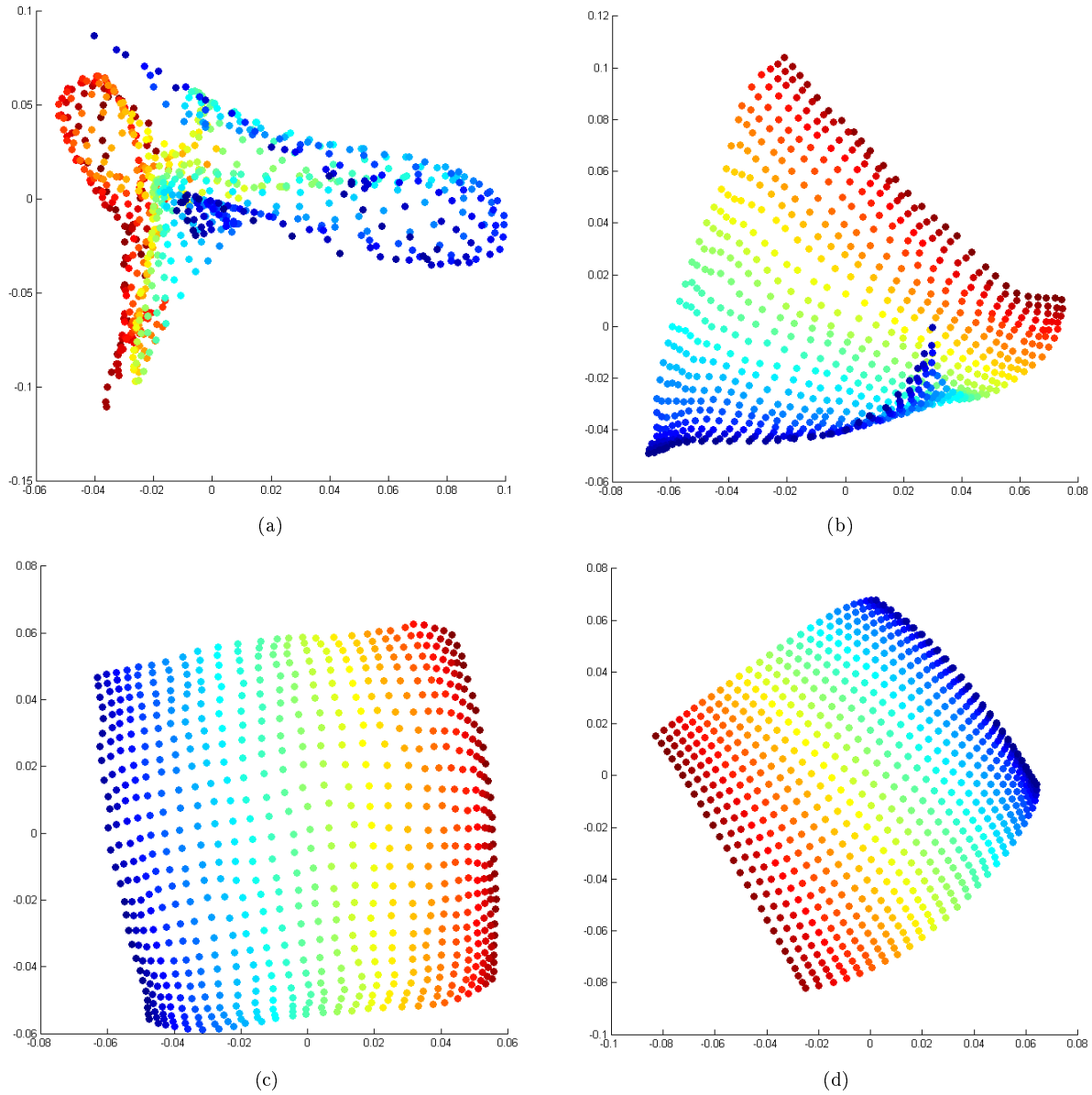


Figure 5.12: Isomap embeddings learned from random projections of the 625 images of shifting squares. (a) 25 random projections; (b) 50 random projections; (c) 25 random projections; (d) full data.

The advantages of random projections extend even to cases where the original data is available in the ambient space $\mathbb{R}^N$. For example, consider a wireless network of cameras observing a static scene. The set of images captured by the cameras can be visualized as living on a low-dimensional manifold in the image space. To perform joint image analysis, the following steps might be executed:

1. **Collate:** Each camera node transmits its respective captured image (of size N) to a central processing unit.

2. **Preprocess:** The central processor estimates the *intrinsic dimension* K of the underlying image manifold.
3. **Learn:** The central processor performs a nonlinear embedding of the data points – for instance, using Isomap [166] – into a K -dimensional Euclidean space, using the estimate of K from the previous step.

In situations where N is large and communication bandwidth is limited, the dominating costs will be in the first transmission/collation step. To reduce the communication expense, one may perform nonlinear image compression (such as JPEG) at each node before transmitting to the central processing. However, this requires a good deal of processing power at each sensor, and the compression would have to be undone during the learning step, thus adding to overall computational costs.

As an alternative, every camera could encode its image by computing (either directly or indirectly) a small number of random projections to communicate to the central processor [47]. These random projections are obtained by linear operations on the data, and thus are cheaply computed. Clearly, in many situations it will be less expensive to store, transmit, and process such randomly projected versions of the sensed images. The method of random projections is thus a powerful tool for ensuring the stable embedding of low-dimensional manifolds into an intermediate space of reasonable size. It is now possible to think of settings involving a huge number of low-power devices that inexpensively capture, store, and transmit a very small number of measurements of high-dimensional data.

5.9 Inference using compressive measurements²⁰

While the compressive sensing²¹ (CS) literature has focused almost exclusively on problems in signal reconstruction/approximation (Section 4.6), this is frequently not necessary. For instance, in many signal processing applications (including computer vision, digital communications and radar systems), signals are acquired only for the purpose of making a detection or classification decision. Tasks such as detection do not require a reconstruction of the signal, but only require estimates of the relevant *sufficient statistics* for the problem at hand.

As a simple example, suppose a surveillance system (based on compressive imaging) observes the motion of a person across a static background. The relevant information to be extracted from the data acquired by this system would be, for example, the identity of the person, or the location of this person with respect to a predefined frame of coordinates. There are two ways of doing this:

- Reconstruct the full data using standard sparse recovery (Section 4.6) techniques and apply standard computer vision/inference algorithms on the reconstructed images.
- Develop an inference test which operates *directly* on the compressive measurements, without ever reconstructing the full images.

A crucial property that enables the design of compressive inference algorithms is the *information scalability* property of compressive measurements. This property arises from the following two observations:

- For certain signal models, the action of a *random* linear function on the set of signals of interest preserves enough information to perform inference tasks on the observed measurements.
- The *number* of random measurements required to perform the inference task usually depends on the nature of the inference task. Informally, we observe that more sophisticated tasks require more measurements.

We examine three possible inference problems for which algorithms that *directly operate on the compressive measurements* can be developed: detection (determining the presence or absence of an information-bearing signal), classification (assigning the observed signal to one of two (or more) signal classes), and parameter estimation (calculating a *function* of the observed signal).

²⁰This content is available online at <<http://cnx.org/content/m37372/1.4/>>.

²¹"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

5.9.1 Detection

In detection one simply wishes to answer the question: is a (known) signal present in the observations? To solve this problem, it suffices to estimate a relevant *sufficient statistic*. Based on a concentration of measure inequality, it is possible to show that such sufficient statistics for a detection problem can be accurately estimated from random projections, where the quality of this estimate depends on the signal to noise ratio (SNR) [45]. We make no assumptions on the signal of interest s , and hence we can build systems capable of detecting s even when it is not known in advance. Thus, we can use random projections for dimensionality-reduction in the detection setting without knowing the relevant structure.

In the case where the class of signals of interest corresponds to a low dimensional subspace, a truncated, simplified sparse approximation can be applied as a detection algorithm; this has been dubbed as IDEA [70]. In simple terms, the algorithm will mark a detection when a large enough amount of energy from the measurements lies in the projected subspace. Since this problem does not require accurate estimation of the signal values, but rather whether it belongs in the subspace of interest or not, the number of measurements necessary is much smaller than that required for reconstruction, as shown in Figure 5.13.

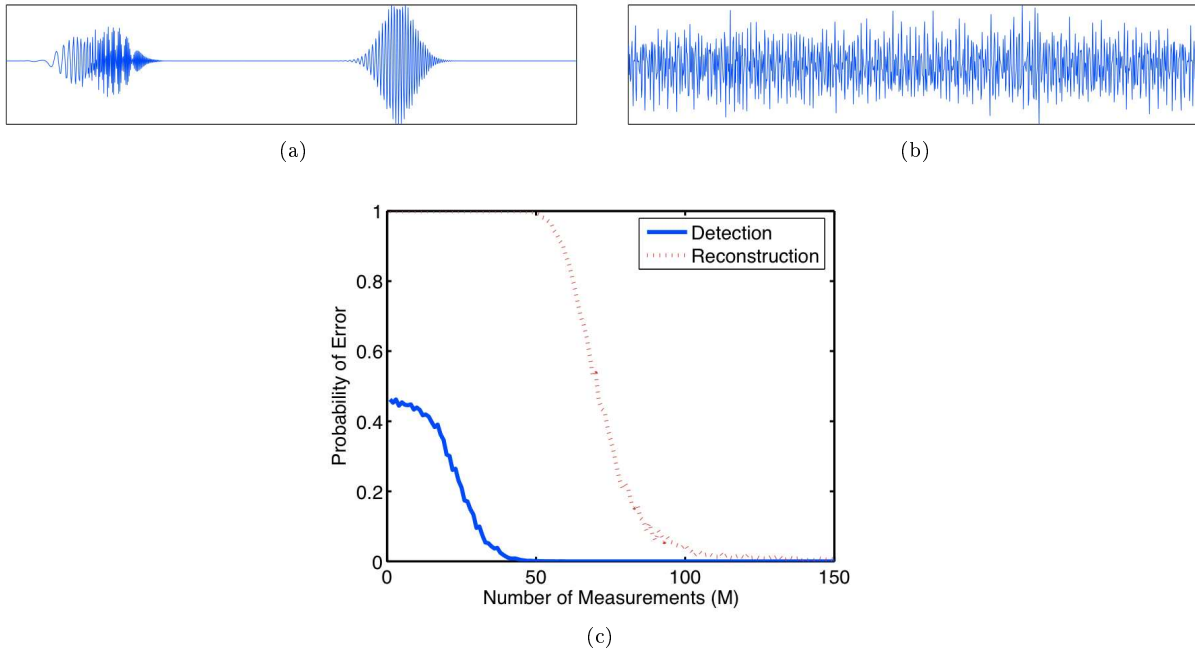


Figure 5.13: Performance for IDEA. (Top) Sample wideband chirp signal and same chirp embedded in strong narrowband interference. (Bottom) Probability of error to reconstruct and detect chirp signals embedded in strong sinusoidal interference ($SIR = -6$ dB) using greedy algorithms. In this case, detection requires $3\times$ fewer measurements and $4\times$ fewer computations than reconstruction for an equivalent probability of success. Taken from [70].

5.9.2 Classification

Similarly, random projections have long been used for a variety of classification and clustering problems. The Johnson-Lindenstrauss Lemma is often exploited in this setting to compute approximate nearest neighbors,

which is naturally related to classification. The key result that random projections result in an isometric embedding allows us to generalize this work to several new classification algorithms and settings [45].

Classification can also be performed when more elaborate models are used for the different classes. Suppose the signal/image class of interest can be modeled as a low-dimensional manifold (Section 5.8) in the ambient space. In such case it can be shown that, even under random projections, certain geometric properties of the signal class are preserved up to a small distortion; for example, interpoint Euclidean (ℓ_2) distances are preserved [9]. This enables the design of classification algorithms in the *projected* domain. One such algorithm is known as the smashed filter [46]. As an example, under equal distribution among classes and a gaussian noise setting, the smashed filter is equivalent to building a nearest-neighbor (NN) classifier in the measurement domain. Further, it has been shown that for a K -dimensional manifold, $M = O(K \log N)$ measurements are sufficient to perform reliable compressive classification. Thus, the number of measurements scales as the dimension of the signal class, as opposed to the *sparsity* of the individual signal. Some example results are shown in Figure 5.14(a).

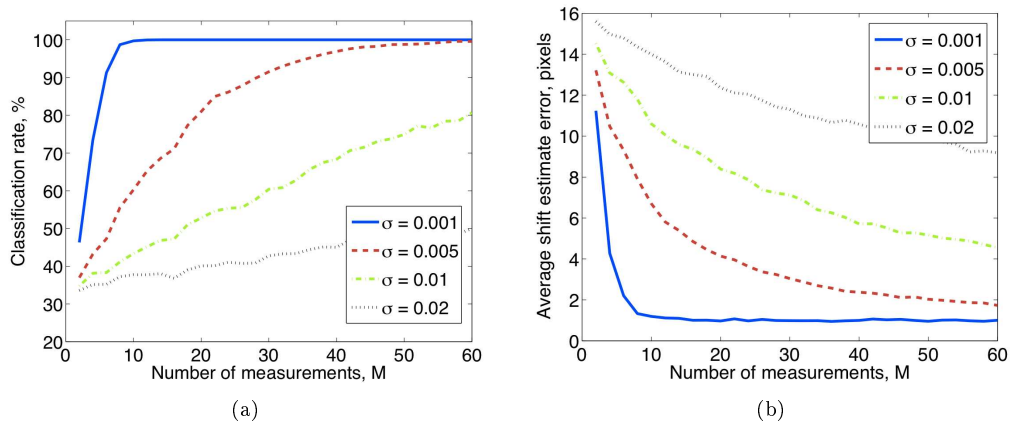


Figure 5.14: Results for smashed filter image classification and parameter estimation experiments. (a) Classification rates and (b) average estimation error for varying number of measurements M and noise levels σ for a set of images of several objects under varying shifts. As M increases, the distances between the manifolds increase as well, thus increasing the noise tolerance and enabling more accurate estimation and classification. Thus, the classification and estimation performances improve as σ decreases and M increases in all cases. Taken from [71].

5.9.3 Estimation

Consider a signal $x \in \mathbb{R}^N$, and suppose that we wish to estimate some function $f(x)$ but only observe the measurements $y = \Phi x$, where Φ is again an $M \times N$ matrix. The data streaming (Section 5.8) community has previously analyzed this problem for many common functions, such as linear functions, ℓ_p norms, and histograms. These estimates are often based on so-called *sketches*, which can be thought of as random projections.

As an example, in the case where f is a *linear* function, one can show that the estimation error (relative to the norms of x and f) can be bounded by a constant determined by M . This result holds for a wide class of random matrices, and can be viewed as a straightforward consequence of the same concentration of measure inequality (Section 3.8) that has proven useful for CS and in proving the JL Lemma [45].

Parameter estimation can also be performed when the signal class is modeled as a low-dimensional manifold. Suppose an observed signal x can be parameterized by a K -dimensional parameter vector θ , where $K \ll N$. Then, it can be shown that with $O(K \log N)$ measurements, the parameter vector can be obtained via *multiscale manifold navigation* in the compressed domain [71]. Some example results are shown in Figure 5.14(b).

5.10 Compressive sensor networks²²

Sparse (Section 2.3) and compressible (Section 2.4) signals are present in many sensor network applications, such as environmental monitoring, signal field recording and vehicle surveillance. Compressive sensing²³ (CS) has many properties that make it attractive in this settings, such as its low complexity sensing and compression, its universality and its graceful degradation. CS is robust to noise, and allows querying more nodes to obey further detail on signals as they become interesting. Packet drops also do not harm the network nearly as much as many other protocols, only providing a marginal loss for each measurement not obtained by the receiver. As the network becomes more congested, data can be scaled back smoothly.

Thus CS can enable the design of generic compressive sensors that perform random or incoherent projections.

Several methods for using CS in sensor networks have been proposed. Decentralized methods pass data throughout the network, from neighbor to neighbor, and allow the decoder to probe any subset of nodes. In contrast, centralized methods require all information to be transmitted to a centralized data center, but reduce either the amount of information that must be transmitted or the power required to do so. We briefly summarize each class below.

5.10.1 Decentralized algorithms

Decentralized algorithms enable the calculation of compressive measurements at each sensor in the network, thus being useful for applications where monitoring agents traverse the network during operation.

5.10.1.1 Randomized gossiping

In randomized gossiping [146], each sensor communicates M random projection of its data sample to a random set of nodes, in each stage aggregating and forwarding the observations received to a new set of random nodes. In essence, a spatial dot product is being performed as each node collects and aggregates information, compiling a sum of the weighted samples to obtain M CS measurements which becomes more accurate as more rounds of random gossiping occur. To recover the data, a basis that provides data sparsity (or at least compressibility) is required, as well as the random projections used. However, this information does not need to be known while the data is being passed.

The method can also be applied when each sensor observes a compressible signal. In this case, each sensor computes multiple random projections of the data and transmits them using randomized gossiping to the rest of the network. A potential drawback of this technique is the amount of storage required per sensor, as it could be considerable for large networks. In this case, each sensor can store the data from only a subset of the sensors, where each group of sensors of a certain size will be known to contain CS measurements for all the data in the network. To maintain a constant error as the network size grows, the number of transmissions becomes $\Theta(kMn^2)$, where k represents the number of groups in which the data is partitioned, M is the number of values desired from each sensor and n are the number of nodes in the network. The results can be improved by using geographic gossiping algorithms [53].

²²This content is available online at <http://cnx.org/content/m37373/1.3/>.

²³"Introduction to compressive sensing" <http://cnx.org/content/m37172/latest/>

5.10.1.2 Distributed sparse random projections

A second method modifies the randomized gossiping approach by limiting the number of communications each node must perform, in order to reduce overall power consumption [187]. Each data node takes M projections of its data, passing along information to a small set of L neighbors, and summing the observations; the resulting CS measurements are sparse, since $N - L$ of each row's entries will be zero. Nonetheless, these projections can still be used as CS measurements with quality similar to that of full random projections. Since the CS measurement matrix formed by the data nodes is sparse, a relatively small amount of communication is performed by each encoding node and the overall power required for transmission is reduced.

5.10.2 Centralized algorithms

Decentralized algorithms are used when the sensed data must be routed to a single location; this architecture is common in sensor networks where low power, simple nodes perform sensing and a powerful central location performs data processing.

5.10.2.1 Compressive wireless sensing

Compressive wireless sensing (CWS) emphasizes the use of synchronous communication to reduce the transmission power of each sensor [4]. In CWS, each sensor calculates a noisy projection of their data sample. Each sensor then transmits the calculated value by analog modulation and transmission of a communication waveform. The projections are aggregated at the central location by the receiving antenna, with further noise being added. In this way, the fusion center receives the CS measurements, from which it can perform reconstruction using knowledge of the random projections.

A drawback of this method is the required accurate synchronization. Although CWS is constraining the power of each node, it is also relying on constructive interference to increase the power received by the data center. The nodes themselves must be accurately synchronized to know when to transmit their data. In addition, CWS assumes that the nodes are all at approximately equal distances from the fusion center, an assumption that is acceptable only when the receiver is far away from the sensor network. Mobile nodes could also increase the complexity of the transmission protocols. Interference or path issues also would have a large effect on CWS, limiting its applicability.

If these limitations are addressed for a suitable application, CWS does offer great power benefits when very little is known about the data beyond sparsity in a fixed basis. Distortion will be proportional to $M^{-2\alpha/(2\alpha+1)}$, where α is some positive constant based on the network structure. With much more a priori information about the sensed data, other methods will achieve distortions proportional to $M^{-2\alpha}$.

5.10.2.2 Distributed compressive sensing

Distributed Compressive Sensing (DCS) provides several models for combining neighboring sparse signals, relying on the fact that such sparse signals may be similar to each other, a concept that is termed joint sparsity [73]. In an example model, each signal has a common component and a local innovation, with the commonality only needing to be encoded once while each innovation can be encoded at a lower measurement rate. Three different joint sparsity models (JSMs) have been developed:

1. Both common signal and innovations are sparse;
2. Sparse innovations with shared sparsity structure;
3. Sparse innovations and dense common signal.

Although JSM 1 would seem preferable due to the relatively limited amount of data, only JSM 2 is computationally feasible for large sensor networks; it has been used in many applications [73]. JSMs 1 and 3 can be solved using a linear program, which has cubic complexity on the number of sensors in the network.

DCS, however, does not address the communication or networking necessary to transmit the measurements to a central location; it relies on standard communication and networking techniques for measurement transmission, which can be tailored to the specific network topology.

5.11 Genomic sensing²⁴

Biosensing of pathogens is a research area of high consequence. An accurate and rapid biosensing paradigm has the potential to impact several fields, including healthcare, defense and environmental monitoring. In this module we address the concept of biosensing based on compressive sensing²⁵ (CS) via the *Compressive Sensing Microarray* (CSM), a DNA microarray adapted to take CS-style measurements.

DNA microarrays are a frequently applied solution for microbe sensing; they have a significant edge over competitors due to their ability to sense many organisms in parallel [113], [153]. A DNA microarray consists of genetic sensors or *spots*, each containing DNA sequences termed *probes*. From the perspective of a microarray, each DNA sequence can be viewed as a sequence of four DNA bases $\{A, T, G, C\}$ that tend to bind with one another in complementary pairs: *A* with *T* and *G* with *C*. Therefore, a DNA subsequence in a target organism's genetic sample will tend to bind or “hybridize” with its complementary subsequence on a microarray to form a stable structure. The target DNA sample to be identified is fluorescently tagged before it is flushed over the microarray. The extraneous DNA is washed away so that only the bound DNA is left on the array. The array is then scanned using laser light of a wavelength designed to trigger fluorescence in the spots where binding has occurred. A specific pattern of array spots will fluoresce, which is then used to infer the genetic makeup in the test sample.

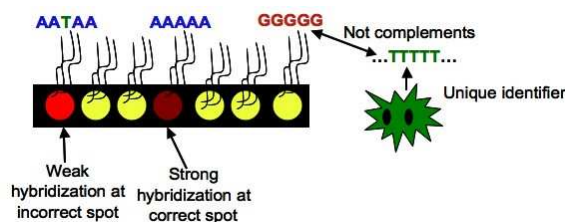


Figure 5.15: Cartoon of traditional DNA microarray showing strong and weak hybridization of the unique pathogen identifier at different microarray spots

There are three issues with the traditional microarray design. Each spot consists of probes that can uniquely identify only one target of interest (each spot contains multiple copies of a probe for robustness.) The first concern with this design is that very often the targets in a test sample have similar base sequences, causing them to hybridize with the wrong probe (see Figure 5.15). These cross-hybridization events lead to errors in the array readout. Current microarray design methods do not address cross-matches between similar DNA sequences.

The second concern in choosing unique identifier based DNA probes is its restriction on the number of organisms that can be identified. In typical biosensing applications multiple organisms must be identified; therefore a large number of DNA targets requires a microarray with a large number of spots. In fact, there are over 1000 known harmful microbes, many with more than 100 strains. The implementation cost and processing speed of microarray data is directly related to its number of spots, representing a significant problem for commercial deployment of microarray-based biosensors. As a consequence readout systems for traditional DNA arrays cannot be miniaturized or implemented using electronic components and require complicated fluorescent tagging.

The third concern is the inefficient utilization of the large number of array spots in traditional microarrays. Although the number of potential agents in a sample is very large, *not all agents* are expected to be present

²⁴This content is available online at <<http://cnx.org/content/m37374/1.3/>>.

²⁵"Introduction to compressive sensing" <<http://cnx.org/content/m37172/latest/>>

in a significant concentration at a given time and location, or in an air/water/soil sample to be tested. Therefore, in a traditionally designed microarray only a small fraction of spots will be active at a given time, corresponding to the few targets present.

To combat these problems, a Compressive Sensing DNA Microarray (CSM) uses “combinatorial testing sensors” in order to reduce the number of sensor spots [130], [155], [157]. Each spot in the CSM identifies a *group* of target organisms, and several spots together generate a unique pattern identifier for a single target. (See also "Group testing and data stream algorithms" (Section 5.3).) Designing the probes that perform this combinatorial sensing is the essence of the microarray design process, and what we aim to describe in this module.

To obtain a CS-type measurement scheme, we can choose each probe in a CSM to be a group identifier such that the readout of each probe is a probabilistic combination of all the targets in its group. The probabilities are representative of each probe’s hybridization affinity (or stickiness) to those targets in its group; the targets that are not in its group have low affinity to the probe. The readout signal at each spot of the microarray is a linear combination of hybridization affinities between its probe sequence and each of the target agents.

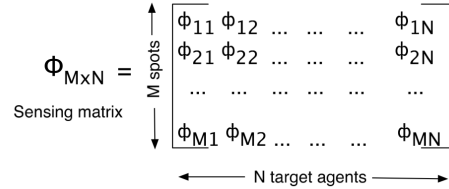


Figure 5.16: Structure of the CSM sensing matrix Φ with M spots identifying N targets

Figure 5.16 illustrates the sensing process. To formalize, we assume there are M spots on the CSM and N targets; we have far fewer spots than target agents. For $1 \leq i \leq M$ and $1 \leq j \leq N$, the probe at spot i hybridizes with target j with affinity $\phi_{i,j}$. The target j occurs in the tested DNA sample with concentration x_j , so that the total hybridization of spot i is $y_i = \sum_{j=1}^N \phi_{i,j} x_j = \phi_i x$, where ϕ_j and x are a row and column vector, respectively. The resulting measured microarray signal intensity vector $y = \{y_i\}_{i=1, \dots, M}$ fits the CS measurement model $y = \Phi x$.

While group testing has previously been proposed for microarrays [154], the sparsity in the target signal is key in applying CS. The chief advantage of a CS-based approach over regular group testing is in its information scalability. We are able to not just detect, but *estimate* the target signal with a reduced number of measurements similar to that of group testing [67]. This is important since there are always minute quantities of certain pathogens in the environment, but it is only their large concentrations that may be harmful to us. Furthermore, we are able to use CS recovery methods such as Belief Propagation (Section 4.10) that decode x while accounting for experimental noise and measurement nonlinearities due to excessive target molecules [157].

Glossary

A

A matrix Φ satisfies the *null space property* (NSP) of order K if there exists a constant $C > 0$ such that,

$$\|h_\Lambda\|_2 \leq C \frac{\|h_{\Lambda^c}\|_1}{\sqrt{K}} \quad (3.3)$$

holds for all $h \in \mathcal{N}(\Phi)$ and for all Λ such that $|\Lambda| \leq K$.

A matrix Φ satisfies the *restricted isometry property* (RIP) of order K if there exists a $\delta_K \in (0, 1)$ such that

$$(1 - \delta_K) \|x\|_2^2 \leq \|\Phi x\|_2^2 \leq (1 + \delta_K) \|x\|_2^2, \quad (3.9)$$

holds for all $x \in \Sigma_K = \{x : \|x\|_0 \leq K\}$.

A random variable X is called *strictly sub-Gaussian* if $X \sim \text{Sub}(\sigma^2)$ where $\sigma^2 = \mathbb{E}(X^2)$, i.e., the inequality

$$\mathbb{E}(\exp(Xt)) \leq \exp(\sigma^2 t^2 / 2) \quad (3.51)$$

holds for all $t \in \mathbb{R}$. To denote that X is strictly sub-Gaussian with variance σ^2 , we will use the notation $X \sim \text{SSub}(\sigma^2)$.

A random variable X is called *sub-Gaussian* if there exists a constant $c > 0$ such that

$$\mathbb{E}(\exp(Xt)) \leq \exp(c^2 t^2 / 2) \quad (3.48)$$

holds for all $t \in \mathbb{R}$. We use the notation $X \sim \text{Sub}(c^2)$ to denote that X satisfies ²⁶.

L

Let $\Phi : \mathbb{R}^N \rightarrow \mathbb{R}^M$ denote a sensing matrix and $\Delta : \mathbb{R}^M \rightarrow \mathbb{R}^N$ denote a recovery algorithm. We say that the pair (Φ, Δ) is *C-stable* if for any $x \in \Sigma_K$ and any $e \in \mathbb{R}^M$ we have that

$$\|\Delta(\Phi x + e) - x\|_2 \leq C \|e\|. \quad (3.11)$$

T

The coherence of a matrix Φ , $\mu(\Phi)$, is the largest absolute inner product between any two columns ϕ_i, ϕ_j of Φ :

$$\mu(\Phi) = \max_{1 \leq i < j \leq N} \frac{|\langle \phi_i, \phi_j \rangle|}{\|\phi_i\|_2 \|\phi_j\|_2}. \quad (3.45)$$

The spark of a given matrix Φ is the smallest number of columns of Φ that are linearly dependent.

²⁶<http://cnx.org/content/m37185/latest/>

Bibliography

- [1] D. Achlioptas. Database-friendly random projections. In *Proc. Symp. Principles of Database Systems (PODS)*, Santa Barbara, CA, May 2001.
- [2] W. Bajwa, J. Haupt, G. Raz, S. Wright, and R. Nowak. Toeplitz-structured compressed sensing matrices. In *Proc. IEEE Work. Stat. Signal Processing*, Madison, WI, Aug. 2007.
- [3] W. Bajwa, J. Haupt, G. Raz, S. Wright, and R. Nowak. Toeplitz-structured compressed sensing matrices. In *Proc. IEEE Work. Stat. Signal Processing*, Madison, WI, Aug. 2007.
- [4] W. Bajwa, J. Haupt, A. Sayeed, and R. Nowak. Compressive wireless sensing. In *Proc. Int. Symp. Inform. Processing in Sensor Networks (IPSN)*, Nashville, TN, Apr. 2006.
- [5] K. Ball. *An Elementary Introduction to Modern Convex Geometry*, volume 31 of *Flavors of Geometry*. MSRI Publ., Cambridge Univ. Press, Cambridge, England, 1997.
- [6] R. Baraniuk, M. Davenport, R. DeVore, and M. Wakin. A simple proof of the restricted isometry property for random matrices. *Const. Approx.*, 28(3):2538211;263, 2008.
- [7] R. Baraniuk, M. Davenport, R. DeVore, and M. Wakin. A simple proof of the restricted isometry property for random matrices. *Const. Approx.*, 28(3):2538211;263, 2008.
- [8] R. Baraniuk and M. Wakin. Random projections of smooth manifolds. *Found. Comput. Math.*, 9(1):518211;77, 2009.
- [9] R. Baraniuk and M. Wakin. Random projections of smooth manifolds. *Found. Comput. Math.*, 9(1):518211;77, 2009.
- [10] D. Baron, S. Sarvotham, and R. Baraniuk. Sudocodes - fast measurement and reconstruction of sparse signals. In *Proc. IEEE Int. Symp. Inform. Theory (ISIT)*, Seattle, WA, Jul. 2006.
- [11] J. Bect, L. Blanc-Feraud, G. Aubert, and A. Chambolle. A γ -unified variational framework for image restoration. In *Proc. European Conf. Comp. Vision (ECCV)*, Prague, Czech Republic, May 2004.
- [12] R. Berinde and P. Indyk. Sequential sparse matching pursuit. In *Communication, Control, and Computing, 2009. Allerton 2009. 47th Annual Allerton Conference on*, page 368211;43. IEEE, 2010.
- [13] R. Berinde, P. Indyk, and M. Ruzic. Practical near-optimal sparse recovery in the l_1 norm. In *Proc. Allerton Conf. Communication, Control, and Computing*, Monticello, IL, Sept. 2008.
- [14] A. Beurling. Sur les int $[U+FFFD]$ ales de fourier absolument convergentes et leur application $[U+FFFD]$ e transformation fonctionnelle. In *Proc. Scandinavian Math. Congress*, Helsinki, Finland, 1938.
- [15] T. Blumensath and M. Davies. Iterative hard thresholding for compressive sensing. *Appl. Comput. Harmon. Anal.*, 27(3):2658211;274, 2009.

- [16] P. Boufounos, M. Duarte, and R. Baraniuk. Sparse signal reconstruction from noisy compressive measurements using cross validation. In *Proc. IEEE Work. Stat. Signal Processing*, Madison, WI, Aug. 2007.
- [17] S. Boyd and L. Vanderberghe. *Convex Optimization*. Cambridge Univ. Press, Cambridge, England, 2004.
- [18] Y. Bresler. Spectrum-blind sampling and compressive sensing for continuous-index signals. In *Proc. Work. Inform. Theory and Applications (ITA)*, San Diego, CA, Jan. 2008.
- [19] Y. Bresler and P. Feng. Spectrum-blind minimum-rate sampling and reconstruction of 2-d multiband signals. In *Proc. IEEE Int. Conf. Image Processing (ICIP)*, Zurich, Switzerland, Sept. 1996.
- [20] V. Buldygin and Y. Kozachenko. *Metric Characterization of Random Variables and Random Processes*. American Mathematical Society, Providence, RI, 2000.
- [21] T. Cai and T. Jiang. Limiting laws of coherence of random matrices with applications to testing covariance structure and construction of compressed sensing matrices. Preprint, 2010.
- [22] E. Cand[**U+FFFD**] The restricted isometry property and its implications for compressed sensing. *Comptes rendus de l'Acad[**U+FFFD**]e des Sciences, S[**U+FFFD**]e I*, 346(9-10):5898211;592, 2008.
- [23] E. Cand[**U+FFFD**] The restricted isometry property and its implications for compressed sensing. *Comptes rendus de l'Acad[**U+FFFD**]e des Sciences, S[**U+FFFD**]e I*, 346(9-10):5898211;592, 2008.
- [24] E. Cand[**U+FFFD**] The restricted isometry property and its implications for compressed sensing. *Comptes rendus de l'Acad[**U+FFFD**]e des Sciences, S[**U+FFFD**]e I*, 346(9-10):5898211;592, 2008.
- [25] E. Cand[**U+FFFD**]nd Y. Plan. Near-ideal model selection by minimization. *Ann. Stat.*, 37(5A):21458211;2177, 2009.
- [26] E. Cand[**U+FFFD**]nd J. Romberg. Sparsity and incoherence in compressive sampling. *Inverse Problems*, 23(3):9698211;985, 2007.
- [27] E. Cand[**U+FFFD**] J. Romberg, and T. Tao. Stable signal recovery from incomplete and inaccurate measurements. *Comm. Pure Appl. Math.*, 59(8):12078211;1223, 2006.
- [28] E. Cand[**U+FFFD**] M. Rudelson, T. Tao, and R. Vershynin. Error correction via linear programming. In *Proc. IEEE Symp. Found. Comp. Science (FOCS)*, Pittsburg, PA, Oct. 2005.
- [29] E. Cand[**U+FFFD**]nd T. Tao. Decoding by linear programming. *IEEE Trans. Inform. Theory*, 51(12):42038211;4215, 2005.
- [30] E. Cand[**U+FFFD**]nd T. Tao. Near optimal signal recovery from random projections: Universal encoding strategies? *IEEE Trans. Inform. Theory*, 52(12):54068211;5425, 2006.
- [31] E. Cand[**U+FFFD**]nd T. Tao. The dantzig selector: Statistical estimation when is much larger than. *Ann. Stat.*, 35(6):23138211;2351, 2007.
- [32] M. Charikar, K. Chen, and M. Farach-Colton. Finding frequent items in data streams. In *Proc. Int. Coll. Autom. Lang. Programm.*, M[**U+FFFD**]ga, Spain, Jul. 2002.
- [33] S. Chen, D. Donoho, and M. Saunders. Atomic decomposition by basis pursuit. *SIAM J. Sci. Comp.*, 20(1):338211;61, 1998.
- [34] A. Cohen, W. Dahmen, and R. DeVore. Instance optimal decoding by thresholding in compressed sensing. In *Int. Conf. Harmonic Analysis and Partial Differential Equations*, Madrid, Spain, Jun. 2008.

- [35] A. Cohen, W. Dahmen, and R. DeVore. Compressed sensing and best k -term approximation. *J. Amer. Math. Soc.*, 22(1):2118211;231, 2009.
- [36] A. Cohen, W. Dahmen, and R. DeVore. Compressed sensing and best k -term approximation. *J. Amer. Math. Soc.*, 22(1):2118211;231, 2009.
- [37] R. Coifman, F. Geshwind, and Y. Meyer. Noiselets. *Appl. Comput. Harmon. Anal.*, 10:278211;44, 2001.
- [38] P. Combettes and J. Pesquet. Proximal thresholding algorithm for minimization over orthonormal bases. *SIAM J. Optimization*, 18(4):13518211;1376, 2007.
- [39] G. Cormode and M. Hadjieleftheriou. Finding the frequent items in streams of data. *Comm. ACM*, 52(10):978211;105, 2009.
- [40] G. Cormode and M. Hadjieleftheriou. Finding the frequent items in streams of data. *Comm. ACM*, 52(10):978211;105, 2009.
- [41] G. Cormode and S. Muthukrishnan. Improved data stream summaries: The count-min sketch and its applications. *J. Algorithms*, 55(1):588211;75, 2005.
- [42] W. Dai and O. Milenkovic. Subspace pursuit for compressive sensing signal reconstruction. *IEEE Trans. Inform. Theory*, 55(5):22308211;2249, 2009.
- [43] S. Dasgupta and A. Gupta. An elementary proof of the johnson-lindenstrauss lemma. Technical report TR-99-006, Univ. of Cal. Berkeley, Comput. Science Division, Mar. 1999.
- [44] I. Daubechies, M. Defrise, and C. De Mol. An iterative thresholding algorithm for linear inverse problems with a sparsity constraint. *Comm. Pure Appl. Math.*, 57(11):14138211;1457, 2004.
- [45] M. Davenport, P. Boufounos, M. Wakin, and R. Baraniuk. Signal processing with compressive measurements. *IEEE J. Select. Top. Signal Processing*, 4(2):4458211;460, 2010.
- [46] M. Davenport, M. Duarte, M. Wakin, J. Laska, D. Takhar, K. Kelly, and R. Baraniuk. The smashed filter for compressive classification and target recognition. In *Proc. IS&T/SPIE Symp. Elec. Imag.: Comp. Imag.*, San Jose, CA, Jan. 2007.
- [47] M. Davenport, C. Hegde, M. Duarte, and R. Baraniuk. Joint manifolds for data fusion. *IEEE Trans. Image Processing*, 19(10):25808211;2594, 2010.
- [48] M. Davenport, J. Laska, P. Boufouons, and R. Baraniuk. A simple proof that random matrices are democratic. Technical report TREE 0906, Rice Univ., ECE Dept., Nov. 2009.
- [49] M. Davenport and M. Wakin. Analysis of orthogonal matching pursuit using the restricted isometry property. *IEEE Trans. Inform. Theory*, 56(9):43958211;4401, 2010.
- [50] R. DeVore. Nonlinear approximation. *Acta Numerica*, 7:518211;150, 1998.
- [51] R. DeVore. Nonlinear approximation. *Acta Numerica*, 7:518211;150, 1998.
- [52] R. DeVore. Deterministic constructions of compressed sensing matrices. *J. Complex.*, 23(4):9188211;925, 2007.
- [53] A. Dimakis, A. Sarwate, and M. Wainwright. Geographic gossip: Efficient aggregation for sensor networks. In *Proc. Int. Symp. Inform. Processing in Sensor Networks (IPSN)*, Nashville, TN, Apr. 2006.
- [54] D. Donoho. Denoising by soft-thresholding. *IEEE Trans. Inform. Theory*, 41(3):6138211;627, 1995.

- [55] D. Donoho. Neighborly polytopes and sparse solutions of underdetermined linear equations. Technical report 2005-04, Stanford Univ., Stat. Dept., Jan. 2005.
- [56] D. Donoho. For most large underdetermined systems of linear equations, the minimal ℓ_1 -norm solution is also the sparsest solution. *Comm. Pure Appl. Math.*, 59(6):7978211;829, 2006.
- [57] D. Donoho. High-dimensional centrally symmetric polytopes with neighborliness proportional to dimension. *Discrete and Comput. Geometry*, 35(4):6178211;652, 2006.
- [58] D. Donoho, I. Drori, Y. Tsaig, and J.-L. Stark. Sparse solution of underdetermined linear equations by stagewise orthogonal matching pursuit. Preprint, 2006.
- [59] D. Donoho and M. Elad. Optimally sparse representation in general (nonorthogonal) dictionaries via ℓ_1 minimization. *Proc. Natl. Acad. Sci.*, 100(5):21978211;2202, 2003.
- [60] D. Donoho and M. Elad. Optimally sparse representation in general (nonorthogonal) dictionaries via minimization. *Proc. Natl. Acad. Sci.*, 100(5):21978211;2202, 2003.
- [61] D. Donoho and C. Grimes. Hessian eigenmaps: Locally linear embedding techniques for high-dimensional data. *Proc. Natl. Acad. Sci.*, 100(10):55918211;5596, 2003.
- [62] D. Donoho, A. Maleki, and A. Montanari. Message passing algorithms for compressed sensing. *Proc. Natl. Acad. Sci.*, 106(45):189148211;18919, 2009.
- [63] D. Donoho and J. Tanner. Neighborliness of randomly projected simplices in high dimensions. *Proc. Natl. Acad. Sci.*, 102(27):94528211;9457, 2005.
- [64] D. Donoho and J. Tanner. Sparse nonnegative solutions of undetermined linear equations by linear programming. *Proc. Natl. Acad. Sci.*, 102(27):94468211;9451, 2005.
- [65] D. Donoho and J. Tanner. Counting faces of randomly-projected polytopes when the projection radically lowers dimension. *J. Amer. Math. Soc.*, 22(1):18211;53, 2009.
- [66] D. Donoho and J. Tanner. Precise undersampling theorems. *Proc. IEEE*, 98(6):9138211;924, 2010.
- [67] D.-Z. Du and F. Hwang. *Combinatorial group testing and its applications*. World Scientific Publishing Co., Singapore, 2000.
- [68] M. Duarte and R. Baraniuk. Kronecker compressive sensing. Preprint, 2009.
- [69] M. Duarte, M. Davenport, D. Takhar, J. Laska, T. Sun, K. Kelly, and R. Baraniuk. Single-pixel imaging via compressive sampling. *IEEE Signal Processing Mag.*, 25(2):838211;91, 2008.
- [70] M. Duarte, M. Davenport, M. Wakin, and R. Baraniuk. Sparse signal detection from incoherent projections. In *Proc. IEEE Int. Conf. Acoust., Speech, and Signal Processing (ICASSP)*, Toulouse, France, May 2006.
- [71] M. Duarte, M. Davenport, M. Wakin, J. Laska, D. Takhar, K. Kelly, and R. Baraniuk. Multiscale random projections for compressive classification. In *Proc. IEEE Int. Conf. Image Processing (ICIP)*, San Antonio, TX, Sept. 2007.
- [72] M. Duarte, M. Wakin, and R. Baraniuk. Fast reconstruction of piecewise smooth signals from random projections. In *Proc. Work. Struc. Parc. Rep. Adap. Signaux (SPARS)*, Rennes, France, Nov. 2005.
- [73] M. F. Duarte, M. B. Wakin, D. Baron, and R. G. Baraniuk. Universal distributed sensing via random projections. In *Int. Workshop on Inform. Processing in Sensor Networks (IPSN)*, page 1778211;185, Nashville, TN, Apr. 2006.

- [74] M. Elad. Why simple shrinkage is still relevant for redundant representations? *IEEE Trans. Inform. Theory*, 52(12):55598211;5569, 2006.
- [75] M. Elad. *Sparse and Redundant Representations: From Theory to Applications in Signal and Image Processing*. Springer, New York, NY, 2010.
- [76] M. Elad, B. Matalon, J. Shtok, and M. Zibulevsky. A wide-angle view at iterated shrinkage algorithms. In *Proc. SPIE Optics Photonics: Wavelets*, San Diego, CA, Apr. 2007.
- [77] M. Elad, B. Matalon, and M. Zibulevsky. Coordinate and subspace optimization methods for linear least squares with non-quadratic regularization. *Appl. Comput. Harmon. Anal.*, 23(3):3468211;367, 2007.
- [78] Y. Erlich, N. Shental, A. Amir, and O. Zuk. Compressed sensing approach for high throughput carrier screen. In *Proc. Allerton Conf. Communication, Control, and Computing*, Monticello, IL, Sept. 2009.
- [79] Y. Erlich, N. Shental, A. Amir, and O. Zuk. Compressed sensing approach for high throughput carrier screen. In *Proc. Allerton Conf. Communication, Control, and Computing*, Monticello, IL, Sept. 2009.
- [80] V. Fedorov. *Theory of Optimal Experiments*. Academic Press, New York, NY, 1972.
- [81] P. Feng. *Universal spectrum blind minimum rate sampling and reconstruction of multiband signals*. Ph. d. thesis, University of Illinois at Urbana-Champaign, Mar. 1997.
- [82] P. Feng and Y. Bresler. Spectrum-blind minimum-rate sampling and reconstruction of multiband signals. In *Proc. IEEE Int. Conf. Acoust., Speech, and Signal Processing (ICASSP)*, Atlanta, GA, May 1996.
- [83] M. Figueiredo and R. Nowak. An em algorithm for wavelet-based image restoration. *IEEE Trans. Image Processing*, 12(8):9068211;916, 2003.
- [84] M. Figueiredo, R. Nowak, and S. Wright. Gradient projections for sparse reconstruction: Application to compressed sensing and other inverse problems. *IEEE J. Select. Top. Signal Processing*, 1(4):5868211;597, 2007.
- [85] A. Garnaev and E. Gluskin. The widths of euclidean balls. *Dokl. An. SSSR*, 277:10488211;1052, 1984.
- [86] M. Gehm, R. John, D. Brady, R. Willett, and T. Schultz. Single-shot compressive spectral imaging with a dual disperser architecture. *Optics Express*, 15(21):140138211;14027, 2007.
- [87] S. Gervsgorin. [U+FFFD]er die abgrenzung der eigenwerte einer matrix. *Izv. Akad. Nauk SSSR Ser. Fiz.-Mat.*, 6:7498211;754, 1931.
- [88] A. Gilbert, S. Guha, P. Indyk, S. Muthukrishnan, and M. Strauss. Near-optimal sparse fourier representations via sampling. In *Proc. ACM Symp. Theory of Comput.*, Montreal, Canada, May 2002.
- [89] A. Gilbert and P. Indyk. Sparse recovery using sparse matrices. *Proc. IEEE*, 98(6):9378211;947, 2010.
- [90] A. Gilbert, Y. Li, E. Porat, and M. Strauss. Approximate sparse recovery: Optimizaing time and measurements. In *Proc. ACM Symp. Theory of Comput.*, Cambridge, MA, Jun. 2010.
- [91] A. Gilbert, S. Muthukrishnan, and M. Strauss. Improved time bounds for near-optimal sparse fourier representations. In *Proc. SPIE Optics Photonics: Wavelets*, San Diego, CA, Aug. 2005.
- [92] A. Gilbert, M. Strauss, J. Tropp, and R. Vershynin. One sketch for all: Fast algorithms for compressed sensing. In *Proc. ACM Symp. Theory of Comput.*, San Diego, CA, Jun. 2007.

- [93] AC Gilbert, MJ Strauss, JA Tropp, and R. Vershynin. One sketch for all: fast algorithms for compressed sensing. In *Proceedings of the thirty-ninth annual ACM symposium on Theory of computing*, page 2378211;246. ACM, 2007.
- [94] I. Gorodnitsky, J. George, and B. Rao. Neuromagnetic source imaging with focuss: A recursive weighted minimum norm algorithm. *Electroencephalography and Clinical Neurophysiology*, 95(4):2318211;251, 1995.
- [95] I. Gorodnitsky and B. Rao. Convergence analysis of a class of adaptive weighted norm extrapolation algorithms. In *Proc. Asilomar Conf. Signals, Systems, and Computers*, Pacific Grove, CA, Nov. 1993.
- [96] E. Hale, W. Yin, and Y. Zhang. A fixed-point continuation method for ℓ_1 -regularized minimization with applications to compressed sensing. Technical report TR07-07, Rice Univ., CAAM Dept., 2007.
- [97] T. Hastie, R. Tibshirani, and J. Friedman. *The Elements of Statistical Learning*. Springer, New York, NY, 2001.
- [98] J. Haupt and R. Nowak. Signal reconstruction from noisy random projections. *IEEE Trans. Inform. Theory*, 52(9):40368211;4048, 2006.
- [99] C. Hegde, M. Wakin, and R. Baraniuk. Random projections for manifold learning. In *Proc. Adv. in Neural Processing Systems (NIPS)*, Vancouver, BC, Dec. 2007.
- [100] M. Herman and T. Strohmer. High-resolution radar via compressed sensing. *IEEE Trans. Signal Processing*, 57(6):22758211;2284, 2009.
- [101] P. Indyk. Explicit constructions for compressed sensing of sparse signals. In *Proc. ACM-SIAM Symp. on Discrete Algorithms (SODA)*, page 308211;33, Jan. 2008.
- [102] P. Indyk and M. Ruzic. Near-optimal sparse recovery in the ℓ_1 norm. In *Proc. IEEE Symp. Found. Comp. Science (FOCS)*, Philadelphia, PA, Oct. 2008.
- [103] S. Jafarpour, W. Xu, B. Hassibi, and R. Calderbank. Efficient and robust compressed sensing using optimized expander graphs. *IEEE Trans. Inform. Theory*, 55:42998211;4308, Sep. 2009.
- [104] T. Jayram and D. Woodruff. Optimal bounds for johnson-lindenstrauss transforms and streaming problems with sub-constant error. In *Proc. ACM-SIAM Symp. on Discrete Algorithms (SODA)*, San Francisco, CA, Jan. 2011.
- [105] S. Ji, Y. Xue, and L. Carin. Bayesian compressive sensing. *IEEE Trans. Signal Processing*, 56(6):23468211;2356, 2008.
- [106] W. Johnson and J. Lindenstrauss. Extensions of lipschitz mappings into a hilbert space. In *Proc. Conf. Modern Anal. and Prob.*, New Haven, CT, Jun. 1982.
- [107] R. Kainkaryam, A. Breux, A. Gilbert, P. Woolf, and J. Schiefelbein. poolmc: Smart pooling of mrna samples in microarray experiments. *BMC Bioinformatics*, 11(1):299, 2010.
- [108] R. Kainkaryam, A. Breux, A. Gilbert, P. Woolf, and J. Schiefelbein. poolmc: Smart pooling of mrna samples in microarray experiments. *BMC Bioinformatics*, 11(1):299, 2010.
- [109] K. Kelly, D. Takhar, T. Sun, J. Laska, M. Duarte, and R. Baraniuk. Compressed sensing for multispectral and confocal microscopy. In *Proc. American Physical Society Meeting*, Denver, CO, Mar. 2007.
- [110] S.-J. Kim, K. Koh, M. Lustig, S. Boyd, and D. Gorinevsky. An interior point method for large-scale ℓ_1 -regularized least squares. *IEEE J. Select. Top. Signal Processing*, 1(4):6068211;617, 2007.

- [111] S. Kirolos, J. Laska, M. Wakin, M. Duarte, D. Baron, T. Ragheb, Y. Massoud, and R. Baraniuk. Analog-to-information conversion via random demodulation. In *Proc. IEEE Dallas Circuits and Systems Work. (DCAS)*, Dallas, TX, Oct. 2006.
- [112] F. Krahmer and R. Ward. New and improved johnson-lindenstrauss embeddings via the restricted isometry property. Preprint, Sept. 2010.
- [113] E. Lander. Array of hope. *Nature Genetics*, 21:38211;4, 1999.
- [114] J. Laska, P. Boufounos, M. Davenport, and R. Baraniuk. Democracy in action: Quantization, saturation, and compressive sensing. Preprint, 2009.
- [115] J. Laska, S. Kirolos, M. Duarte, T. Ragheb, R. Baraniuk, and Y. Massoud. Theory and implementation of an analog-to-information convertor using random demodulation. In *Proc. IEEE Int. Symposium on Circuits and Systems (ISCAS)*, New Orleans, LA, May 2007.
- [116] M. Ledoux. *The Concentration of Measure Phenomenon*. American Mathematical Society, Providence, RI, 2001.
- [117] S. Levy and P. Fullagar. Reconstruction of a sparse spike train from a portion of its spectrum and application to high-resolution deconvolution. *Geophysics*, 46(9):12358211;1243, 1981.
- [118] B. Logan. *Properties of High-Pass Signals*. Ph. d. thesis, Columbia University, 1965.
- [119] M. Lustig, D. Donoho, and J. Pauly. Rapid mr imaging with compressed sensing and randomly under-sampled 3dft trajectories. In *Proc. Annual Meeting of ISMRM*, Seattle, WA, May 2006.
- [120] M. Lustig, J. Lee, D. Donoho, and J. Pauly. Faster imaging with randomly perturbed, under-sampled spirals and reconstruction. In *Proc. Annual Meeting of ISMRM*, Miami, FL, May 2005.
- [121] M. Lustig, J. Lee, D. Donoho, and J. Pauly. k-t sparse: High frame rate dynamic mri exploiting spatio-temporal sparsity. In *Proc. Annual Meeting of ISMRM*, Seattle, WA, May 2006.
- [122] M. Lustig, J. Santos, J. Lee, D. Donoho, and J. Pauly. Application of compressed sensing for rapid mr imaging. In *Proc. Work. Struc. Parc. Rep. Adap. Signaux (SPARS)*, Rennes, France, Nov. 2005.
- [123] D. MacKay. Information-based objective functions for active data selection. *Neural Comput.*, 4:5908211;604, 1992.
- [124] S. Mallat. *A Wavelet Tour of Signal Processing*. Academic Press, San Diego, CA, 1999.
- [125] S. Mallat. *A Wavelet Tour of Signal Processing*. Academic Press, San Diego, CA, 1999.
- [126] S. Mallat. *A Wavelet Tour of Signal Processing*. Academic Press, San Diego, CA, 1999.
- [127] S. Mallat and Z. Zhang. Matching pursuits with time-frequency dictionaries. *IEEE Trans. Signal Processing*, 41(12):33978211;3415, 1993.
- [128] R. Marcia, Z. Harmany, and R. Willett. Compressive coded aperture imaging. In *Proc. IS&T/SPIE Symp. Elec. Imag.: Comp. Imag.*, San Jose, CA, Jan. 2009.
- [129] S. Mendelson, A. Pajor, and N. Tomczack-Jaegermann. Uniform uncertainty principle for bernoulli and subgaussian ensembles. *Const. Approx.*, 28(3):2778211;289, 2008.
- [130] O. Milenkovic, R. Baraniuk, and T. Simunic-Rosing. Compressed sensing meets bioinformatics: A novel dna microarray design. In *Proc. Work. Inform. Theory and Applications (ITA)*, San Diego, CA, Jan. 2007.

- [131] M. Mishali and Y. C. Eldar. Blind multi-band signal reconstruction: Compressed sensing for analog signals. *IEEE Trans. Signal Processing*, 57(3):9938211;1009, 2009.
- [132] M. Mishali and Y. C. Eldar. From theory to practice: Sub-nyquist sampling of sparse wideband analog signals. *IEEE J. Select. Top. Signal Processing*, 4(2):3758211;391, 2010.
- [133] C. De Mol and M. Defrise. A note on wavelet-based inversion algorithms. *Contemporary Math.*, 313:858211;96, 2002.
- [134] S. Muthukrishnan. *Data Streams: Algorithms and Applications*, volume 1 of *Found. Trends in Theoretical Comput. Science*. Now Publishers, Boston, MA, 2005.
- [135] S. Muthukrishnan. *Data Streams: Algorithms and Applications*, volume 1 of *Found. Trends in Theoretical Comput. Science*. Now Publishers, Boston, MA, 2005.
- [136] D. Needell and J. Tropp. Cosamp: Iterative signal recovery from incomplete and inaccurate samples. *Appl. Comput. Harmon. Anal.*, 26(3):3018211;321, 2009.
- [137] D. Needell and J. Tropp. Cosamp: Iterative signal recovery from incomplete and inaccurate samples. *Appl. Comput. Harmon. Anal.*, 26(3):3018211;321, 2009.
- [138] D. Needell and J. Tropp. Cosamp: Iterative signal recovery from incomplete and inaccurate samples. *Appl. Comput. Harmon. Anal.*, 26(3):3018211;321, 2009.
- [139] D. Needell and R. Vershynin. Uniform uncertainty principle and signal recovery via regularized orthogonal matching pursuit. *Found. Comput. Math.*, 9(3):3178211;334, 2009.
- [140] D. Needell and R. Vershynin. Signal recovery from incomplete and inaccurate measurements via regularized orthogonal matching pursuit. *IEEE J. Select. Top. Signal Processing*, 4(2):3108211;316, 2010.
- [141] R. Nowak and M. Figueiredo. Fast wavelet-based image deconvolution using the em algorithm. In *Proc. Asilomar Conf. Signals, Systems, and Computers*, Pacific Grove, CA, Nov. 2001.
- [142] B. Olshausen and D. Field. Emergence of simple-cell receptive field properties by learning a sparse representation. *Nature*, 381:6078211;609, 1996.
- [143] S. Osher, M. Burger, D. Goldfarb, J. Xu, and W. Yin. An iterative regularization method for total variation-based image restoration. *SIAM J. Multiscale Modeling and Simulation*, 4(2):4608211;489, 2005.
- [144] Y. Pati, R. Rezaifar, and P. Krishnaprasad. Orthogonal matching pursuit: Recursive function approximation with applications to wavelet decomposition. In *Proc. Asilomar Conf. Signals, Systems, and Computers*, Pacific Grove, CA, Nov. 1993.
- [145] W. Pennebaker and J. Mitchell. *JPEG Still Image Data Compression Standard*. Van Nostrand Reinhold, 1993.
- [146] M. Rabbat, J. Haupt, A. Singh, and R. Nowak. Decentralized compression and predistribution via randomized gossiping. In *Proc. Int. Symp. Inform. Processing in Sensor Networks (IPSN)*, Nashville, TN, Apr. 2006.
- [147] R. Robucci, L. Chiu, J. Gray, J. Romberg, P. Hasler, and D. Anderson. Compressive sensing on a cmos separable transform image sensor. In *Proc. IEEE Int. Conf. Acoust., Speech, and Signal Processing (ICASSP)*, Las Vegas, NV, Apr. 2008.
- [148] J. Romberg. Compressive sensing by random convolution. *SIAM J. Imag. Sci.*, 2(4):10988211;1128, 2009.

- [149] J. Romberg. Compressive sensing by random convolution. *SIAM J. Imag. Sci.*, 2(4):10988211;1128, 2009.
- [150] M. Rosenfeld. *The Mathematics of Paul Erdős II*, chapter In praise of the Gram matrix, page 3188211;323. Springer, Berlin, Germany, 1996.
- [151] S. Roweis and L. Saul. Nonlinear dimensionality reduction by locally linear embedding. *Science*, 290(5500):23238211;2326, 2000.
- [152] S. Sarvotham, D. Baron, and R. Baraniuk. Compressed sensing reconstruction via belief propagation. Technical report TREE-0601, Rice Univ., ECE Dept., 2006.
- [153] M. Schena, D. Shalon, R. Davis, and P. Brown. Quantitative monitoring of gene expression patterns with a complementary dna microarray. *Science*, 270(5235):4678211;470, 1995.
- [154] A. Schliep, D. Torney, and S. Rahmann. Group testing with dna chips: Generating designs and decoding experiments. In *Proc. Conf. Comput. Systems Bioinformatics*, Stanford, CA, Aug. 2003.
- [155] M. Sheikh, O. Milenkovic, S. Sarvotham, and R. Baraniuk. Compressed sensing dna microarrays. Technical report TREE-0706, Rice Univ., ECE Dept., May 2007.
- [156] M. Sheikh, S. Sarvotham, O. Milenkovic, and R. Baraniuk. Dna array decoding from nonlinear measurements by belief propagation. In *Proc. IEEE Work. Stat. Signal Processing*, Madison, WI, Aug. 2007.
- [157] M. Sheikh, S. Sarvotham, O. Milenkovic, and R. Baraniuk. Dna array decoding from nonlinear measurements by belief propagation. In *Proc. IEEE Work. Stat. Signal Processing*, Madison, WI, Aug. 2007.
- [158] N. Shental, A. Amir, and O. Zuk. Identification of rare alleles and their carriers using compressed sequencing. *Nucleic Acids Research*, 38(19):e179, 2009.
- [159] N. Shental, A. Amir, and O. Zuk. Identification of rare alleles and their carriers using compressed sequencing. *Nucleic Acids Research*, 38(19):e179, 2009.
- [160] J. P. Slavinsky, J. Laska, M. Davenport, and R. Baraniuk. The compressive multiplexer for multi-channel compressive sensing. In *Proc. IEEE Int. Conf. Acoust., Speech, and Signal Processing (ICASSP)*, Prague, Czech Republic, May 2011.
- [161] T. Strohmer and R. Heath. Grassmanian frames with applications to coding and communication. *Appl. Comput. Harmon. Anal.*, 14(3):2578211;275, Nov. 2003.
- [162] D. Takhar, J. Laska, M. Wakin, M. Duarte, D. Baron, K. Kelly, and R. Baraniuk. A compressed sensing camera: New theory and an implementation using digital micromirrors. In *Proc. IS&T/SPIE Symp. Elec. Imag.: Comp. Imag.*, San Jose, CA, Jan. 2006.
- [163] D. Takhar, J. Laska, M. Wakin, M. Duarte, D. Baron, S. Sarvotham, K. Kelly, and R. Baraniuk. A new compressive imaging camera architecture using optical-domain compression. In *Proc. IS&T/SPIE Symp. Elec. Imag.: Comp. Imag.*, San Jose, CA, Jan. 2006.
- [164] D. Taubman and M. Marcellin. *JPEG 2000: Image Compression Fundamentals, Standards and Practice*. Kluwer, 2001.
- [165] H. Taylor, S. Banks, and J. McCoy. Deconvolution with the norm. *Geophysics*, 44(1):398211;52, 1979.
- [166] J. Tenenbaum, V. de Silva, and J. Landford. A global geometric framework for nonlinear dimensionality reduction. *Science*, 290:23198211;2323, 2000.

- [167] R. Tibshirani. Regression shrinkage and selection via the lasso. *J. Royal Statist. Soc B*, 58(1):2678211;288, 1996.
- [168] R. Tibshirani. Regression shrinkage and selection via the lasso. *J. Royal Statist. Soc B*, 58(1):2678211;288, 1996.
- [169] M. Tipping. Sparse bayesian learning and the relevance vector machine. *J. Machine Learning Research*, 1:2118211;244, 2001.
- [170] M. Tipping and A. Faul. Fast marginal likelihood maximization for sparse bayesian models. In *Proc. Int. Conf. Art. Intell. Stat. (AISTATS)*, Key West, FL, Jan. 2003.
- [171] J. Tropp and A. Gilbert. Signal recovery from partial information via orthogonal matching pursuit. *IEEE Trans. Inform. Theory*, 53(12):46558211;4666, 2007.
- [172] J. Tropp and A. Gilbert. Signal recovery from partial information via orthogonal matching pursuit. *IEEE Trans. Inform. Theory*, 53(12):46558211;4666, 2007.
- [173] J. Tropp, J. Laska, M. Duarte, J. Romberg, and R. Baraniuk. Beyond nyquist: Efficient sampling of sparse, bandlimited signals. *IEEE Trans. Inform. Theory*, 56(1):5208211;544, 2010.
- [174] J. Tropp, J. Laska, M. Duarte, J. Romberg, and R. Baraniuk. Beyond nyquist: Efficient sampling of sparse, bandlimited signals. *IEEE Trans. Inform. Theory*, 56(1):5208211;544, 2010.
- [175] J. Tropp, M. Wakin, M. Duarte, D. Baron, and R. Baraniuk. Random filters for compressive sampling and reconstruction. In *Proc. IEEE Int. Conf. Acoust., Speech, and Signal Processing (ICASSP)*, Toulouse, France, May 2006.
- [176] J. Tropp, M. Wakin, M. Duarte, D. Baron, and R. Baraniuk. Random filters for compressive sampling and reconstruction. In *Proc. IEEE Int. Conf. Acoust., Speech, and Signal Processing (ICASSP)*, Toulouse, France, May 2006.
- [177] J. A. Tropp. Norms of random submatrices and sparse approximation. *C. R. Acad. Sci. Paris, Ser. I*, 346(238211;24):12718211;1274, 2008.
- [178] J. A. Tropp. On the conditioning of random subdictionaries. *Appl. Comput. Harmon. Anal.*, 25:1–24, 2008.
- [179] J. Trzasko and A. Manduca. Highly undersampled magnetic resonance image reconstruction via homotopic -minimization. *IEEE Trans. Med. Imaging*, 28(1):1068211;121, 2009.
- [180] V. Vapnik. *The Nature of Statistical Learning Theory*. Springer-Verlag, New York, NY, 1999.
- [181] R. Varga. *Ger353;gorin and His Circles*. Springer, Berlin, Germany, 2004.
- [182] R. Venkataramani and Y. Bresler. Further results on spectrum blind sampling of 2-d signals. In *Proc. IEEE Int. Conf. Image Processing (ICIP)*, Chicago, IL, Oct. 1998.
- [183] A. Wagadarikar, R. John, R. Willett, and D. Brady. Single disperser design for coded aperture snapshot spectral imaging. *Appl. Optics*, 47(10):B448211;51, 2008.
- [184] M. Wakin, J. Laska, M. Duarte, D. Baron, S. Sarvotham, D. Takhar, K. Kelly, and R. Baraniuk. An architecture for compressive imaging. In *Proc. IEEE Int. Conf. Image Processing (ICIP)*, Atlanta, GA, Oct. 2006.
- [185] M. Wakin, J. Laska, M. Duarte, D. Baron, S. Sarvotham, D. Takhar, K. Kelly, and R. Baraniuk. Compressive imaging for video representation and coding. In *Proc. Picture Coding Symp.*, Beijing, China, Apr. 2006.

- [186] C. Walker and T. Ulrych. Autoregressive recovery of the acoustic impedance. *Geophysics*, 48(10):1338-1350, 1983.
- [187] W. Wang, M. Garofalakis, and K. Ramchandran. Distributed sparse random projections for refinable approximation. In *Proc. Int. Symp. Inform. Processing in Sensor Networks (IPSN)*, Cambridge, MA, Apr. 2007.
- [188] R. Ward. Compressive sensing with cross validation. *IEEE Trans. Inform. Theory*, 55(12):5773-5782, 2009.
- [189] L. Welch. Lower bounds on the maximum cross correlation of signals. *IEEE Trans. Inform. Theory*, 20(3):397-399, 1974.
- [190] P. Wojtaszczyk. Stability and instance optimality for gaussian measurements in compressed sensing. *Found. Comput. Math.*, 10(1):18-33, 2010.
- [191] S. Wright, R. Nowak, and M. Figueiredo. Sparse reconstruction by separable approximation. *IEEE Trans. Signal Processing*, 57(7):2479-2493, 2009.
- [192] W. Yin, S. Osher, D. Goldfarb, and J. Darbon. Bregman iterative algorithms for ℓ_1 -minimization with applications to compressed sensing. *SIAM J. Imag. Sci.*, 1(1):143-168, 2008.

Index of Keywords and Terms

Keywords are listed by the section with that keyword (page numbers are in parentheses). Keywords do not necessarily appear in the text of the page. They are merely associated with that section. *Ex.* apples, § 1.1 (1) **Terms** are referenced by the page they appear on. *Ex.* apples, 1

- A** Analog to digital converter (ADC), § 5.5(69)
 Analysis, § 2.2(13)
 Approximation, § 2.4(17)
 Atoms, § 2.2(13)
- B** Basis, § 2.2(13)
 Bayesian methods, § 4.10(64)
 Belief propagation, § 4.10(64), § 5.11(90)
 Best K-term approximation, § 2.4(17)
 Biosensing, § 5.11(90)
 Bounded noise, § 4.3(47)
 Bregman iteration, § 4.7(55)
- C** Classification, § 5.9(85)
 Coherence, § 3.6(32)
 Combinatorial algorithms, § 4.9(62)
 Compressibility, § 2.3(14), § 2.4(17)
 Compressive imaging, § 5.6(72)
 Compressive multiplexer, § 3.5(30)
 Compressive sensing, § 3.1(21)
 Compressive signal processing, § 5.9(85)
 Concentration inequalities, § 3.8(35)
 Concentration of measure, § 3.5(30), § 3.8(35), § 3.9(38)
 Convex optimization, § 4.7(55)
 Convex relaxation, § 4.1(43)
 CoSaMP, § 4.8(58)
 Count-median, § 4.9(62)
 Count-min, § 4.9(62)
 Cross-polytope, § 4.5(53)
- D** Dantzig selector, § 4.3(47)
 Data streams, § 4.9(62), § 5.3(68)
 Democracy, § 3.5(30)
 Detection, § 5.9(85)
 Deterministic constructions, § 3.5(30)
 Dictionary, § 2.2(13)
 Dimensionality reduction, § 3.1(21)
 Distributed compressive sensing, § 5.10(88)
 DNA microarray, § 5.11(90)
 Dual frame, § 2.2(13)
- E** Electroencephalography (EEG)
 Magnetoencephalography (MEG), § 5.4(69)
 ell_p norms, § 2.1(11)
 Error correcting codes, § 4.10(64), § 5.2(67)
 Estimation, § 5.9(85)
 Explicit constructions, § 3.5(30)
- F** Frame, § 2.2(13)
- G** Gaussian noise, § 4.3(47)
 Gershgorin circle theorem, § 3.6(32)
 Greedy algorithms, § 4.8(58)
 Group testing, § 4.9(62), § 5.3(68), § 5.11(90)
- H** Hilbert space, § 2.2(13)
 Hyperspectral imaging, § 5.7(76)
- I** Inner products, § 2.1(11)
 Instance optimality, § 3.2(22)
 Instance-optimality, § 4.4(51)
 Instance-optimality in probability, § 4.4(51)
- J** Johnson-Lindenstrauss lemma, § 3.3(24), § 3.9(38)
- L** L1 minimization, § 4.2(45), § 4.3(47), § 4.4(51), § 4.5(53), § 4.7(55)
 Lasso, § 4.7(55)
- M** Magnetic resonance imaging (MRI), § 5.4(69)
 Manifolds, § 5.8(81)
 Matching pursuit, § 4.8(58)
 Measurement bounds, § 3.3(24)
 Model selection, § 5.1(67)
 Modulated wideband converter, § 3.5(30)
 Multiband signal, § 5.5(69)
- N** Noise-free recovery, § 4.2(45)
 Nonlinear approximation, § 2.4(17)
 Norms, § 2.1(11)
 Null space property, § 3.2(22), § 3.4(28)
- O** Orthogonal matching pursuit, § 4.8(58)
 Orthonormal basis, § 2.2(13)

- P** p norms, § 2.1(11)
Parseval frame, § 2.2(13)
Phase transition, § 4.5(53)
Power law decay, § 2.4(17)
Probabilistic guarantees, § 4.4(51)
- R** Random convolution, § 3.5(30)
Random demodulator, § 3.5(30), § 5.5(69)
Random filtering, § 3.5(30)
Random matrices, § 3.5(30), § 3.8(35), § 3.9(38)
Random projections, § 5.8(81)
Relevance vector machine, § 4.10(64)
Restricted isometry property, § 3.3(24), § 3.4(28), § 3.5(30), § 3.6(32), § 3.9(38)
- S** Sensing matrices, § 3.1(21)
Sensing matrix design, § 3.1(21)
Sensor networks, § 5.10(88)
Shrinkage, § 4.7(55)
Signal recovery in noise, § 4.3(47)
Single-pixel camera, § 5.6(72), § 5.7(76)
Sketching, § 4.9(62), § 5.3(68)
Spark, § 3.2(22), § 3.6(32)
Sparse Bayesian learning, § 4.10(64)
Sparse linear regression, § 5.1(67)
Sparse recovery, § 4.1(43), § 4.9(62)
Sparse recovery algorithms, § 4.6(54)
Sparse signal recovery, § 4.2(45), § 4.3(47)
Sparsity, § 2.3(14)
Stability in compressed sensing, § 3.3(24)
Stagewise orthogonal matching pursuit, § 4.8(58)
Strictly sub-Gaussian distributions, § 3.7(33)
Sub-Gaussian distributions, § 3.7(33), § 3.9(38)
Sub-Gaussian matrices, § 3.5(30), § 3.9(38)
Sub-Gaussian random variables, § 3.8(35)
Synthesis, § 2.2(13)
- T** The ℓ_0 norm, § 2.1(11)
Tight frame, § 2.2(13)
- U** Uniform guarantees, § 3.2(22), § 4.4(51)
Uniform uncertainty principle, § 3.3(24)
Universality, § 3.5(30)
- V** Vector spaces, § 2.1(11)
Vectors, § 2.1(11)
- W** Welch bound, § 3.6(32)
- ℓ ℓ_0 minimization, § 4.1(43)
 ℓ_1 minimization, § 4.1(43)

Attributions

Collection: *Introduction to Compressive Sensing*

Edited by: Marco F. Duarte

URL: <http://cnx.org/content/col11355/1.2/>

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "The Shannon-Whitaker Sampling Theorem"

By: Michael A Lexa, Ronald DeVore

URL: <http://cnx.org/content/m15146/1.2/>

Pages: 1-2

Copyright: Michael A Lexa, Ronald DeVore

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Stable Signal Representations"

By: Michael A Lexa, Ronald DeVore, Shriram Sarvotham

URL: <http://cnx.org/content/m15144/1.1/>

Pages: 3-5

Copyright: Michael A Lexa, Ronald DeVore, Shriram Sarvotham

License: <http://creativecommons.org/licenses/by/2.0/>

Module: "Optimal Encoding"

By: Ronald DeVore, Shriram Sarvotham

URL: <http://cnx.org/content/m15139/1.1/>

Pages: 5-6

Copyright: Ronald DeVore, Shriram Sarvotham

License: <http://creativecommons.org/licenses/by/2.0/>

Module: "Kolmogorov Entropy"

By: Ronald DeVore, Shriram Sarvotham

URL: <http://cnx.org/content/m15137/1.1/>

Page: 7

Copyright: Ronald DeVore, Shriram Sarvotham

License: <http://creativecommons.org/licenses/by/2.0/>

Module: "Optimal Encoding of Bandlimited Signals"

By: Ronald DeVore, Shriram Sarvotham

URL: <http://cnx.org/content/m15140/1.2/>

Pages: 8-9

Copyright: Ronald DeVore, Shriram Sarvotham

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Introduction to vector spaces"

By: Marco F. Duarte, Mark A. Davenport

URL: <http://cnx.org/content/m37167/1.6/>

Pages: 11-13

Copyright: Marco F. Duarte, Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Bases and frames"
 By: Marco F. Duarte, Mark A. Davenport
 URL: <http://cnx.org/content/m37165/1.6/>
 Pages: 13-14
 Copyright: Marco F. Duarte, Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Sparse representations"
 By: Marco F. Duarte, Mark A. Davenport
 URL: <http://cnx.org/content/m37168/1.5/>
 Pages: 14-16
 Copyright: Marco F. Duarte, Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Compressible signals"
 By: Marco F. Duarte, Mark A. Davenport
 URL: <http://cnx.org/content/m37166/1.5/>
 Pages: 17-20
 Copyright: Marco F. Duarte, Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Sensing matrix design"
 By: Mark A. Davenport
 URL: <http://cnx.org/content/m37169/1.6/>
 Page: 21
 Copyright: Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Null space conditions"
 By: Marco F. Duarte, Mark A. Davenport
 URL: <http://cnx.org/content/m37170/1.6/>
 Pages: 22-24
 Copyright: Marco F. Duarte, Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "The restricted isometry property"
 By: Mark A. Davenport
 URL: <http://cnx.org/content/m37171/1.6/>
 Pages: 24-28
 Copyright: Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "The RIP and the NSP"
 By: Mark A. Davenport
 URL: <http://cnx.org/content/m37176/1.5/>
 Pages: 28-30
 Copyright: Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Matrices that satisfy the RIP"
 By: Mark A. Davenport
 URL: <http://cnx.org/content/m37177/1.5/>
 Pages: 30-32
 Copyright: Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Coherence"

By: Marco F. Duarte

URL: <http://cnx.org/content/m37178/1.5/>

Pages: 32-33

Copyright: Marco F. Duarte

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Sub-Gaussian random variables"

By: Mark A. Davenport

URL: <http://cnx.org/content/m37185/1.6/>

Pages: 33-35

Copyright: Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Concentration of measure for sub-Gaussian random variables"

By: Mark A. Davenport

URL: <http://cnx.org/content/m32583/1.7/>

Pages: 35-38

Copyright: Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Proof of the RIP for sub-Gaussian matrices"

By: Mark A. Davenport

URL: <http://cnx.org/content/m37186/1.4/>

Pages: 38-41

Copyright: Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Signal recovery via ℓ_1 minimization"

Used here as: "Signal recovery via ℓ_1 -norm minimization"

By: Mark A. Davenport

URL: <http://cnx.org/content/m37179/1.5/>

Pages: 43-45

Copyright: Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Noise-free signal recovery"

By: Mark A. Davenport

URL: <http://cnx.org/content/m37181/1.6/>

Pages: 45-47

Copyright: Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Signal recovery in noise"

By: Mark A. Davenport

URL: <http://cnx.org/content/m37182/1.5/>

Pages: 47-51

Copyright: Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Instance-optimal guarantees revisited"
 By: Mark A. Davenport
 URL: <http://cnx.org/content/m37183/1.6/>
 Pages: 51-53
 Copyright: Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "The cross-polytope and phase transitions"
 By: Mark A. Davenport
 URL: <http://cnx.org/content/m37184/1.5/>
 Pages: 53-54
 Copyright: Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Sparse recovery algorithms"
 By: Chinmay Hegde
 URL: <http://cnx.org/content/m37292/1.3/>
 Pages: 54-55
 Copyright: Chinmay Hegde
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Convex optimization-based methods"
 By: Wotao Yin
 URL: <http://cnx.org/content/m37293/1.5/>
 Pages: 55-58
 Copyright: Wotao Yin
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Greedy algorithms"
 By: Chinmay Hegde
 URL: <http://cnx.org/content/m37294/1.4/>
 Pages: 58-62
 Copyright: Chinmay Hegde
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Combinatorial algorithms"
 By: Mark A. Davenport, Chinmay Hegde
 URL: <http://cnx.org/content/m37295/1.3/>
 Pages: 62-64
 Copyright: Mark A. Davenport, Chinmay Hegde
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Bayesian methods"
 By: Chinmay Hegde, Mona Sheikh
 URL: <http://cnx.org/content/m37359/1.4/>
 Pages: 64-66
 Copyright: Chinmay Hegde, Mona Sheikh
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Linear regression and model selection"
 By: Mark A. Davenport
 URL: <http://cnx.org/content/m37360/1.3/>
 Page: 67
 Copyright: Mark A. Davenport
 License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Sparse error correction"

By: Mark A. Davenport

URL: <http://cnx.org/content/m37361/1.3/>

Pages: 67-68

Copyright: Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Group testing and data stream algorithms"

By: Mark A. Davenport

URL: <http://cnx.org/content/m37362/1.4/>

Page: 68

Copyright: Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Compressive medical imaging"

By: Mona Sheikh

URL: <http://cnx.org/content/m37363/1.4/>

Page: 69

Copyright: Mona Sheikh

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Analog-to-information conversion"

By: Mark A. Davenport, Jason Laska

URL: <http://cnx.org/content/m37375/1.4/>

Pages: 69-72

Copyright: Mark A. Davenport, Jason Laska

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Single-pixel camera"

By: Marco F. Duarte, Mark A. Davenport

URL: <http://cnx.org/content/m37369/1.4/>

Pages: 72-76

Copyright: Marco F. Duarte, Mark A. Davenport

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Hyperspectral imaging"

By: Marco F. Duarte

URL: <http://cnx.org/content/m37370/1.4/>

Pages: 76-81

Copyright: Marco F. Duarte

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Compressive processing of manifold-modeled data"

By: Marco F. Duarte

URL: <http://cnx.org/content/m37371/1.6/>

Pages: 81-85

Copyright: Marco F. Duarte

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Inference using compressive measurements"

By: Marco F. Duarte

URL: <http://cnx.org/content/m37372/1.4/>

Pages: 85-88

Copyright: Marco F. Duarte

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Compressive sensor networks"

By: Marco F. Duarte

URL: <http://cnx.org/content/m37373/1.3/>

Pages: 88-89

Copyright: Marco F. Duarte

License: <http://creativecommons.org/licenses/by/3.0/>

Module: "Genomic sensing"

By: Mona Sheikh

URL: <http://cnx.org/content/m37374/1.3/>

Pages: 90-91

Copyright: Mona Sheikh

License: <http://creativecommons.org/licenses/by/3.0/>

Introduction to Compressive Sensing

Course notes for Fall 2011

About Connexions

Since 1999, Connexions has been pioneering a global system where anyone can create course materials and make them fully accessible and easily reusable free of charge. We are a Web-based authoring, teaching and learning environment open to anyone interested in education, including students, teachers, professors and lifelong learners. We connect ideas and facilitate educational communities.

Connexions's modular, interactive courses are in use worldwide by universities, community colleges, K-12 schools, distance learners, and lifelong learners. Connexions materials are in many languages, including English, Spanish, Chinese, Japanese, Italian, Vietnamese, French, Portuguese, and Thai. Connexions is part of an exciting new information distribution system that allows for **Print on Demand Books**. Connexions has partnered with innovative on-demand publisher QOOP to accelerate the delivery of printed course materials and textbooks into classrooms worldwide at lower prices than traditional academic publishers.