

INTERDISCIPLINARY MATHEMATICS INSTITUTE

2010:09

Compressed Sensing and
Electron Microscopy

P. Binev, W. Dahmen, R. DeVore, P.
Lamby, D. Savu and R. Sharpley

IMI

PREPRINT SERIES

COLLEGE OF ARTS AND SCIENCES
UNIVERSITY OF SOUTH CAROLINA

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 2010		2. REPORT TYPE		3. DATES COVERED 00-00-2010 to 00-00-2010	
4. TITLE AND SUBTITLE Compressed Sensing and Electron Microscopy				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) University of South Carolina, College of Arts and Sciences, Interdisciplinary Mathematics Institute, Columbia, SC, 29208				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT Compressed Sensing (CS) is a relatively new approach to signal acquisition which has as its goal to minimize the number of measurements needed of the signal in order to guarantee that it is captured to a prescribed accuracy. It is natural to inquire whether this new subject has a role to play in Electron Microscopy (EM). In this paper, we shall describe the foundations of Compressed Sensing and then examine which parts of this new theory may be useful in EM.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 53	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Compressed Sensing and Electron Microscopy

Peter Binev, Wolfgang Dahmen, Ronald DeVore, Philipp Lamby,
Daniel Savu, and Robert Sharpley *

Abstract

Compressed Sensing (CS) is a relatively new approach to signal acquisition which has as its goal to minimize the number of measurements needed of the signal in order to guarantee that it is captured to a prescribed accuracy. It is natural to inquire whether this new subject has a role to play in Electron Microscopy (EM). In this paper, we shall describe the foundations of Compressed Sensing and then examine which parts of this new theory may be useful in EM.

AMS Subject Classification: 94A12, 65C99, 68P30, 41A25, 15A52

Key Words: compressed sensing, electron microscopy, sparsity, optimal encoding and decoding.

1 Introduction

Images formed from modern electron microscopes play a central role in the analysis of the composition and structure of materials [23, 32]. In particular, processing the data from STEM (scanning transmission electron microscopes [22, 2]), is becoming increasingly important for the analysis of biological and other soft materials at fine resolution. However, the effective and realistic imaging of fine scale structures requires high density sampling with good signal to noise ratios and consequently a significant number of electrons must be applied per unit area. This intrusion into the sampled material can result in structural changes or even a destruction of the observed portion. Thus, a critical issue in electron microscopy is the electron dosage needed to produce a suitable quality image. Higher dose scans can damage the specimen while lower dose scans result in high noise content in the signal. It is therefore a central question to determine *how low can one keep the dose* while still being able to faithfully extract the information held at the highest physically possible resolution level. This calls for the development of specially tailored imaging techniques for electron microscopy that are able to go beyond the confines of currently used off-the-shelf tools.

*This research was supported in part by the College of Arts and Sciences at the University of South Carolina, the ARO/DoD Contracts W911NF-05-1-0227 and W911NF-07-1-0185; the Office of Naval Research Contract ONR-N00014-08-1-1113; the NSF Grants DMS-0915104 and DMS-0915231; the Special Priority Program SPP 1324, funded by DFG.

Compressed Sensing (CS) is an emerging new discipline which offers a fresh view of signal/image acquisition and reconstruction. The goal of compressed sensing is to acquire a signal with the *fewest number of measurements*. This is accomplished through innovative methods for sampling (encoding) and reconstruction (decoding). The purpose of this paper is to describe the main elements of compressed sensing with an eye toward their possible use in Electron Microscopy (EM). In fact, correlating “low dose” with “fewest possible measurements” triggers our interest in exploring the potentially beneficial use of CS-concepts in EM.

In the following section, we shall give the rudiments of Compressed Sensing. We tailor our presentation to the acquisition and reconstruction of images since this matches the goals of EM. The subsequent sections of this paper will discuss possible uses of CS in Electron Microscopy. More specifically, we shall address two scenarios. The first applies to high resolution EM acquisition for materials with crystalline-like lattice structure, and the second corresponds to a much lower resolution level, which is a typical setting for electron tomography.

2 The foundations of compressed sensing

The ideas of compressed sensing apply to both image and signal acquisition and their reconstruction. Since our main interest is to discuss whether these ideas have a role to play in Electron Microscopy, we shall restrict our discussion to image acquisition.

Typical digital cameras acquire an image by measuring the number of photons that impinge on a collection device at an array of physical locations (pixels). The resulting array of pixel values is then compressed by using a change of basis from pixel representation to another representation such as discrete wavelets or discrete cosines. In this new representation, most basis coefficients are small and are quantized to zero. The positions and quantized values of the remaining coefficients can be described by a relatively small bitstream.

Since the compressed bitstream uses far fewer bits than the original pixel array, it is natural to ask whether one could have - in the very beginning - captured the image with fewer measurements; for example a number of measurements which is comparable to the number of bits retained. Compressed Sensing answers this question in the affirmative and describes what these measurements should look like. It also develops a quantitative theory that explains the efficiency (distortion rate) for these new methods of sampling.

The main ingredients of this new theory for sensing are: (i) a new way of modeling real world images by using the concept of sparsity, (ii) new ideas on how to sample images, (iii) innovative methods for reconstructing the image from the samples. Each of these components can shed some light on Electron Microscopy and indeed may improve the methodology of EM acquisition and processing. To understand these possibilities we first describe the primary components of CS.

2.1 Models classes for images

A digitized image is an array of N pixel values which can be represented by a matrix with real entries. We can also think of each digitized image as a vector $f \in \mathbb{R}^N$ obtained by scanning the pixel values in a specified order; usually this is the first row from left to right and then the second row left to right and so on. We shall treat the components of f as real numbers, although sensors would quantize these real numbers to a certain number of bits (typically eight or sixteen). One should view N as very large. As the resolution of sensors improves, N will grow.

If all possible vectors $f \in \mathbb{R}^N$ could appear as the pixel array of an image, there would be no hope for compression or fast acquisition. However, it is generally agreed that the images that are of interest represent a small number of the mathematically possible f . How can we justify this claim when we do not have a precise definition of real world images? We present the two most common arguments.

Firstly, one can carry out the following experiment. Randomly assign pixel values and display the resulting image. Each such image is a mathematically allowable image occurring with equal probability. One will see that all of the resulting images will have no apparent structure and do not match our understanding of real world images. Thus, real world images are such a small percentage of the mathematically possible images that we never even see one by this experiment.

A second more mathematical argument is to recognize that the pixel values that occur in a real world image have some regularity. This is not easy to see with the pixel representation of the image so we shall make a basis transformation to draw this out. The pixel representation can be thought of as representing the vector f in terms of the canonical basis functions $e_i \in \mathbb{R}^N$, $i = 1, \dots, N$, where the vector e_i is one in the i -th position but zero in all other entries. So $f = \sum_{i=1}^N p(i)e_i$ with $p(i)$ the corresponding pixel value. There are of course many other natural bases $\{b_1, b_2, \dots, b_N\}$ (with $b_j \in \mathbb{R}^N$) that could also be used to represent f . Two that are commonly used for images are the discrete Fourier and a discrete wavelet bases. We can write our image vector f in terms of these basis elements, $f = \sum_{i=1}^N x(i)b_i$. Notice that the coefficient vector $x = Bf$ for a suitable change of basis $N \times N$ matrix B . The vector x is again in $\mathbb{R}^N$. If one carries out this change of basis for real world images to either of the above mentioned bases, then one observes that most of the coefficients $x(i)$ are zero or very small.

Figures 1-2 are an illustration of this fact. The 512×512 raw image in Figure 1 a) is of an M1 catalyst, a phase of mixed-metal oxide in the system Mo-V-Nb-Te-O from EM. Although this image looks to have very regular structure, a magnification of the image (Figure 1 b)) demonstrates that there is little regularity at the pixel level.

If we look at the histogram of pixel values there is no particular structure (Figure 2 a)). However, if we write this image in a wavelet representation (Haar system, for example), then we see that the histogram of coefficients noticeably peak at zero, meaning that most coefficients in this basis representation are either zero or very small (Figure 2 b)). This behavior is typical of all real world images.

It is useful to give this second argument a more mathematical formulation. For this, we introduce the concepts of sparsity and compressibility. We say a vector $x \in \mathbb{R}^N$ has *sparsity* k if at most k of the entries in x are nonzero. We denote by Σ_k the set of all

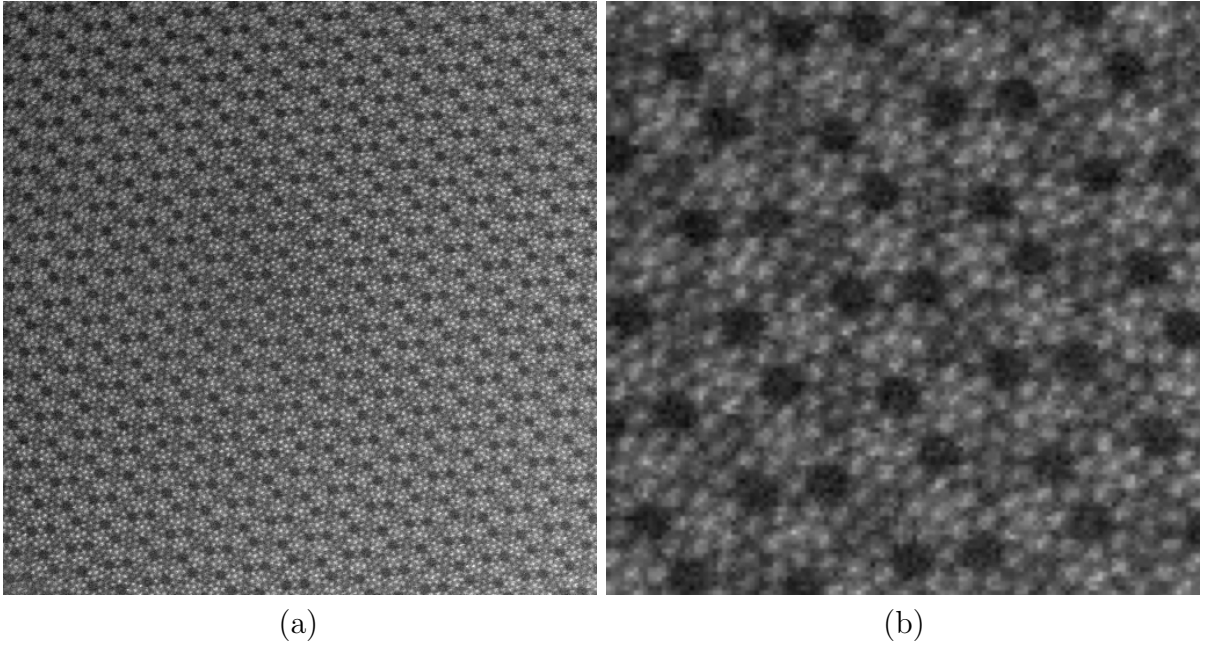


Figure 1: Electron microscopy images: (a) 512 by 512 M1 catalyst in the Mo-V-Nb-Te-O family of mixed oxides; (b) four times magnification of cropped northwest corner of EM image.

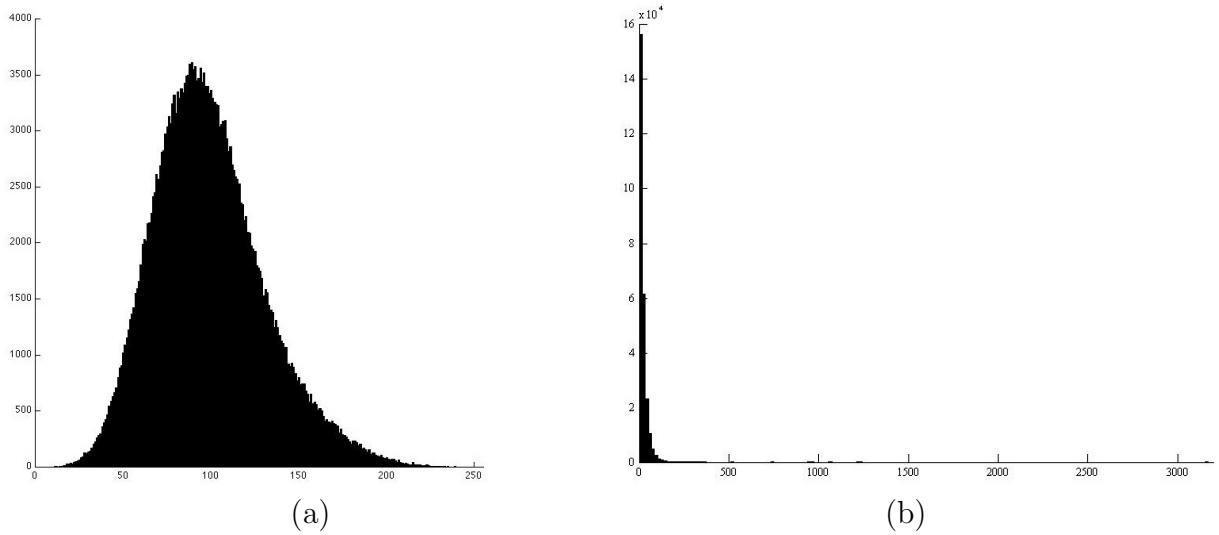


Figure 2: Comparison of histograms of coefficients in the pixel and wavelet bases for the M1 catalyst demonstrating sparsity in the wavelet basis: (a) standard image histogram of pixel values; (b) histogram of wavelet coefficients showing sparsity.

vectors $x \in \mathbb{R}^N$ which have sparsity k . Notice that Σ_k is not a linear space since we have not specified the location of the nonzero entries. For example, each of the coordinate vectors e_i , $i = 1, \dots, N$, is 1-sparse but their linear span is all of $\mathbb{R}^N$.

A vector $x \in \Sigma_k$ is much simpler than a general vector $x \in \mathbb{R}^N$ since it can be described by $2k$ pieces of information, namely, the k positions i where $x(i) \neq 0$ (called the *support of x*) and the values of x at these positions. Thus, if an image f has a coefficient vector $x = Bf$ which is in Σ_k for some small value of k , then this f is considerably simpler than a general vector in $\mathbb{R}^N$. Notice that we are not saying that f itself is sparse (this would correspond to only a small number of pixel values of f are nonzero). Rather, we are saying that after we transform f to a suitable basis, the resulting basis coefficients are sparse. For example, if f has a periodic structure, then transforming to a Fourier basis would result in a sparse representation.

Of course, it is a very idealized assumption to say that $x = Bf$ is sparse. Real images do not give sparse coefficients sequences x because the images have detail at fine scale and also the image may be corrupted by sensor noise. What is true is that real world images are usually well approximated by sparse sequences and it is indeed this fact and this fact alone that allows them to be successfully compressed by transform methods.

We shall next give a precise mathematical formulation for the notion of being well approximated. To do this, we must first agree upon a method to measure distortion. In engineering disciplines, the measurement of distortion is almost exclusively done in the least squares sense. Given our original image $f = (f(i))_{i=1}^N$ and given some compressed version $\hat{f} = (\hat{f}(i))_{i=1}^N$ of f , the least squares distortion between these two images is ¹

$$\|f - \hat{f}\| := \left(\sum_{i=1}^N |f(i) - \hat{f}(i)|^2 \right)^{1/2}. \quad (2.1)$$

The smaller this quantity is the better we think that $\hat{f}$ represents f . If our basis (b_i) is an orthonormal system $x = Bf$ and $\hat{x} = B\hat{f}$ are the coefficient sequences for f and $\hat{f}$ respectively, then

$$\|f - \hat{f}\|_{\ell_2} := \left(\sum_{i=1}^N |x(i) - \hat{x}(i)|^2 \right)^{1/2}. \quad (2.2)$$

Thus, we can also measure the distortion after we have transformed to a new basis.

Measuring distortion in the above least squares norm, while customary, is not the only possibility and it may indeed be that other norms better describe the intended application. The least squares norm is a special case of the ℓ_p norms (or quasi-norms) of sequences defined by

$$\|x\|_{\ell_p} := \|x\|_{\ell_p^N} := \begin{cases} \left(\sum_{j=1}^N |x_j|^p \right)^{1/p}, & 0 < p < \infty, \\ \max_{j=1, \dots, N} |x_j|, & p = \infty. \end{cases} \quad (2.3)$$

Notice that for $p = 2$ we have the least squares norms used above to measure distortion.

¹In mathematics the symbol $:=$ is used to mean that the quantity nearest the colon $:$ is defined by the quantity nearest the equal sign $=$

Any of these norms could equally well be used to measure distortion but this is not our main reason for introducing them. Rather, we want to point out that when these norms are applied to the basis coefficients x of our image f , they give a measure of how nice f is, as we shall now explain. Let us first notice two properties of the ℓ_p norms.

Rearrangement Property: *The ℓ_p norm depends only on the size of the entries in x and not where they appear in the sequence. If we rearrange the entries in the sequence we get a new vector but it has exactly the same ℓ_p norm as x .*

Monotonicity Property: *If we fix the vector x then $\|x\|_{\ell_p} \leq \|x\|_{\ell_q}$ whenever $q \leq p$. For example, the vector $x = (1, 1, \dots, 1)$ has least squares norm $N^{1/2}$ but $\|x\|_{\ell_1} = N$. Thus, a requirement placed on x of the form $\|x\|_{\ell_p} \leq 1$ is stronger (harder to satisfy) as p gets smaller.*

We have made the claim that compression of an image is possible if its transformed coefficients $x = Bf$ can be well approximated by sparse vectors. In order to make this claim precise, we introduce the error in approximating a general vector x by the elements of Σ_k . Although we are primarily interested in such approximation in the least squares norm, we can make the definition for any norm $\|\cdot\|_X$ and in particular for the ℓ_p norms just introduced. Namely, given a sparsity level k , we define

$$\sigma_k(x)_X := \inf_{z \in \Sigma_k} \|x - z\|_X. \quad (2.4)$$

Thus, $\sigma_k(x)_X$ measures how well we can approximate x by the elements of Σ_k if we decide to measure the error in $\|\cdot\|_X$. This process is referred to as *k-term approximation* and $\sigma_k(x)_X$ is the *error of k-term approximation in X*.

It is particularly simple to understand *k-term approximation* in the ℓ_p norms. The best approximation to x is obtained by finding the set $\Lambda_k := \Lambda_k(x)$ of k coordinates where the entries $|x(i)|$, $i \in \Lambda_k$, are largest. Then the vector x_{Λ_k} , which agrees with x on the coordinates of Λ_k and is otherwise zero, will be in Σ_k and is a best approximation to x (in any of the ℓ_p norms) from Σ_k . The error of *k-term approximation* is then

$$\sigma_k(x)_{\ell_p} = \left(\sum_{i \notin \Lambda_k} |x(i)|^p \right)^{1/p}. \quad (2.5)$$

That is, $\|x - x_{\Lambda_k}\|_{\ell_p} = \|x_{\Lambda_k^c}\|_{\ell_p} = \sigma_k(x)_{\ell_p}$, where Λ_k^c denotes the set complement of Λ_k . This approximation process should be considered as *adaptive* since the indices of those coefficients which are retained vary from one image to another. Note that while the set Λ_k is not unique because of possible ties in the size of coefficients, the error $\sigma_k(x)_{\ell_p}$ is unique.

Let us return to the case of measuring error in the least squares norm (the case $p = 2$). Given an image f and the coefficients $x = Bf$, a typical encoding scheme for compression (see [15], [11], [16], [33]) is to list the absolute value $|x(i)|$ of the coefficients in decreasing size. Thus, we determine $i_1, i_2, \dots, i_N$ such that $|x(i_1)| \geq |x(i_2)| \geq \dots \geq |x(i_N)|$. The first information we would encode (or send to a client) about f is the position i_1 and the value $x(i_1)$. This would be followed by i_2 and $x(i_2)$ and so on. In actuality, we cannot send complete information about $x(i_1)$ because it is a real number and would possibly need an infinite number of bits to exactly describe it. Instead, one sends a fixed number of bits (the

lead bits in its binary representation) so that it is captured to a sufficiently high accuracy. This process is called quantization of the coefficients. The smaller coefficients need fewer bits to capture them to the same accuracy and, once the magnitude of a coefficient is beneath the quantization level, no bits are sent at all (the coefficient is quantized to zero). There are various ways to encode the positions of the coefficients that try to take advantage of the fact that large coefficients tend to organize themselves in certain clusters corresponding to specific image locations (e.g. edges).

The reason the above compression scheme is so efficient lies in the observation we made earlier that for real images f , the coefficient sequence x has relatively few large coefficients. Said in another way, the k term approximation error $\sigma_k(x)_{\ell_2}$ tends to zero quite fast. Let us dig into this a little deeper.

One way to understand how the entries of x tend to zero, when rearranged in decreasing order, is to examine the ℓ_p norm of x . Indeed, if ϵ is the size of the largest coordinate $x(i)$, with $i \notin \Lambda_k$, then we have for $p \leq 2$

$$\sigma_k(x)_{\ell_2}^2 = \sum_{i \notin \Lambda_k} |x(i)|^2 \leq \epsilon^{2-p} \sum_{i \notin \Lambda_k} |x(i)|^p \leq \epsilon^{2-p} \|x\|_{\ell_p}^p. \quad (2.6)$$

There is a simple way to estimate ϵ . Since all the coordinates $x(i)$, $i \in \Lambda_k$, are larger than ϵ , we must have

$$\epsilon^p k \leq \sum_{i \in \Lambda_k} |x(i)|^p \leq \|x\|_{\ell_p}^p. \quad (2.7)$$

When this is combined with (2.6), we obtain the fundamental inequality

$$\sigma_k(x)_{\ell_2} \leq \|x\|_{\ell_p} k^{1/2-1/p}. \quad (2.8)$$

Thus, the smaller we can make p then the faster the decay of $\sigma_k(x)$ and the better that the image f can be compressed.

We have to add some further explanation and words of caution to the above. Since N is finite, every vector from $\mathbb{R}^N$ is in ℓ_p and $\|x\|_{\ell_p}$ is finite. So as we decrease p , the decay rate $k^{1/2-1/p}$ will get better but we have to also consider the increase in $\|x\|_{\ell_p}$. Usually, this has a natural solution as the following examples will point out. First consider the case where $x(i) = i^{-1}$, $i = 1, \dots, N$. If $p > 1$, then $\|x\|_{\ell_p}$ has a reasonable bound. If $p = 1$, then $\|x\|_{\ell_p} \approx \log N$ and if $p < 1$ then $\|x\|_{\ell_p} \approx N^{1-p}$. So the natural demarcation occurs when $p = 1$. This demarcation also becomes obvious if we let $N = \infty$ because then the sequence x is not in ℓ_1 but in every ℓ_p , $p > 1$.

One can also see this demarcation for natural images. Consider again the EM image f of Figure 1 and its wavelet coefficients $x = Bf$. If we compute $\sigma_k(x)_{\ell_2}$ and display $\sigma_k(x)_{\ell_2}$ versus k on a log - log plot as in Figure 3, we see an approximate straight line whenever k is not too close to the dimension N of f . This slope of this line is the log of the right side of (2.8) and gives for this particular image that $1/2 - 1/p = -0.1007$. The negative slope $\alpha = 0.2014$ of this line was estimated by a least squares fit. This value thereby determines the natural value of p , which for this example is $p = 1.6647$. The ordinate-intercept on the plot for this example gives an estimate of the norm $\|x\|_{\ell_p} = 0.1738$. The linear fit breaks down as k nears the dimension of the vector f . If we took a finer

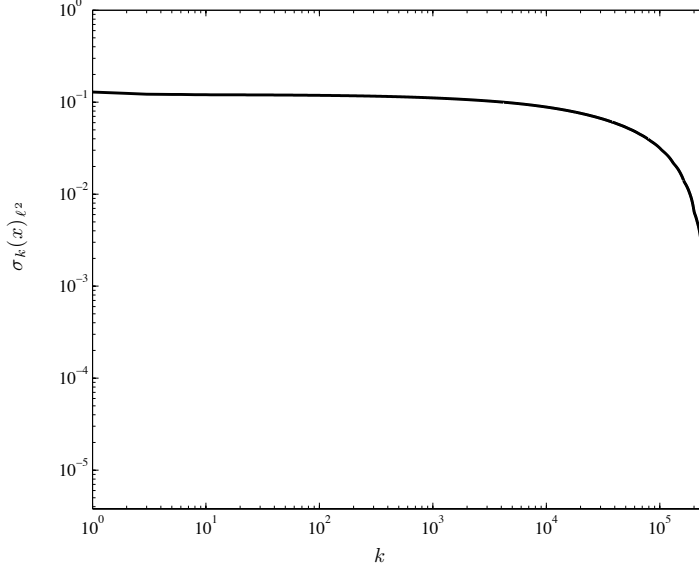


Figure 3: Log-log plot of $\sigma_k(x)_{\ell^2}$ versus k to demonstrate the near linear behavior over the primary range of values. A least squares linear fit of the plot provides an exponential decay rate of $\alpha/2 = 0.1007$ for $\sigma_k(x)_{\ell^2}$.

resolution of the underlying image (more pixels), then the value of N would increase and the linear fit would hold for a larger range of k .

In summary, we have shown in this section that a natural model for real world images is given by sparsity or more generally compressibility (the decay rate of $\sigma_k(x)_{\ell^2}$, $k = 1, 2, \dots$). This decay rate is determined by the smallest value of p for which $\|x\|_{\ell_p}$ does not depend on N . The smaller the value of p , the more compressible the image.

2.2 Sampling

Let us now agree to model images by sparsity or compressibility as described above. Given the extra information that $x = Bf$ is sparse or well approximated by sparse vectors ($\sigma_k(x) \rightarrow 0$ suitably fast), can we say what would be an ideal way to sample f ? To enter such a discussion, we first have to agree on what we would allow as a sample.

CS allows as a measurement any inner product of f with a vector v of our choosing. The result is a real number which is recorded as our sample. In actuality, this number would be quantized but we neglect that aspect at this stage of our discussion. Notice what such an inner product looks like for the image f . We multiply each pixel value $f(i) = p(i)$, $i = 1, \dots, N$, by a real number $v(i)$ and add them all together. The simplest case to understand is when the entries $v(i)$ of v are either 0 or 1. In this case, the inner product counts the total number of photons corresponding to pixels given a one and does not count any others. This is quite easy to implement in a sensor by using micro mirror arrays. So in contrast to a measurement being one pixel value as is the case for digital cameras, now a measurement is a sum of pixel values, the positions of which are of our choosing.

If we make n inner product measurements, then we can represent the totality of

samples by the application of an $n \times N$ matrix A to f . The n entries in Af are the samples we have taken of f . We can also view this as matrix multiplication on the basis coefficient vector x . Namely, since $f = B^{-1}x$, we have that $Af = \Phi x$ where $\Phi = AB^{-1}$. This clarifies the problem. Since our model is that x is either sparse or well approximated by sparse vectors, our problem is to find an appropriate $n \times N$ matrix Φ (called the *CS matrix*) such that

$$y = \Phi x, \quad (2.9)$$

captures enough information about x so that we can approximate x well (or perhaps even determine x exactly in the case it is a sparse vector) from the vector y .

When using CS matrices to sample f , it will not be obvious how to extract the information that the measurement vector y holds about x (respectively f). This is the problem of decoding. A decoder Δ is a mapping from $\mathbb{R}^n$ into $\mathbb{R}^N$. The vector $x^* := \Delta(y) = \Delta(\Phi x)$ is our approximation to x extracted from the information y . We use x^* to create the image $f^* := \sum_{i=1}^N x^*(i) b_i$ as our approximation to f . In contrast to the sensing matrices Φ , we allow the decoder Δ to be nonlinear and perhaps computationally more intensive. We shall discuss good decoders in the following section. The remainder of this section will concentrate on what are good CS-matrices Φ .

How should we evaluate an encoder-decoder pair (Φ, Δ) ? Although not exclusively, by far most research has focused on the ability of such an encoder-decoder pair (Φ, Δ) to recover x exactly when it is sparse. One can show (see [12]) that there are such pairs that recover each $x \in \Sigma_k$ by using only $n = 2k$ measurement which is obviously the smallest number of samples that could work. However, these pairs have a glaring deficiency in that they are unstable. Namely, if we perturb a sparse vector x slightly, the system (Φ, Δ) will give an x^* which is not close to x . Such systems are obviously not useful in practice. It is known that one cannot avoid the instability if one does not enlarge the number of measurements some. The instability problem can be fixed at the expense of requiring slightly more samples. For example, a typical theorem says that there are stable pairs (Φ, Δ) such that whenever $x \in \Sigma_k$, with $k \leq an/\log(N/k)$ for a specified constant a , then $x^* = x$. We will describe such sensing systems in due course but first we formulate a better way to evaluate an encoder-decoder pair.

From both a theoretical and a practical perspective, it is highly desirable to have pairs (Φ, Δ) that are robust in the sense that they are effective even when the vector x is not assumed to be sparse. The question arises as to how we should measure the effectiveness of such an encoder-decoder pair (Φ, Δ) for non-sparse vectors. In [12] we have proposed to measure such performance in a metric $\|\cdot\|_X$ by the largest value of k for which

$$\|x - \Delta(\Phi x)\|_X \leq C_0 \sigma_k(x)_X, \quad \forall x \in \mathbb{R}^N, \quad (2.10)$$

with C_0 a constant independent of k, n, N . We say that a pair (Φ, Δ) which satisfies property (2.10) is *instance-optimal* of order k with constant C_0 . Notice that such an instance-optimal pair will automatically preserve vectors x with sparsity k . Indeed, such a vector has $\sigma_k(x)_X = 0$ and so (2.10) shows that $\Delta(\Phi x) = x$.

Our goal regarding instance-optimality has two formulations. We could be given a value of k and ask to design a pair (Φ, Δ) such that instance optimality of order k holds and the number of rows n in Φ is as small as possible. Another view is that the size n

is fixed for the matrix and we ask what is the largest value of k such that Φ is instance optimal of order k . These two formulations are equivalent.

Instance-optimality heavily depends on the norm employed to measure error. Let us illustrate this by two contrasting results from [12]:

- (i) If $\|\cdot\|_X$ is the ℓ_1 -norm, it is possible to build encoding-decoding pairs (Φ, Δ) which are instance-optimal of order k with a suitable constant C_0 whenever $n \geq ck \log(N/k)$ provided c and C_0 are sufficiently large. Therefore, in order to obtain the accuracy of k -term approximation, the number n of samples needs only to exceed k by the small factor $c \log(N/k)$. We shall speak of the range of k which satisfy $k \leq an / \log(N/k)$ as the *large range* since it is known to be the largest range of k for which instance-optimality can hold.
- (ii) In the case $\|\cdot\|_X$ is the ℓ_2 -norm, if (Φ, Δ) is any encoding-decoding pair which is instance-optimal of order $k = 1$ with a fixed constant C_0 , then the number of measurement n is always larger than aN where $a > 0$ depends only on C_0 . Therefore, the number of non-adaptive measurements has to be very large in order to compete with even one single adaptive measurement. In other words, instance optimality in the least squares norm is not viable. However, as we shall describe in a moment the situation in the least squares norm is not all that bleak.

What are the matrices Φ which give the largest range of instance-optimality for ℓ_1 ? Unfortunately, all constructions of such matrices are given by using stochastic processes. Perhaps the simplest to understand is the Bernoulli random family. If we fix n and N , we can construct a family of matrices with entries $\pm 1/\sqrt{n}$ as follows. We take a fair coin and flip it. If it lands heads we place $+1/\sqrt{n}$ in the $(1, 1)$ position; if it is tails we place $-1/\sqrt{n}$ as this first entry. We then repeat the coin flip to decide on the $(1, 2)$ entry and so on. It is known that with overwhelmingly high probability (but not with certainty) this matrix will satisfy instance optimality for the large range of k .

The unfortunate part of this construction is that if we construct a fixed matrix with $\pm 1/\sqrt{n}$ entries using coin flips, we cannot check whether this matrix actually satisfies instance optimality of order k . So we have to accept the fact that the result of our encoding-decoding may not represent f well. However, this happens with extremely low probability. From this view, it is also possible to remedy the lack of instance optimality in ℓ_2 . Namely, if we use the same Bernoulli matrices then the following probabilistic results hold. Given any x , if we draw the matrix Φ from the Bernoulli family at random then using this Φ together with an appropriate decoder (see [12, 13, 17]) will result in an approximation x^* to x which with high probability satisfies instance optimality in ℓ_2 :

$$\|x - x^*\|_{\ell_2} \leq C_0 \sigma_k(x)_{\ell_2}, \quad (2.11)$$

for the large range of k .

What is different between the ℓ_1 and ℓ_2 instance optimality results? In ℓ_1 instance optimality, when we draw a matrix we are sure that with high probability it will work for all x . So it is either good or bad for all x . On the other hand, in the ℓ_2 case, if we draw the matrix first, no matter how fortunate we are with the draw, our adversary could

find an x for which the instance-optimality fails. However, once x is fixed, for the vast majority of these random matrices we will have instance optimality for this particular x .

There is nothing magical in the Bernoulli family given above other than it is the easiest to describe. It can be replaced by other random variables (for example independent draws of suitably normalized Gaussians) to fill out the matrix. As long as the underlying random variable is sub-Gaussian, the results stated above hold equally well for these other random constructions (see [17]). In fact, a sufficient probabilistic property, shared by all the instances mentioned above, is the following concentration property. For the vast majority of draws Φ from such a family of random matrices one has

$$|\|\Phi x\|_{\ell_2}^2 - \|x\|_{\ell_2}^2| \leq \delta \|x\|^2, \quad (2.12)$$

where $\delta \in (0, 1]$ is fixed. More precisely, the probability that the above inequality fails decays like $be^{-c\delta^2 n}$ with fixed constants b, c depending on the particular random family. As an important consequence, one can show that most elements Φ of a family of random matrices satisfying (2.12) enjoy the so called *restricted isometry property* (RIP) of order k

$$(1 - \eta)\|x\|_{\ell_2} \leq \|\Phi x\|_{\ell_2} \leq (1 + \eta)\|x\|_{\ell_2}, \quad \forall x \in \Sigma_k, \quad (2.13)$$

where $\eta \in (0, 1)$ and k is from the large range. This latter property means that any submatrix of Φ consisting of any k columns of Φ is nearly orthogonal. The RIP is an important analytic property of matrices and is useful because a good RIP guarantees that the matrix will be good in CS.

Finally, let us point out another favorable property of CS matrices constructed by stochastic methods. Suppose that our sample y of x is contaminated by noise, as it almost certainly would be in the design of any practical sensor. Then instead of y we observe $y + e$, where e is a noise vector. Again, applying appropriate decoding techniques to such noisy observations, for instance, those based on greedy or thresholding algorithms (see [13, 27]), we obtain a vector $\bar{x} = \Delta(y + e)$ which satisfies

$$\|x - \bar{x}\|_{\ell_2} \leq C_0 [\sigma_k(x) + \|e\|_{\ell_2}], \quad (2.14)$$

again with high probability. So, as long as the noise level is relatively low, we can retain the performance of k -term approximation. We shall explain below for which type of decoder favorable relations like (2.14) hold.

2.3 Decoding in Compressed Sensing

As we have already noted, the decoding of the information $y = \Phi x$ to get a good approximation to x is not a trivial problem. Historically, the fact that random matrices encode enough information to stably capture sparse vectors was known from the 1970's (see [24, 20]). This fact was not used in designing sensors since it was not known how to reasonably decode this information. It was only recently through the work of Candes and Donoho that practically efficient decoders emerged (see, [9, 7, 18]). Practical decoding still remains an active research area.

To begin the discussion of decoding let us see why the problem is difficult. Given any $x \in \mathbb{R}^N$ and the samples $y = \Phi x$, there are actually many other vectors z which give

the same information, i.e. $\Phi z = y$. Indeed, Φ maps the large dimensional space $\mathbb{R}^N$ into the small dimensional space $\mathbb{R}^n$ and so there is a lot of collapsing of information. For example, any vector η in the null space $\mathcal{N} = \mathcal{N}(\Phi)$ of Φ is mapped into zero and this null space has dimension at least $N - n$. Confronted with this fact, one should be skeptical that compressed sensing can actually work as stated above. However, what is saving the day is our model assumption that x is sparse (or that it is well approximated by sparse vectors). For example, the random matrices Φ used in compressed sensing have the property that with high probability no suitably sparse vectors are in this null space save for the zero vector. Namely, if the matrix Φ has size $n \times N$ with $n \geq ck \log(N/k)$, then no vector from Σ_{2k} is in the null space of Φ . This geometrical fact is behind the amazing performance of these matrices.

In designing a decoder we want to take advantage of the above geometry. Assume for a moment that $y = \Phi x$ for a sparse x . While

$$\Phi z = y \tag{2.15}$$

is a highly underdetermined system of equations for the unknown z , we know there is only one sparse solution (namely x) and we want to find it. The question is how we should proceed.

A standard approach in solving underdetermined system like (2.15) is to use least squares minimization. This procedure looks at all of the z that solve (2.15) and chooses the one that has smallest least squares norm, i.e. $\|z\|_{\ell_2}$ is smallest. It is easy to find this least squares z by using the Moore-Penrose pseudo-inverse. However, it fails to be the sparse solution and in fact is generally not even a good approximation to the sparse x .

If the reader will recall our discussion of ℓ_p spaces, then the sparse solution we want is the z that satisfies (2.15) which has smallest support. Solving this minimization would definitely find x but it turns out that this minimization problem is a difficult combinatorial problem (NP hard in the language of complexity). So this minimization cannot be made into a practical decoder.

Therefore we are caught between the ℓ_2 solution which does not capture sparsity and the ℓ_0 solution which cannot be numerically executed. One may try to replace ℓ_0 by an ℓ_p with p close to zero. But this leads to a nonconvex optimization problem (whenever $p < 1$) which has its own numerical difficulties. A compromise is to consider ℓ_1 minimization which is a convex optimization problem that can be solved by linear programming. This gives us the decoder Δ , where

$$\Delta(y) := \underset{\Phi z = y}{\operatorname{argmin}} \|z\|_{\ell_1}. \tag{2.16}$$

An intuitive idea of why ℓ_1 -minimization promotes sparsity may be obtained from Figure 4 illustrating the way (2.16) works in the case $N = 2, n = 1$ of a single equation in one unknown when the solution set is a line in the plane. Compressed sensing would position the line so that gradually inflating an initially small ℓ_1 -ball the solution set is touched first by a vertex on one of the coordinate axis, thereby picking a solution with a single nonzero entry.

When this decoder is combined with the random matrices of the previous section, we obtain compressed sensing pairs (Φ, Δ) which perform near optimally for capturing sparse vectors x and also gives the highest range k of instance-optimality.

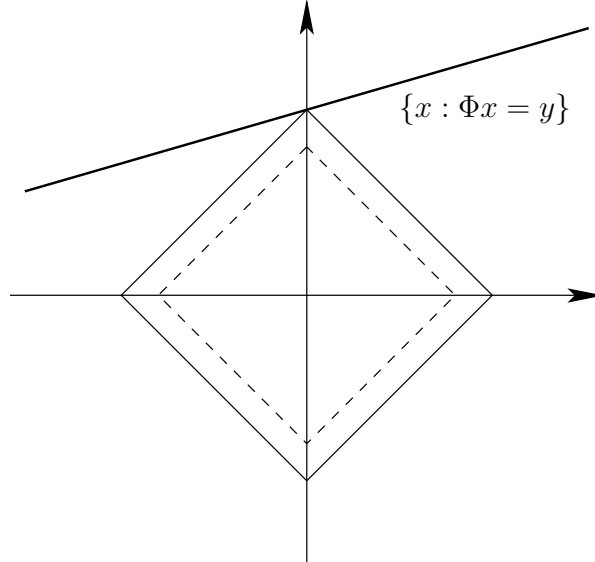


Figure 4: Geometric idea of ℓ_1 minimization.

The general case is not that obvious but, for instance, using Bernoulli random matrices together with the ℓ_1 minimization decoder (2.16) gives a CS system which is instance-optimal in probability for ℓ_2 and the large range of k (see [17]).

There is continued interest in improving the decoding in CS. One goal is to find alternatives to ℓ_1 minimization which may be numerically faster and still provide the same recovery performance as ℓ_1 minimization. This has led to alternatives such as *iterative reweighted least squares* and *greedy algorithms*. Reweighted least squares has its initial goal to capture the ℓ_1 minimizer in (2.16) by solving simpler least squares problems (see [14]). Let us refer the reader to [13] for an extensive discussion of greedy decoders that are computationally very simple. However, for the latter methods to work well it seems to be important that the sensing matrix Φ satisfies RIP (2.13) for rather small values of η while ℓ_1 minimization is quantitatively less stringent on η .

2.4 Dealing with Noise

Decoding in Compressed Sensing may be viewed as one specific instance of an ill-posed problem (since uniqueness is lacking) and ℓ_1 -minimization appears as a *regularization* method promoting sparsity. The methodology of ℓ_1 minimization itself and related variants – in a regularization context – existed long before the emergence of Compressed Sensing and was used for a myriad of inversion/estimation/optimization problems, like deconvolution, deblurring, denoising as well as for general problems in statistical estimation.

For example, it plays a role in what is called Total Variation (TV) denoising used in image processing and more general inverse problems, see e.g. [30]. We recall that the BV-norm $\|g\|_{\text{BV}}$ of a function g measures in a certain way the jump discontinuities of g . We refer the reader to standard analysis texts for its definition.

It will be instructive to digress for a moment and briefly sketch some relevant facts

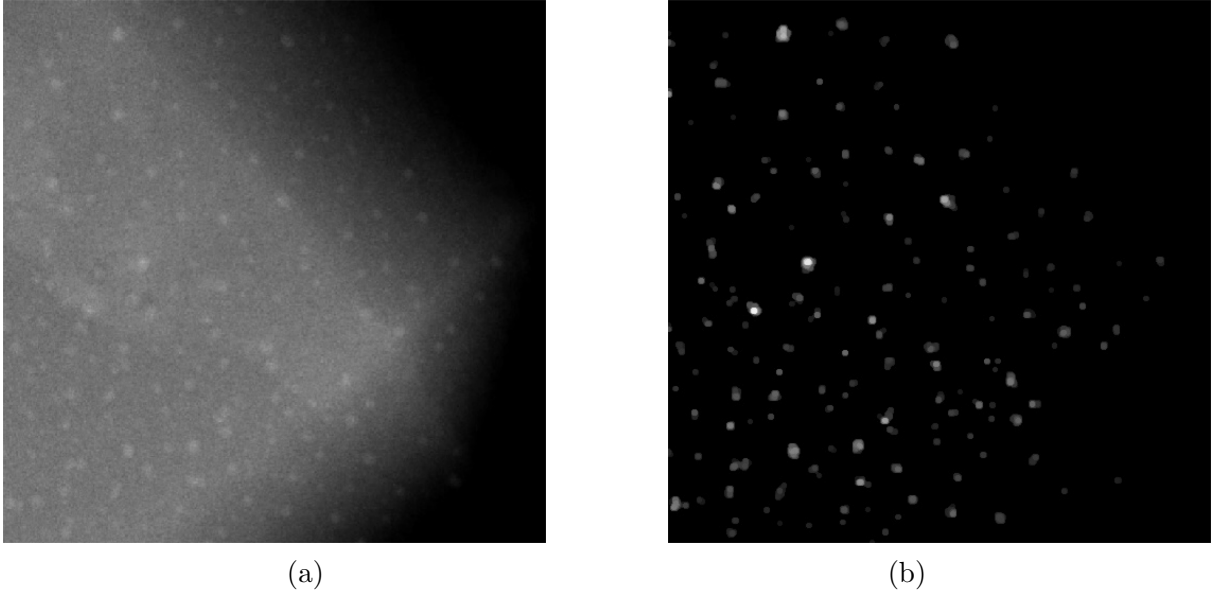


Figure 5: Denoising and feature extraction using multiresolution thresholding and morphological operations: (a) single frame tomographic image from a tilt series (courtesy of Nigel Browning); (b) processed image.

concerning the issue of denoising when f is fully observed, which is of course not the case in CS and EM where we only see f through the sensor via the measurements y . Suppose an image function f is corrupted by additive noise e so that one observes $\bar{f} = f + e$ rather than f . The motivation behind TV denoising is that the image f has structure and finite bounded variation whereas the noise e will not. This leads one to choose a number $\lambda > 0$ and try to approximate the corrupted $\bar{f}$ by a function $\hat{f}$ of bounded variation through the extremal problem

$$\hat{f} := \operatorname{argmin}_g \left\{ \|\bar{f} - g\|^2 + \lambda |g|_{\text{BV}} \right\}. \quad (2.17)$$

The function $\hat{f}$ is called a *total variation (TV) denoising* of $\bar{f}$. Under some models for the noise, one can predetermine the best choice of λ . In most cases λ is found experimentally.

There is a closely related extremal problem in terms of basis coefficients. We write $\bar{f} = \sum_{j=1}^N \bar{x}_j \phi_j$ in terms of our chosen transform basis (e.g. a wavelet basis). If the basis is orthonormal and the noise is white, the perturbation of f translates into a corresponding perturbation of its expansion coefficients x . The analogue of (2.17) is then to solve (see [10])

$$\hat{x} := \operatorname{argmin}_{z \in \mathbb{R}^N} \left\{ \|\bar{x} - z\|_{\ell_2}^2 + \lambda \|z\|_{\ell_1} \right\}. \quad (2.18)$$

This is very easy to implement numerically (in terms of what is called soft thresholding) and gives results close to the denoising of (2.17). For example, using the simplest image processing techniques provides results such as that in Figure 5.

In fact, a heuristic argument why an ℓ_1 penalization helps denoising is that the noise manifests itself in all of the coefficients of x and we want to retain only the large coefficients since they will be guaranteed to be part of the signal. In the present situation of wavelet

expansions, there is a rigorous explanation of its effectiveness. Namely, the ℓ_1 -norm of the wavelet coefficients (normalized in L_2) turns out to be equivalent to a Besov norm that is very close to the BV-norm in (2.17).

Returning now to Compressed Sensing, the sensor noise manifests itself in measurements y and not in x per se. If we assume that this noise $e \in \mathbb{R}^n$ is additive then in place of y , we observe the vector $\tilde{y} = y + e$. As we have noted earlier, for standard families of random matrices, when we decode $\tilde{y}$ by the decoder (2.16) based on ℓ_1 minimization, we receive an $\bar{x}$ which satisfies the ℓ_2 instance optimality estimate (2.14). Notice that in this case there is not a noise reduction (it appears in full force on the right side of (2.14)) but it has not been amplified by the decoding. An alternative often used in the literature (see [7]) is to decode by

$$x^* := \operatorname{argmin} \{ \|z\|_{\ell_1} : z \in \mathbb{R}^N, \|\Phi z - \tilde{y}\|_{\ell_2} \leq \epsilon \}, \quad (2.19)$$

where ϵ is a bound for $\|e\|_{\ell_2}$: $\|e\|_{\ell_2} \leq \epsilon$. The disadvantage of this approach is that it requires an a priori bound for the noise level. Problem (2.19) in turn, is essentially equivalent to the formulation

$$x^* := \operatorname{argmin} \{ \|\Phi z - \tilde{y}\|_{\ell_2}^2 + \lambda \|z\|_{\ell_1} : z \in \mathbb{R}^N \}, \quad (2.20)$$

where λ is related to ϵ and hence to the noise level. Thus, we are back in a situation similar to (2.18). It is not hard to show that solving (2.19) or (2.20) is a decoder realizing (2.14), provided the noise level is known.

There are several strategies to actually solve the optimization problems (2.19) or (2.20). One option is to employ convex optimization techniques. Another is to employ iterative methods involving soft thresholding in each step (as used in a single step for (2.18)). Such concepts have been analyzed for coefficient sequences x appearing in frame representations. Corresponding assumptions do not quite hold in the Compressed Sensing context and one therefore generally experiences a very slow convergence. A certain improvement is offered by variants of such iterations such as *Bregman-iteration*, see [5, 31, 37, 30]. As mentioned above, greedy techniques also lead to decoders satisfying (2.14), however, under much more stringent conditions on η for RIP, see [13, 27]. To ensure their validity for the above examples of random matrix families, the sparsity range k , although asymptotically still in the large range, needs to be more constrained.

2.5 Summary

Let us briefly summarize the essential points of the above findings in order to provide some orientation for the second half of this paper where we discuss possible uses of compressed sensing techniques in EM.

- Objects/images of size N , that are (nearly) k -sparse in some basis, can be recovered (with high accuracy) through a number n of linear (nonadaptive) measurements that is not much larger than the sparsity level, namely

$$n \geq ck \log(N/n). \quad (2.21)$$

- Appropriate measurements that work for the above large range of k make heavy use of randomness.
- The decoders are highly nonlinear and computationally expensive.
- Decoding may be arranged not to amplify noise. But, by themselves they would generally not reduce the noise level. One generally has to be able to tell between noise and best k -term approximation error. Thus, noise control should be part of the acquisition process, a point to be picked up later in more detail.

A small number of random measurements coupled with an appropriate nonlinear decoding allows one to capture a sparse image with relatively few non-adaptive measurements. However, to reconstruct the image using the decoder will require the knowledge of the basis in which the image has a sparse representation. In fact, precise information on sparse representations may relax demands on the measurement side when properly incorporated in the decoding part. The following well-known example (from [8]) illustrates this fact and will later guide one of our approaches.

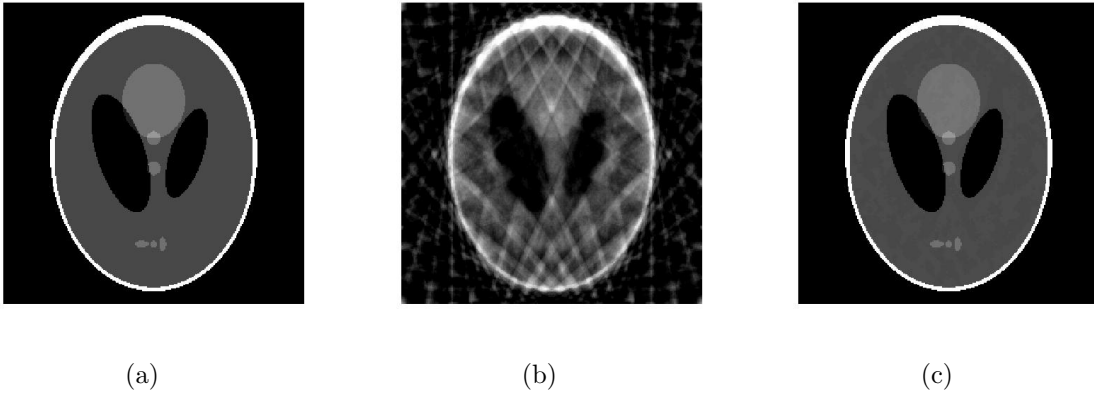


Figure 6: (a) Logan-Shepp phantom; (b) minimum energy reconstruction from 22 projections; (c) TV-regularized reconstruction. The images have been produced using the ℓ_1 -MAGIC code [6].

Figure 6 shows a very simple digital $N \times N$ -image f representing a “piecewise constant” function taking only very few different grey level values associated with 11 ellipses. The pixel representation of f is, of course, a large object. Yet the actual information content is rather small and is held by the pixels demarking the grey level boundaries and the grey levels themselves. In the experiment the discrete Fourier transform

$$\mathcal{F}f(\omega_x, \omega_y) = \sum_{k=0}^{N-1} \sum_{l=0}^{N-1} f(k, l) e^{-2\pi i(\omega_x k + \omega_y l)/N} \quad (2.22)$$

was computed for $\omega = (\omega_x, \omega_y) \in \{0, \dots, N-1\}^2$, but only the coefficients $\omega \in \Omega$ lying approximately on 22 equidistributed radial lines were retained. In practice this would correspond to the available Fourier data from 22 parallel projections. The minimum

energy reconstruction that puts all the other Fourier coefficients to zero and then inverts the discrete Fourier transform yields the Figure 6(b) showing the expected well-known aliasing artifacts caused by subsampling. This reconstruction is actually equivalent to the solution of the problem

$$\hat{f}_{\Omega, \ell_2} = \arg \min \{ \|g\|_{\ell_2} : \mathcal{F}g(\omega) = \mathcal{F}f(\omega), \omega \in \Omega \}. \quad (2.23)$$

However, the retained Fourier information turns out to be sufficient for even reconstructing f exactly when employing a different decoder. In fact, using the prior information that f is piecewise constant, one can look for g which has the same Fourier data $\mathcal{F}g(\omega)$ as f and in addition minimizes the (discrete) total variation:

$$\hat{f}_{\Omega, \text{TV}} = \arg \min \{ \|g\|_{\text{TV}} : \mathcal{F}g(\omega) = \mathcal{F}f(\omega), \omega \in \Omega \}. \quad (2.24)$$

to observe that $\hat{f}_{\Omega, \text{TV}} = f$.

3 What Could CS Buy for Electron Microscopy?

Electron Microscopy differs significantly from other types of image acquisition such as digital cameras. A detailed account of the hardware and physical models for EM would go far beyond the scope of this article. We refer the reader to other contributions in this volume for more discussion on this issue. However, in order to bring forward possible directions in which ideas from Compressed Sensing may offer improvements in EM data acquisition and its image reconstruction, we shall give an idealization of EM imaging in two settings that arise in practice.

To begin the discussion, we need some description of the materials to be studied by EM with the goal of deriving a model class for the images that are to be reconstructed. There are many possibilities here but we shall concentrate on two of these which will be sufficient to illustrate the directions in which we see that CS may have a useful impact.

Model Class 1: Our first example is concerned with the classical case of an extremely thin specimen of at most a few hundred atomic layers thickness. The atoms align themselves in columns and the goal of EM, in particular of STEM, is to determine the position of these columns and the (interpreted) atomic number associated to each of these columns. Ideally, a column consists of atoms of the same type but aberrations of this occur and are important to detect. In any given specimen the possible atoms are drawn from a small class of possibilities (typically no more than five). If there are a total of N columns in the portion of the material under observation, then we can think of the ideal specimen as determined by the set of N positions $\tilde{p}_i$ of these columns and the N (interpreted) atomic number Z_i of the atoms in the given column. Here, without loss of generality, we can think of $\tilde{p}_i$ as a point in the unit square $[0, 1]^2$. Due to atomic vibration, the positions $\tilde{p}_i$ are viewed as stochastic variables with a mean p_i and a probabilistic distribution describing its deviation about the mean. The electron beam is assumed to be positioned parallel to the columns. In a simplistic model deviation from this ideal case could be considered as noise. However, the quantification of the possible local deviations is important and one should try to capture them with a finer model as we propose in Phases 2 and 3 of the experiments considered in Section 3.1.

Even in the simple setting, we want to demarcate between two types of materials depending on their sensitivity to beam intensity.

Type 1: For these materials, the specimen is not significantly altered during the acquisition process provided the beam intensity is low enough. Moreover, we assume that after a suitable relaxation time the material returns to its natural state. Therefore, one can think of rescanning the material many times as long as the beam intensity is low enough. Strontium Titanite and M1 catalysts are examples of materials of this type. Of course, if the maximum allowable beam intensity is small, then the measurements are very noisy, as will be discussed in more detail below.

Type 2: For this type of material, the totality of exposure to electron beams determines whether the material is altered. Zeolites are a typical example. In this case, the arbitrarily rescanning of the specimen is not possible. However, a fixed number of low intensity scans may be utilized in place of one higher intensity scan.

Model Class 2: In the second model class, we assume that material is truly three dimensional. If N is now the number of atoms in the portion of the material under observation, then the position vectors p_i , $i = 1, \dots, N$, are three dimensional. The EM sensing is a form of electron tomography. In this case, one is interested in the 3D-structure of the material under inspection. The resolution is much lower than in the first model class and far from physical resolution limits. The reasons for this are that the material is generally more beam sensitive and that more scans are needed to resolve the three dimensional structure.

For the scenario of Model Class 2 that we shall focus on here, one is primarily interested in the distribution and geometric formation of heavy material clusters immersed in some carrier material whose atomic structure is far from resolved, see e.g. Figure 5. The methodology of approaching this problem is described, for instance, in [36]. The quality of tomographic reconstructions increases with the number of projections. However, we are again faced with the problem of the beam sensitivity of the material which places a limit on the number of projections that can be used.

We shall discuss each of these model classes in detail below. We shall denote by p the position of an atomic column (Model Class 1) or atom (Model Class 2) and by $\mathcal{P}$ the set of all positions in the given specimen. But before beginning such a discussion, we first make some general remarks on EM acquisition and the imaging of the sensor measurements. These remarks will be expanded upon in later sections when we examine each model class in more detail.

The imaging of materials in EM is not a simple process and it seems there is no agreed upon description of the image that would be obtained from a perfect sensor and decoder. However, the following relation exists between the sensor and the material specimen. The electron beam width is smaller than the atomic spacing. A typical setting in STEM is that the beam width is a fraction of an angstrom while the atomic spacing is at least 3 angstroms. When the beam is centered near a particular atom or atom column, the beam produces an intensity distribution at the collector that is proportional to the square of the (interpreted) atomic number Z , i.e. Z^2 . The proportionality constant depends on the distance between the center of the beam and the atom (or atom column). This proportionality constant decays as the center of the beam moves away from the atom.

In Model Class 1, for beam resistant materials like Strontium Titanite or M1 catalysts

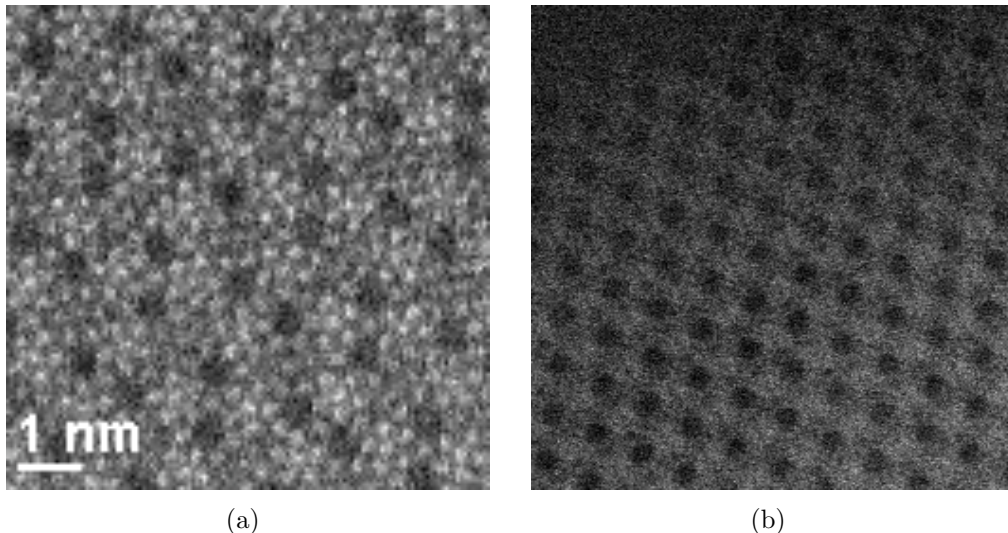


Figure 7: Low dose STEM micrographs of different types of materials (courtesy of Doug Blom): (a) M1 catalyst; (b) zeolite.

the physically feasible level of resolution can be nearly exploited and high resolution images are available, see the left image in Figure 7. But more beam sensitive material like zeolites, that still exhibit similar lattice structures, pose much more of a challenge, see the right part of Figure 7.

Rastering of the beam across the sample enables certain electron imaging and spectroscopic techniques such as mapping by energy dispersive X-ray (EDX) spectroscopy, electron energy loss spectroscopy (EELS) and annular dark-field imaging (ADF). These signals can be obtained simultaneously, allowing direct correlation of image and spectroscopic data. By using a STEM and a high-angle annular detector, it is possible to obtain atomic resolution images where the contrast is directly related to the atomic number ($\approx Z^2$). This is in contrast to conventional high resolution electron microscopy, which uses phase-contrast, and therefore produces results which need simulation to aid in interpretation. Therefore, we shall confine our discussion primarily to HAADF-STEM (High-Angle Annular Dark Field Scanning Transmission Electron Microscopy).

EM, in particular HAADF-STEM will be increasingly important especially in biology. However, the corresponding materials tend to be very beam sensitive so that only very low dosage is applicable without destroying the specimen. As a result one faces extremely low signal to noise ratios. The question is to what extent suitably tailored imaging techniques are able to resolve or at least ameliorate this dilemma, for instance, exploiting CS ideas towards minimizing the number of necessary measurements viz. lowering (or spreading the application of) the total dose.

3.1 High Resolution 2D Images: Model Class 1

In this section, we shall be concerned with images that arise in EM of materials from Model Class 1 (see Figure 7). We shall discuss a model for the ideal STEM images for specimens from this class and argue that these images are sparse with respect to a suitable

dictionary. This will enable the use of ideas from CS for both encoding and decoding.

3.1.1 Image Model and Data Acquisition

Images produced by electron microscopes offer only an indirect reflection of reality. The image is generated from the information extracted by the sensor, namely, the distribution of the intensity of electron scattering at a detector when the beam is centered at chosen locations of the material. We want to understand how this image relates to the atomic description of the material sample and thereby derive a model for such resulting images.

Any image has the form

$$\hat{f} = \sum_P \hat{f}_P \chi_P, \quad (3.1)$$

where χ_P is the characteristic function of the pixel support P and the sum runs over all pixels in the image. In the case of EM, it remains to understand the nature of the intensities $\hat{f}_P$ and how they relate to the atomic structure in the material sample. We shall think of the $\hat{f}_P$ as noisy versions of an ideal pixel value f_P which would result from perfect sensing and decoding.

In STEM for Model Class 1, the electron beam is (nearly) parallel to the atomic columns. The beam is positioned at an equally spaced rectangular grid of points (the raster positions) that we denote by $\mathcal{G}_h$, where h denotes the horizontal and vertical spacing. At each raster position the beam produces an intensity at the detector and results in the assignment of the pixel intensity $\hat{f}_P$ in the image. Thus, the pixels size is the same as the grid spacing h and we can (with only a slight abuse of notation) also index the pixels by $\mathcal{G}_h$. By varying the raster positions, the size of the image can be varied from a very small number of pixels in a frame (256×256) to over 64 million pixels per image (8192×8192).

In STEM mode, the electron dose onto the sample can be controlled in a variety of ways. The number of electrons per unit time can be varied by changing the demagnification of the electron source through the strength of the first condenser lens. The dwell time of the probe is typically varied between $7\mu s$ and $64\mu s$ per pixel in practice, although a much larger range is possible. Finally, the magnification of the image sets the area of the specimen exposed to the electrons and thereby affects the dose per unit area onto the specimen.

We wish to derive a model for the ideal images that would be obtained from the above EM imaging of materials from Model Class 1. Our first goal is to understand the intensity $f(x)$ we should obtain when the beam is placed at position x . Notice that $f(x)$ is a function defined on a continuum of positions. While the position of the electron beam is fixed at a given sampling, the atomic column has a variable position due to atomic vibration and thus the intensity is a random variable. The real number $f(x)$ is the expected intensity at position x obtained by averaging with respect to the underlying probability measure describing atomic position. A model for this expected intensity proposed in [1] is given by

$$f(x) = \sum_{p \in \mathcal{P}} x_p B_p(x), \quad (3.2)$$

where B_p is a bump function (which will require further description), p is the mean position of the atomic column, and the values of x_p are proportional to the squares of the

atomic numbers for the column.

The bump function B_p depends on the nature of the atomic column and the alignment of the atoms within it, but the atoms may neither align exactly in a column nor may the electron beam be perfectly aligned with the column. A first approximation to B_p would be a function which decays from the mean position p of the atomic column in an elliptical pattern. This could be modeled as a tensor product $B_p(x) = G_1(a_1 \cdot (x-p))G_2(a_2 \cdot (x-p))$, where the functions G_1, G_2 are Gaussians with different variances and the vectors $a_1, a_2 \in \mathbb{R}^2$ are an orthogonal pair giving the axes of the ellipse. The nature of a_1, a_2 depends among other things on the alignment of the electron beam with atomic structure. Perhaps this ansatz is still too simplistic. At this point, taking Gaussians, is just a guess and it is not clear at all what a good choice for B_p would be. One could, for instance, think of estimating B_p from images thereby employing a data dependent ansatz. The development of better models for B_p is considered in Section 3.1.4 and is also the subject of future work.

The “ideal” intensity distribution $f(x)$ would for any x in the image plane result from the recorded interaction of the electron beam centered at x with the atomic structure of the material. The images $\hat{f}$ we display in EM are then noisy versions of the *ideal pixelizations* $f^\mathcal{G}$ of the ideal intensity function f for a given pixel lattice $\mathcal{G}$. In other words, the pixel values f_P of $f^\mathcal{G}$ are obtained by averaging f over the pixel $P \in \mathcal{G}$

$$f_P = \frac{1}{|P|} \int_P f, \quad (3.3)$$

and

$$f^\mathcal{G} = \sum_{P \in \mathcal{G}} f_P \chi_P. \quad (3.4)$$

As mentioned before we view the actual values $\hat{f}_P$ as noisy versions of the f_P .

3.1.2 Sparsity for Model Class 1

We first claim that the images $\hat{f}$ we obtain in EM applications to Model Class 1 are in a certain sense sparse so that an application of CS techniques is justified. Such an image is ideally a pixelization of the intensity function f . Thus, if f has a sparse representation with respect to a suitable dictionary then $\hat{f}$ (which as we recall we view as a long vector) will have a sparse representation with respect to the vectors obtained by pixelization of the dictionary elements. So we confine ourselves for the most part to a discussion of the sparsity of f .

It is clear that the ideal image f of (3.2) has sparsity determined by the N positions $p \in \mathcal{P}$ and the N intensities x_p and the number of possible bump functions B_p . If this were all of the information we had then the question of sparsity would be in doubt because of the myriad of possibilities for the positions p . However, as is well known, in an ideal setting, the positions of the atomic columns are aligned along a two dimensional lattice. For instance, Figure 7 displays typical STEM images of M1 catalysts and zeolites, respectively. In both cases the atomic lattice structure is clearly visible. A perfect periodicity is prevented by environmental effects as well as by deficiencies in the material. Nevertheless, the near periodicity amounts to a lowered information content and a form of

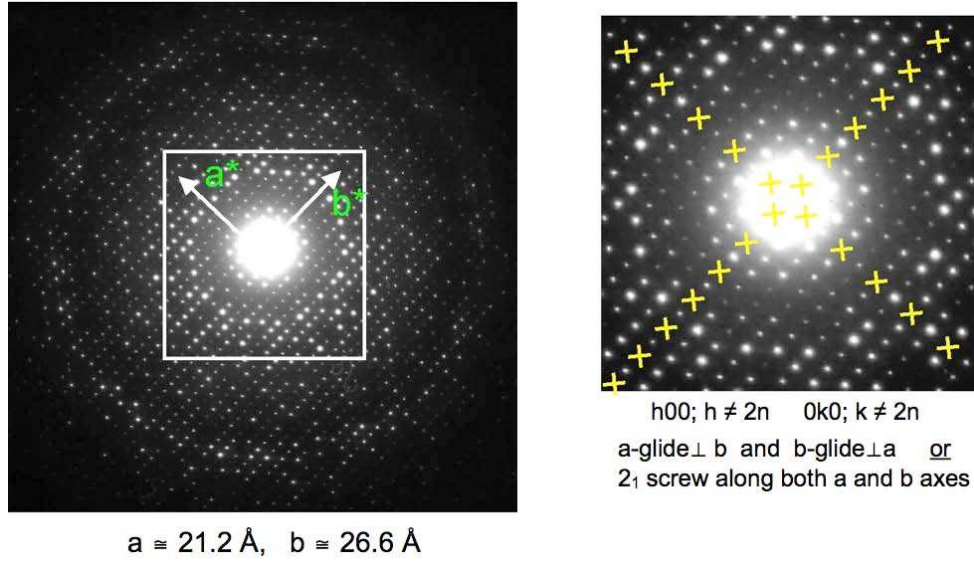


Figure 8: Diffraction pattern in reciprocal space (courtesy of Thomas Vogt)

sparsity. This can indeed be observed by looking at corresponding diffraction patterns in the so called *reciprocal space*, see Figure 8. We could view deviations from an exact lattice structure to be noise or we could add this to our model for f and still retain sparsity (see the examples in Section 3.1.4 for phases 2 and 3).

One can observe from these images that the number of atom columns is quite small. Namely, the area reflecting 60%, say, of the intensity of a typical B_p would be of the order of 10 to 15 pixels in diameter, say, taking the area of the voids into account, the number k of actual column positions could range between 0.1% to 1% of the image size.

Another form of sparsity occurs in the values x_p . In a column of homogeneous atoms, the value of x_p can be taken as Z^2 with the proportionality constant incorporated in the bump function B_p . Thus the number of possible values of x_p would be limited to the number of different atoms. In reality, there are deviations from this homogeneous structure and indeed it is of interest to identify these deviations in the imaging. However, the number of deviations is still small. Thus, ideally the range of x_p is finite and small. This type of sparsity is usually not incorporated into CS models and it is therefore interesting to pursue on a theoretical level how this type of sparsity can be exploited.

This very rough analysis indicates that the sparsity of f in Model 1 is small and therefore f has relatively small information content. It remains to make this analysis more rigorous and to identify precise dictionaries which exhibit the sparsity of the EM images. After pixelization, the sparsity of f translates into a sparsity for the pixelized image $\hat{f}$. For now, we indicate the sparsity level of $\hat{f}$ by k and assume that k is much smaller than the number of pixels $\#(\mathcal{G}_h)$ and turn to the question of how this sparsity can be exploited in EM measurements. It remains to give a rigorous description of k using the remarks on lattice structure and spacing given above.

3.1.3 Measurements for Model Class 1

In traditional STEM imaging, a *measurement* consists of counting and registering the number of collected hits received by the detector as a result of a given positioning of the electron gun. Such a count assigns an intensity value $\hat{f}_P$ to the corresponding pixel in the image. If the material has very little beam sensitivity, a high electron dose per pixel could be applied and gives rise to high resolution images close to the physical resolution limits. However, very beam sensitive materials with a low maximum dose threshold require severe dose restrictions which typically gives rise to noisy images. Thus, we are in a situation where ideas of Compressed Sensing may become interesting since CS says we should be able to capture the image well with roughly k measurements rather than $\#(\mathcal{G}_h)$ measurements. Namely, when measurements are expensive – here damaging – high quality results could possibly be obtainable with a number of measurements comparable to the information content of the signal.

The caveat to the above discussion is that the meaning of measurement in the CS theory is different than the conventional STEM measurement since it requires the sensor to simultaneously test many (or most) locations at once and record the total number of hits not worrying about their pixel location. Let us first discuss how this might be accomplished with current sensors. In what follows the pixel size h will represent the finest level of resolution the recovery procedure is striving for. Therefore, the positions of the atomic columns can only be resolved within a tolerance h and hence will be identified from now on with a subset $\mathcal{P}$ of the *fine grid* $\mathcal{G}_h$. Of course, h is bounded from below by physical constraints and targeting this lower resolution limit would be ideal.

Since we are striving for low dose applications one might use for the actual measurements a larger pixel size $H \geq h$ permitting larger scanning increments. This would give rise to the pixel values $\hat{f}_P$, $P \in \mathcal{G}_H$, from which one would still try to recover the positions $\mathcal{P}$ as well as the coefficients x_p , $p \in \mathcal{P} \subset \mathcal{G}_h$. The very low dose per pixel would entail very low signal to noise ratios for $\hat{f}_P$, so that an accurate recovery of a high resolution image could only be tackled by working with several such coarse frames with a primary focus on denoising. In fact, such a line is pursued in a different way detailed in [4] heavily exploiting the near repetitiveness in images like those in Figure 7.

CS theory however instructs us to proceed in a different way. To describe this alternate strategy, recall from our discussion above that the value $\hat{f}_P$ obtained in the imaging process can be interpreted as

$$\hat{f}_P = f_P + e_P, \quad (3.5)$$

where f_P represents the ideal pixel value that would be obtained through very high dose in (hypothetical) absence of beam damage, and where e_P is a local fluctuation that depends on the applied electron dose and is, in relative terms, the larger the smaller the dose is. Since we are aiming at applying possibly low dose, each single value $\hat{f}_P$, acquired in the above fashion, would give little information.

CS theory tells us that we should make measurements of the following form. We select a set of random locations $S \subset \mathcal{G}_h$ and measure the conglomerate sum

$$Y_S := \sum_{P \in S} \hat{f}_P. \quad (3.6)$$

Thus, the measurement Y_S , rather than being a single pixel value is now a large sum of randomly chosen pixel values. We make many selections for S and record Y_S for each of these. For traditional imaging, this approach has been implemented at Rice University (see [19]) and is known as a “single pixel camera”.

For STEM, this could be implemented during the scanning process by randomly activating or shutting off the electron gun according to, say, a Bernoulli distribution with equal weights. Then, instead of counting the number of electron hits corresponding to each position, we rather count the totality of collected hits from the entire scan. If this turns out to be a useful concept, one can envision new sensors that accomplish one scan in a more economical fashion by simultaneously sensing several selected locations.

There should be serious skepticism concerning the possible advantages of the above approach since in one CS measurement, we are required to touch approximately half of the pixel locations. If this is repeated k times then each pixel location has been hit on average with k times half the individual dosage. So for a fair comparison the individual dosage must be very small and an individual pixel value (which we do not record) would be very noisy. For materials of Type 1, this problem is circumvented by the fact that in a given CS measurement we can choose a dosage at each pixel close to the maximal dosage without damaging the material provided there is a significant relaxation time. This does not hold for materials of Type 2, however we argue that we can expect better signal to noise ratio in CS measurements as compared to traditional STEM measurements. Indeed, in a CS measurement we record the sum of all hits and so the noise will be averaged out in a sum like Y_S and the law of large numbers says that this averaging gives a reduction of noise in a given CS measurement because the number of pixels is much larger than the number of measurements n .

In order to expand on this discussion, we dig deeper into the structure of CS measurements and its relationship to the expected sparsity of the EM image. Let $\tilde{\Phi}$ be a random $n \times \#(\mathcal{G}_H)$ -matrix whose entries are drawn independently and assigned the values $\tilde{\phi}_{i,P} \in \{0, \sqrt{2/n}\}$, $i = 1, \dots, n$, $P \in \mathcal{G}_H$, with equal probability. Now, in these terms, denoting for every $p \in \mathcal{G}_h$ by $(B_{P,p})_{P \in \mathcal{G}_H}$ the vector of pixel values of the corresponding B_p , namely $B_{P,p} = \left(\int_P B_p(x) dx \right) / |P|$, the n CS measurements can be written, in view of (3.2), (3.3), and (3.5), as

$$y_i = \sum_{P \in \mathcal{G}_H} \tilde{\phi}_{i,P} \hat{f}_P = \sum_{P \in \mathcal{G}_H} \tilde{\phi}_{i,P} e_P + \sum_{p \in \mathcal{P}} x_p \left(\sum_{P \in \mathcal{G}_H} \tilde{\phi}_{i,P} B_{P,p} \right), \quad i = 1, \dots, n. \quad (3.7)$$

From this information, we would like to find the positions $p \in \mathcal{P} \subset \mathcal{G}_h$, as well as, ideally, the type of bump function B_p , and the coefficients x_p . Having already restricted $\mathcal{P}$ to be a subset of $\mathcal{G}_h$, in order to make this task tractable, it remains to impose some structure on the B_p as discussed earlier in Section 3.1.1. We shall discuss several such choices in connection with first experiments in the subsequent section.

Having chosen the B_p , we are left with the following linear algebra problem. Given the n measurements y_i , $i = 1, \dots, n$, we search for a sparse solution $\hat{f} = \sum_{p \in \mathcal{G}_h} x_p b_p$ to (3.7) where the vector b_p is the pixelization of $B_p = B(\cdot - p)$

$$b_p := (B_{P,p})_{P \in \mathcal{G}_H}, \quad p \in \mathcal{G}_h. \quad (3.8)$$

In other words, with $y = (y_i)_{1 \leq i \leq n}$, we are looking for a sparse solution to the system of equations

$$y = \Phi x, \quad \text{where} \quad \Phi = \tilde{\Phi} \mathbf{B}_h \quad (3.9)$$

and where $\mathbf{B}_h$ is the matrix whose rows are the vectors b_p , $p \in \mathcal{G}_h$. This is the same form as decoding in CS.

Now, two questions immediately arise: first, what could such *aggregated measurements* buy us? and second, how many such measurements are needed for a good recovery of f and which dose should be applied at each coarse pixel?

As for the first question, recall from (3.7) that the measurements y_i consist of two parts

$$y_i = \sum_{P \in \mathcal{G}_H} \tilde{\phi}_{i,P} f_P + \sum_{P \in \mathcal{G}_H} \tilde{\phi}_{i,P} e_P, \quad (3.10)$$

where the first sum involves the ideal pixel averages while the second sum represents noise generated by the aggregated pixel fluctuations, see (3.5). If the fluctuations e_P had zero mean one would expect the accumulated noise contribution $\sum_{P \in \mathcal{G}_H} \tilde{\phi}_{i,P} e_P$ to be actually as small as the noise associated with the local detector response for the total accumulated electron dose. Thus in summary, the data y_i should have a relatively low noise level, even for materials of Type 2. The very low dose deployed at every activated pixel position should in addition speed up the scanning procedure, so that motion of the specimen should have a diminished effect.

Despite the possible gain in signal to noise ratio in our CS measurements over traditional STEM measurements, we need to note that another difficulty arises. Namely, although Φ is a random matrix, it is not of the standard type to which the CS theory applies directly. In fact, the smaller h - the better the resolution - the more coherent are the columns of the matrix $\mathbf{B}_h$ and hence of Φ . This significantly hampers the correct identification of the positions of the atomic columns. Here, however, there is a redeeming factor in the form of the sparsity of f . The number of atom positions in the image is expected to range between 1% and 0.1%, and these positions are spread out because of the interatomic distances. We could therefore try to determine the set $\mathcal{P} \subset \mathcal{G}_h$ through several stages. At a first stage one could choose a coarser grid $\mathcal{G}_{\bar{h}}$ for some $h < \bar{h} \leq H$ in order to determine a less accurate position set $\mathcal{P}_{\bar{h}} \subset \mathcal{G}_{\bar{h}}$. Since the corresponding matrix $\mathbf{B}_{\bar{h}}$ has less coherent columns the sparse recovery problem is now easier to solve. Once this coarse resolution has been found, we can revisit the sparse inversion problem with a smaller value for $\bar{h}$ by restricting the possible positions to be near the ones we have found before. Proceeding iteratively, we could improve our resolution of the positions $\mathcal{P}$ while maintaining a favorable RIP condition, see (2.13). We shall elaborate more on such strategies in a forthcoming paper. An alternative strategy for coping with large coherence will be indicated below in connection with numerical experiments.

Finally, let us discuss the dosage limits on an application of a CS measurement. If D_C is the critical dosage applicable without damaging the specimen, then in a given CS application, for materials of Type 1, we can apply a dosage close to D_C at each application. Since the number of CS measurements is not restricted in this case, we can expect extremely high quality imaging. For materials of Type 2, however, we would be restricted to a dosage of D_C/n per pixel, where n is the total number of CS measurements

to be taken. So the advantage has to occur in the signal to noise ratio as discussed above. Whether this will ultimately turn out to be sufficient or can even be lowered further has to be seen through detailed simulations and also through experiments.

In summary, we are encouraged to further explore the above perspectives for the reasons already outlined. Another favorable feature of such an aggregated data acquisition process would be that the effect of specimen movement is reduced and thermal relaxation is strongly supported since at each instance the dose is very low. Ultimately, whether CS ideas give a significant improvement of EM remains to be proven through simulation and experiment.

3.1.4 Inversion and Sparse Recovery Techniques

So far the discussion has addressed the question whether CS techniques offer potential benefits in the above high resolution STEM scenario. This is primarily a matter of properly understanding and modeling the data acquisition in connection with the sparsity properties of the image. The second major issue, however, is the ability to actually recover the sparse signal from measurements of the form described above. In this section we focus entirely on this latter issue which we plan to approach in several stages reflected by three types of computational experiments presented below as Phases 1, 2, and 3. Our main objective here is to bring out the essential tasks and difficulties faced when practically applying sparse recovery techniques and compressed sensing concepts in electron microscopy. We shall only indicate the main findings and refer to a more detailed discussion in forthcoming work. It will be seen that several technical problems arise for which we offer one possible solution. Therefore, at this point we rather want to point to specific numerical aspects that should be taken into account in building up a complete inversion model. In this sense our experiments should be considered as steps towards a “proof of concept”.

The validation of the various experiments is based on STEM images obtained by our colleague Doug Blom and computer simulated images produced by him and Tom Vogt’s student Sonali Mitra. The simulation in Figure 9 is based on the Frozen Phonon Model. It is obtained in a two stage process, in which the responses to probes consisting of a single electron are calculated first and then a Gaussian blur is applied with $\sigma = 2.83$ corresponding to a full width at half maximum $FWHM = 0.7nm$ of the intensity of the probe wave. It is important to note that, while the coefficients x_p corresponding to Mo and V atomic columns are large, the nineteen coefficients corresponding to the Oxygen atomic columns are by one order of magnitude smaller and could be misinterpreted as noise, if the noise level gets high. The large dynamic range in connection with noise therefore poses a particular challenge for recovery techniques.

In the first phase of our experiments this image is replaced by an idealized version, but the original simulated image serves as the target image for the second phase.

The idealized target image shown in Figure 10 is based on a very simplistic model and its goal is to check the principal applicability of different minimization routines developed for compressed sensing in the present specific context. We assume that the image (in this case synthetic) is given by (3.2) and the bump function B is known. In particular, we set $B(u) := e^{-u^2/\sigma^2}$. In favor of a higher level of incoherence, we choose $\sigma := 4.8$ which is

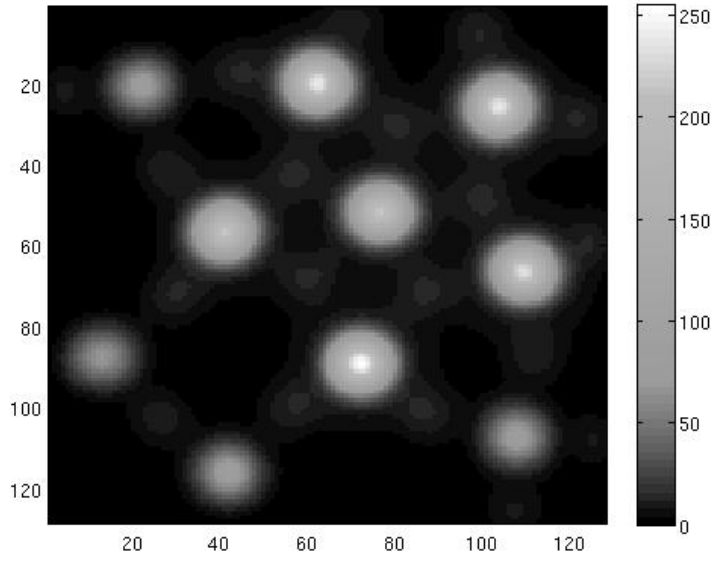


Figure 9: 128×128 computer simulated STEM image of Mo_5V_{14} -oxide; it features 29 atomic columns (6 for Mo , 4 for V , 19 for O). (see [25]).

slightly smaller than the one that will best fit the image in Figure 9.

We now turn to recovery procedures for measurements from the above two target images. Again, in the presence of noise the natural candidate for sparse decoding would be

$$\Delta(y) = \underset{\|\Phi z - y\|_{\ell_2} \leq \epsilon}{\operatorname{argmin}} \|z\|_{\ell_1}, \quad (3.11)$$

where ϵ is the estimated noise level and (the random matrix) Φ is given by (3.9). We shall be using several currently available numerical methods which, however, do not treat (3.11) directly but refer to the related problem (2.20), where the penalty parameter λ needs to be properly chosen depending on the noise level ϵ .

In both experimental phases the finer “high resolution” grid $\mathcal{G}_h$ is set to 128×128 . Thus, we search for the bumps centered at 29 positions $p \in \mathcal{G}_h$ and their intensity values x_p . Likewise in both phases our measurements are taken from a low 64×64 resolution version of the respective target image to which we add different levels of positive Gaussian noise. More precisely, if n is the number of measurements, $y \in \mathbb{R}^n$ is the vector of measurements, and ζ is the desired noise level, we add $\mathcal{N}(\mu_{\text{noise}}, \sigma_{\text{noise}})$, where

$$\mu_{\text{noise}} := \zeta \frac{\|y\|_{\ell_2}}{\sqrt{n}}, \quad \sigma_{\text{noise}} := \frac{\mu_{\text{noise}}}{3}.$$

At this point, the added noise does not necessarily reflect the physical acquisition process but merely quantifies the ability of the decoders to deal with noise.

Phase 1: Denoting as before the set of searched positions by $\mathcal{P}$, we take as the idealized image $f := \sum_{p \in \mathcal{P}} x_p B(\cdot - p)$ and sample the corresponding discrete function $f^{\mathcal{G}_H} =: f^H$

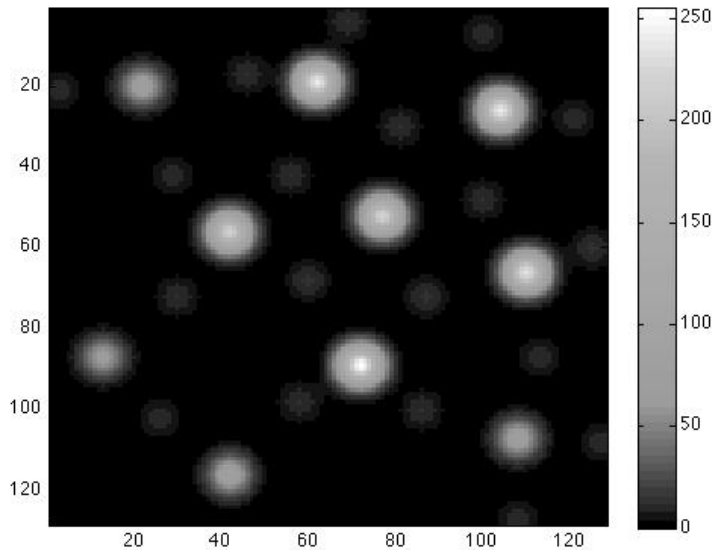


Figure 10: Idealized version of the image in Figure 9 resulting from linear combinations of Gaussians ($\sigma = 4.8$).

defined in (3.4) on a 64×64 grid $\mathcal{G}_H$. We run several reconstruction passes using different numbers n of measurements y_i ranging between 180 and 900. The low number of 180 measurements is much smaller than the total number of 4096 pixels in the low resolution version of the target image. This relation between the number of measurements and signal size is well in line with Compressed Sensing philosophy.

So far, we have resorted only to existing algorithms used in CS to validate our concepts. However, in view of the high coherence of the dictionary $\{B(\cdot - p)\}_{p \in \mathcal{G}_h}$, the matrix Φ from (3.9) is far from satisfying the typical properties required by the sparsity recovering algorithms currently employed in CS applications. Therefore, not all the algorithms we have tested have been able to solve the extremal problem (2.20) in a satisfactory way.

We report here on some results produced by two algorithms, NESTA and SESOP, see [3, 28, 29, 38]. NESTA is an iterative solver for (2.20) which, however, works (in the spirit of a homotopy method) on “nearby” minimization problems of the form

$$\operatorname{argmin}_x \lambda h_\mu(x) + \frac{1}{2} \|\Phi x - b\|_2^2, \quad (3.12)$$

where λ is a fixed penalty parameter that can be chosen by the user. The convex Huber function h_μ (see [3]), ensuring a smooth objective functional, has the parameter μ lowered during the iteration process as much as possible to come close to the original objective functional in (2.20), i.e. h_μ approximates the ℓ_1 -norm when μ tends to zero. For our images the convergence of the method is very slow and requires at least 10^5 outer iterations to receive a meaningful solution. One of the advantages of this method, however, is that it is able to localize well the regions containing the active pixels from the set $\mathcal{P}$. However, realizing the ideally sparse solution is in our case problematic and might require a prohibitively large computational effort.

Therefore, even in the simplistic scenario of Phase 1 the method needs to be adjusted to the specific structure of the problem. In order to speed up the localization process we devised a two stage method. At the first stage we treat the global problem and use it to identify the regions R of energy concentration. Then, at the second stage, we treat localized problems. More precisely, we define $b^R := \Phi \tilde{x}^R$, where $\tilde{x}^R$ is the restriction of the current approximate solution to a single region R , and solve independently the problems (3.12) with $b = b^R$ and the nonzero entries of x restricted only to R . In these local problems we choose higher values of λ in order to promote sparsity since we expect to find a single nonzero entry. Alternatively, one can also simply calculate the local bump $\sum_{p \in R} \tilde{x}_p B(u - p)$ and attribute the energy only to the pixel p that is closest to the point of its maximum. We concatenate the received local solutions and use this as initial guess for a further global iteration revisiting the first stage. This two stage process is iterated, if necessary. The parameter λ is carefully increased during this solution process to enhance the sparsity while maintaining stability, somewhat against the common heuristic approach of other continuation methods, see [3]. As an end result we set the value at the local maxima as a weighted sum of the coefficients in its region. A typical result of this procedure is shown in Figure 11 using 600 measurements with 3% Gaussian noise added to each y_i . The results show very good localization of the positions in $\mathcal{P}$ and a relative ℓ_2 error of 8% for the values of x_p , (3.2), namely 24 out of 29 positions of the atomic columns are recovered correctly, while the other 5 Oxygen positions are recovered within one pixel.

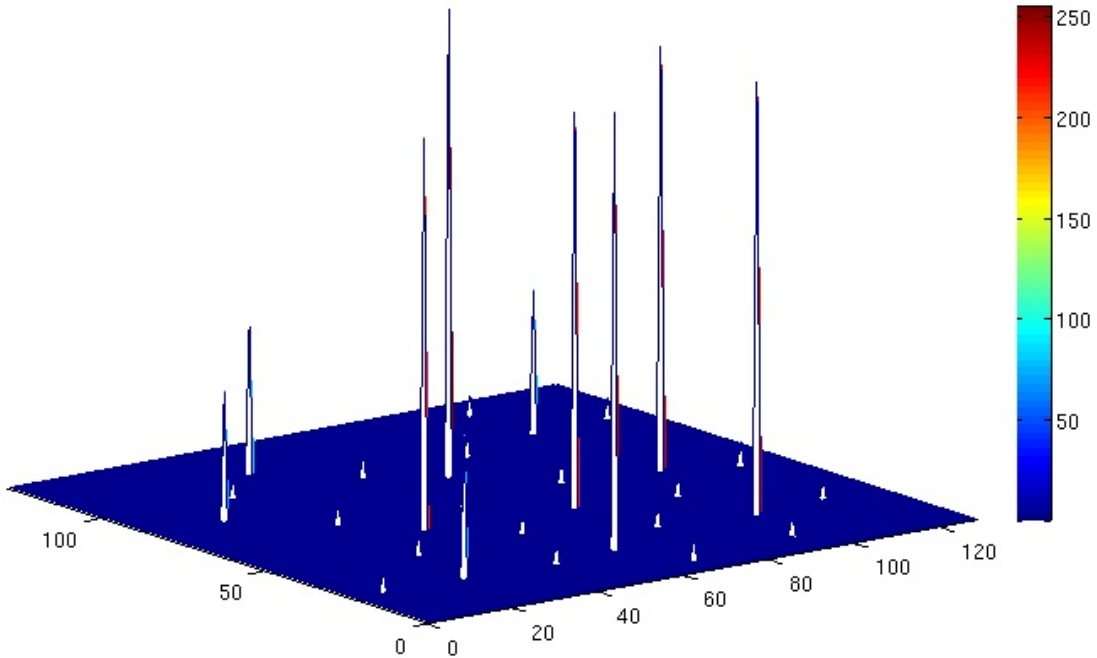


Figure 11: Coefficient reconstruction for the idealized version of Mo_5V_{14} -oxide in Figure 10 with 3% positive Gaussian noise added to each of the 600 measurements.

Similar results are obtained directly by SESOP, even without any adjustments, i.e. no additional localization stage is used. SESOP (see [38]) is also an iterative method, solving

the problem

$$\operatorname{argmin}_x \lambda \|x\|_1 + \|\Phi x - b\|_2^2, \quad (3.13)$$

by adjusting the parameter λ from some good initial guess. Its convergence is in this case faster and the method is more robust regarding the choice of the initial value of λ . However, it is more sensitive to higher levels of noise.

The recovery results produced by SESOP for Phase 1 are displayed in Tables 1 and 2 for several numbers n of measurements, ranging from 180 to 900. Table 1 records the respective relative ℓ_2 -errors

$$E(f) := \frac{\|\tilde{f}^h - f^h\|_2}{\|f^h\|_2} \quad (3.14)$$

on the high resolution (128×128) -grid, where $\tilde{f}^h$ is the approximation to the high resolution of the pixelization f^h via (3.4) of the target image presented in Figure 10. In Table 2 we list the relative ℓ_2 -errors

$$E(c) := \frac{\|\tilde{x} - x\|_2}{\|x\|_2} \quad (3.15)$$

of the recovered coefficients $\tilde{x}$ in the coefficient space, (3.9). We found that a good recovery of the high resolution image, i.e. an acceptably small $E(f)$, can be obtained from as little as 180 measurements. However, good stability, i.e. the accurate detection of the positions, seems to require a higher number of measurements. Specifically, as perhaps expected, $E(c)$ turns out to be much more sensitive towards noise which is seen in Table 2.

Number of measurements	180	300	400	750	900
Added noise: 0%	0.86%	0.45%	0.33%	0.16%	0.11%
1%	6.86%	4.84%	4.02%	2.84%	2.30%
2%	12.52%	8.67%	7.44%	4.96%	4.39%
3%	15.85%	11.60%	9.75%	6.88%	5.95%

Table 1: Relative error $E(f)$ of the SESOP recovery of a 128×128 high resolution idealized image of Mo_5V_{14} -oxide (Figure 10), based on measurements taken from 64×64 grid.

Number of measurements	180	300	400	750	900
Added noise: 0%	9.46%	1.38%	0.90%	0.39%	0.24%
1%	34.87%	23.94%	13.35%	11.78%	10.13%
2%	67.81%	28.61%	23.11%	18.44%	16.81%
3%	87.62%	35.13%	29.04%	25.33%	20.69%

Table 2: Relative error $E(c)$ of the SESOP recovery of the coefficients x_p of an idealized image of Mo_5V_{14} -oxide, based on measurements taken from 64×64 grid.

Phase 2: Our second experiment explores the case in which the local bumps B_p vary and are unknown. As mentioned before, our target is to recover the intensity distribution from the simulated 128×128 STEM image for Mo_5V_{14} -oxide displayed in Figure 9. The

function f^H is obtained by locally averaging the simulated distribution on a 64×64 grid. The physical model underlying the frozen-phonon simulation suggests that the bumps are no longer strictly radial and differ from each other.

The attempts to solve the problem by introducing a specific universal bump function B do not lead to a satisfactory solution. Due to the fact that $B(\cdot - p)$ approximates B_p with an ℓ_2 error as large as 10%, both NESTA and SESOP produced solutions with relative errors of the order of 25% or more. The proper identification of the active coefficients x_p , $p \in \mathcal{P}$ becomes even more difficult due to the high coherence of the dictionary.

To bypass both obstructions we propose again a two-stage strategy. At the first stage, instead of working with the above dictionary whose average bump spread resembles the intensity distribution around an atom position, we choose a dictionary whose elements $\tilde{B}(\cdot - p)$, $p \in \mathcal{G}_h$, are more localized and can therefore individually approximate each of the bumps by a local linear combination. In principle, different choices for the localized trial functions $\tilde{B}$ are conceivable. For example, splines have particularly good local approximation properties. For convenience, in the present experiments we choose $\tilde{B}$, as before in Phase 1, to be a Gaussian, but this time with a *concentration* parameter σ which is less than the one that would have been used to fit the intensity distribution around an atom position, see Table 3.

The approximation resulting from solving the optimization problem is then of the form

$$\tilde{f}(u) = \sum_{q \in \mathcal{G}_h} \tilde{x}_q \tilde{B}(u - q). \quad (3.16)$$

We now expect that the dominating coefficients $\tilde{x}_q$ will form disjoint clusters P_p each of which would learn the actual bumps B_p in a satisfactory manner. Moreover, the sparsely distributed “macro-bumps” are determined by

$$x_p B_p(u) := \sum_{q \in P_p} \tilde{x}_q \tilde{B}(u - q), \quad (3.17)$$

where x_p results from normalizing all the bumps B_p in ℓ_2 .

The second stage of the scheme consists therefore in identifying the “centers” p of the $\sum_{q \in P_p} \tilde{x}_q \tilde{B}(u - q)$, yielding our approximation of the positions in $\mathcal{P}$ and incidentally the coefficients x_p of the resulting macro-bumps.

The result of the first stage of our scheme is displayed in Figure 12.

Now the solution is much less sparse but the recovery schemes work much better due to the lower coherence of the more localized dictionary. For this less pronounced sparsity level NESTA offers an advantage in that the penalty parameter λ promoting sparsity is fixed throughout the computation. In SESOP, however, it is adjusted during the iteration, mostly in favor of a higher sparsity, which in this case may lead to misfits.

In Table 3 we present the relative ℓ_2 -errors $E(f)$ of the NESTA reconstruction of the Molybdenum Vanadate computer simulated image from Figure 9. The underlying numbers n of measurements range from 200 to 900. Note that the values of the concentration parameter σ now change with the number of measurements, reflecting the fact that a larger number of measurements allows us to handle a larger number of coefficients $\tilde{x}_q$ per cluster.

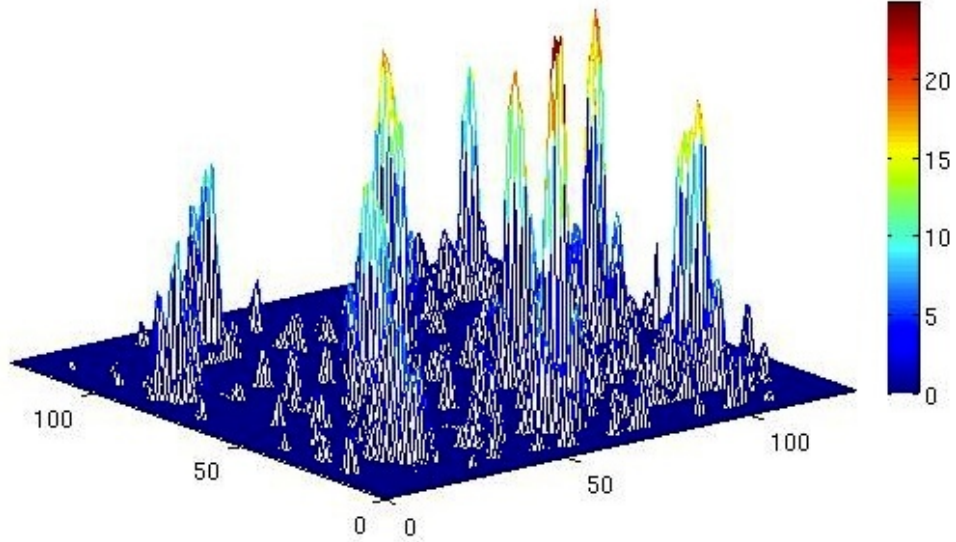


Figure 12: Recovery of the coefficients $\tilde{x}_p$ from (3.16) produced by NESTA for first stage of the Phase 2 experiment, based on 900 measurements with 1% additive noise.

Number of measurements		200	300	500	700	900
Concentration parameter		5.95	5.45	5.35	5.28	5.24
Added noise:	0%	6.89%	3.25%	1.60%	0.91%	0.75%
	1%	11.71%	8.43%	7.42%	5.57%	4.85%
	2%	16.36%	13.77%	12.16%	9.04%	7.50%
	3%	21.03%	16.85%	14.58%	11.85%	10.04%

Table 3: Relative ℓ_2 -errors $E(f)$ for first stage of the Phase 2 experiment for the NESTA recovery of a 128×128 high resolution computer simulated STEM image of Mo_5V_{14} -oxide, presented in Figure 9.

As explained above, image reconstructions already result from (3.16) computed at the first stage of the scheme. In Figure 13 we show the image, corresponding to the coefficients $\tilde{x}_q$, displayed in Figure 12. The concentration parameter of the bumps $\tilde{B}$ was set to 5.24, which produces intensity distributions with a slightly smaller diameter than the ones in the image, the relative ℓ_2 -error $E(f)$ is 4.85%.

The recovered approximate positions p of the atomic columns and the corresponding coefficients x_p are displayed in Figure 14. Here we cannot compare the recovered values x_p with those behind Figure 9 because they are not known to us. Comparing Figure 14 with Figure 11, we see, however, that the positions of the heavy atoms are recovered well while those of the oxygen atoms are less accurate or even missed.

In the above experiments the added noise represents the accumulated perturbation resulting from the summation $\sum_{P \in \mathcal{G}_H} \tilde{\phi}_{i,P} e_P$ in (3.7). Thus, so far, we have imposed that this accumulated noise is of a level up to 3% with which the recovery procedures can cope. We conclude Phase 2 now with a test of the accumulative effect on the final noise level by perturbing the individual pixel intensities by a considerably higher noise level. In the first

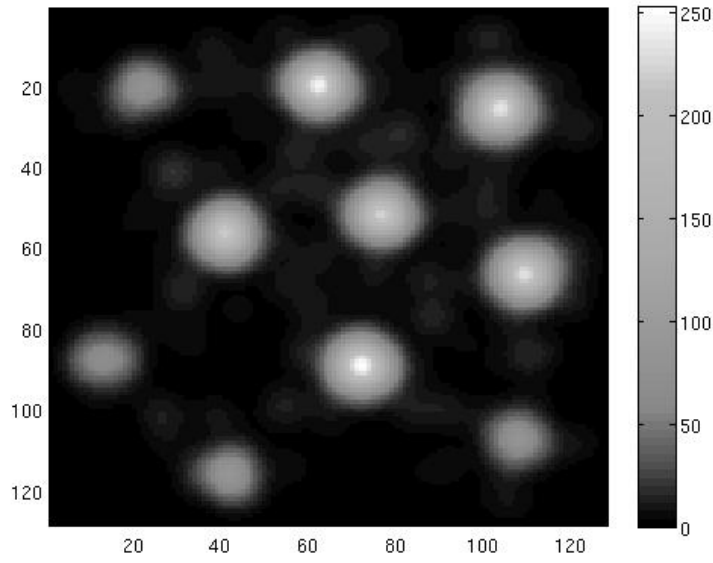


Figure 13: Recovery of the Mo_5V_{14} -oxide computer simulated image produced by NESTA, based on 900 measurements to which 1% Gaussian noise is added.

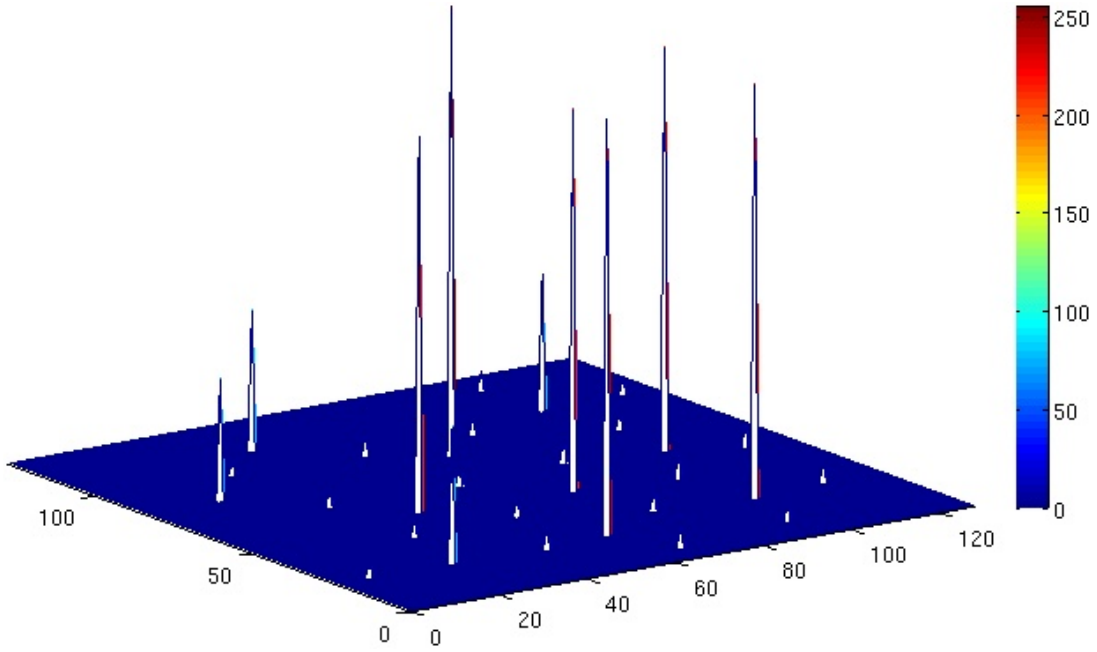


Figure 14: Recovery of the Mo_5V_{14} -oxide computer simulated image produced by NESTA, based on 900 measurements to which 1% Gaussian noise is added.

experiment the pixel intensities are perturbed by $\mathcal{N}(0, 3.2)$ and by Poisson pixel noise with parameter equal with the square root of the pixel intensity to a total noise level of 12.06% for individual pixels. The corresponding noise level of the measurements y_i turns out to be only of the level of 0.74%. The recovered image, based on 200 measurements,

is shown in Figure 15.

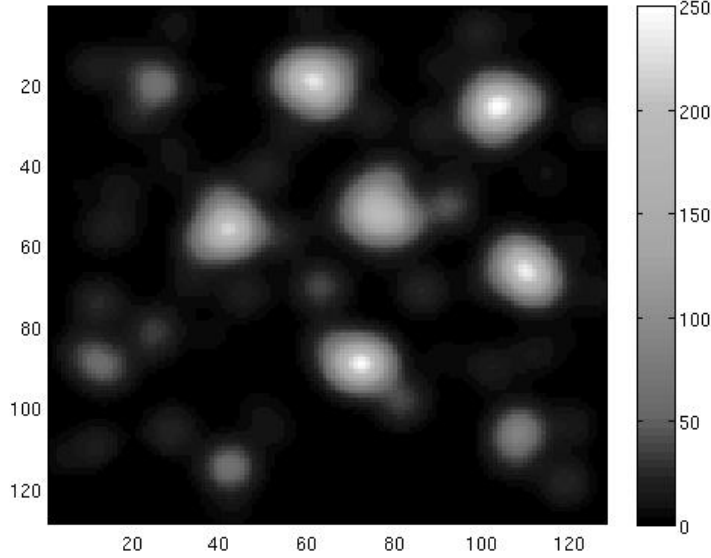


Figure 15: Recovery of the Mo_5V_{14} -oxide computer simulated image produced by NESTA, based on 200 measurements with pixel noise of level 12.06% corresponding to a noise level of 0.74% for the measurements.

The second experiment is analogous but involves a significantly higher level of pixel noise, namely 31.43% produced by Gaussian fluctuations from $\mathcal{N}(0, 12)$ and by Poisson noise with same variance as in the first experiment. The corresponding noise level of the measurements y_i turns out to be only of the level of 1.57%. The NESTA recovery from 900 measurements is shown in Figure 16.

To summarize our findings for the Phase 2 experiments, it should be emphasized that we are using a *flexible* representation of the bumps based upon local linear combinations of translations of the basic function $\tilde{B}$ with carefully chosen parameter σ . This method determines a good reconstruction of the intensity distributions around all the heavy Molybdenum and Vanadium atoms present in the computer simulated image, based on as little as 200 measurements and up to 3% additive noise. However, the correct recovery for the lighter Oxygen atoms succeeds only through a higher number of measurements depending on the added level of noise.

Phase 3: To test this concept further we use now an actual *micrograph* of the M1 phase of $Mo-V-Te-O$ catalyst with a calibration of 0.011796 nm/pixel. It should be emphasized that the purpose of the following experiment is not to improve on the quality of the micrograph, but to further validate the above CS-concepts.

In this third phase of our experiments, we address an additional effect related to spatial uncertainty during the process of data acquisition. Instead of using (3.7) as a measurement model we employ the following model. Let $\{f_p : p \in \mathcal{G}_h\}$ denote the pixel

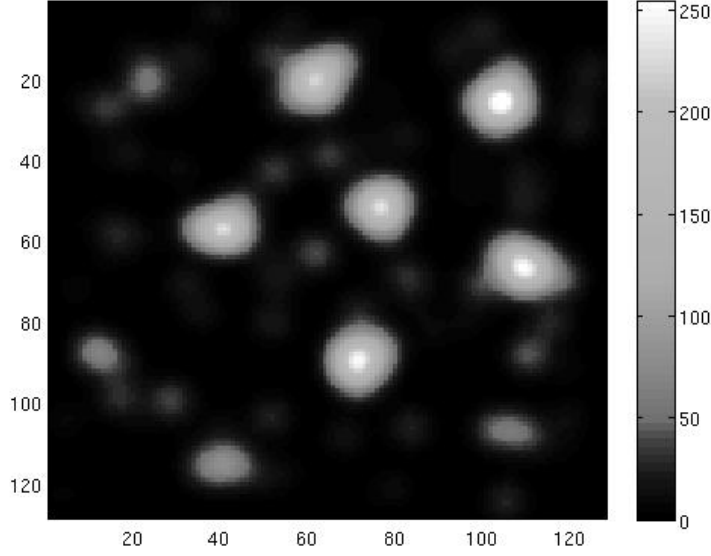


Figure 16: Recovery of the Mo_5V_{14} -oxide computer simulated image produced by NESTA, based on 900 measurements with pixel noise of level 31.43% corresponding to a noise level of 1.57% for the measurements.

intensities of the given 128×128 STEM image shown on Figure 17. Note that, although the micrograph is a high resolution image, the pixel intensities still provide noisy information due to distorted intensities and positions. Therefore, it makes no sense to try to reproduce the original image on $\mathcal{G}_h$ exactly. Instead we wish to see whether an aggregated CS-type data acquisition can extract a reasonably smoothed intensity distribution that still reflects relevant physical information about the material.

To this end, let for any z in the image domain $g(z)$ denote the local bilinear interpolant of the data f_p . Now set

$$y_i = \sum_{P \in \mathcal{G}_H} \tilde{\phi}_{i,P} g(P + s(i, P)), \quad (3.18)$$

where $s(i, P)$ is a random spatial fluctuation with mean $(0, 0)$. In our particular experiments the fluctuation is confined to the square region $[-(h + H/2), (h + H/2)]^2$. Note that the expectation of $\{g(P + s(i, P)) : P \in \mathcal{G}_H\}$ is a slightly blurred version of $\{f_p : p \in \mathcal{G}_h\}$ which is close to $\{f_P : P \in \mathcal{G}_H\}$. The corresponding expected image derived from the above random fluctuations is shown in Figure 18.

We explore the same methodology, described already in Phase 2, using flexible representations of the local bumps B_p via linear combinations of translations on $\mathcal{G}_h$ of an appropriately chosen basis function $\tilde{B}$. Applying now the NESTA recovery, based on 1200 measurements, to the data from (3.18) yields the image displayed in Figure 19. The window parameter for the Gaussian representing $\tilde{B}$ is set to $\sigma = 5.6$. The relative ℓ_2 -error between this recovery and the expected 64×64 image is 11.81%.

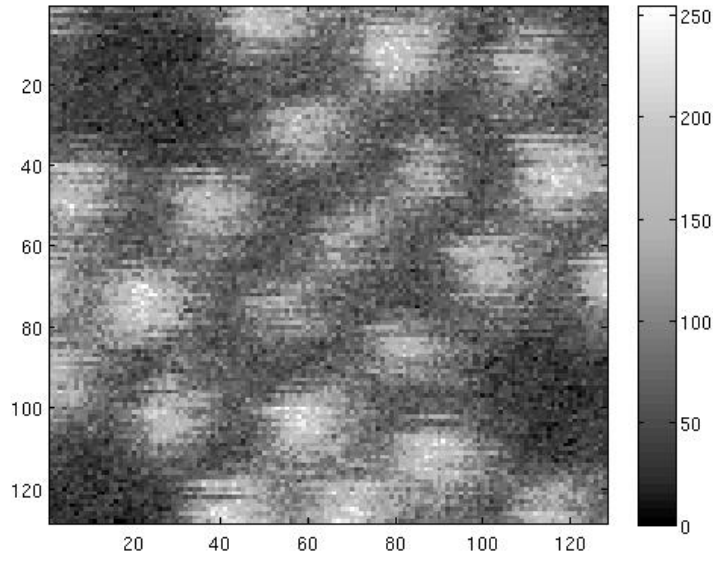


Figure 17: 128×128 patch from a micrograph of M1 phase catalyst Mo-V-Te-O used in the third phase of our experiments.

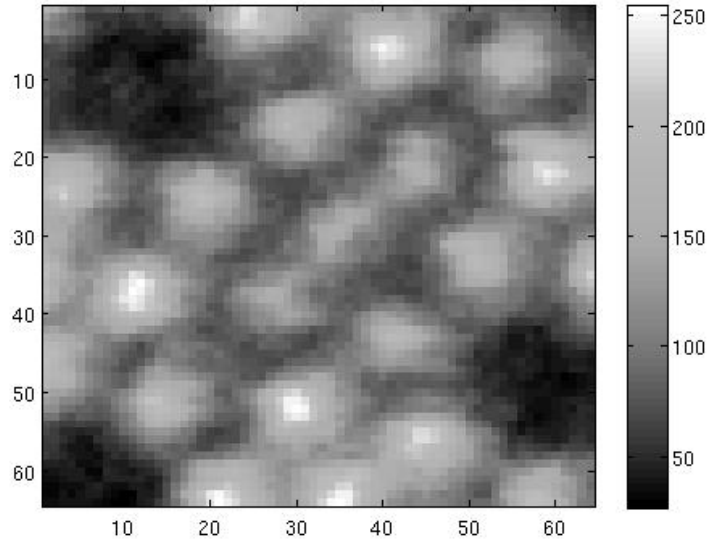


Figure 18: 64×64 image of the expected values from random fluctuations of the original micrograph in Figure 17.

Summary: Let us briefly summarize our findings as follows. The experiments seem to indicate that Compressed Sensing techniques allow one to exploit sparsity in the context of Model Class 1. Although the standard favorable assumptions on sensing matrices do not hold in this context, in absence of noise we have obtained *exact recovery* of the sparse

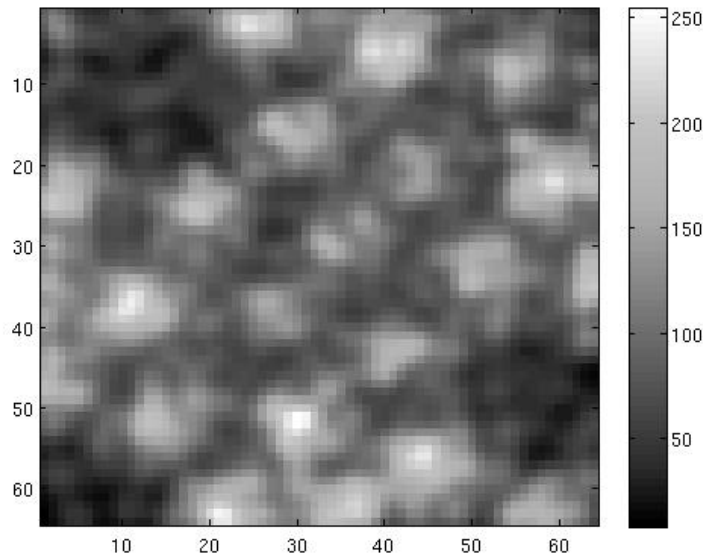


Figure 19: NESTA recovery for the Phase 3 experiment, based on 1200 measurements and $\sigma = 5.6$.

coefficients in the idealized model used in Phase 1. Adding noise to the measurements adversely affects the stable identification of the positions of atomic columns because the involved dictionaries are highly coherent. A certain remedy lies in splitting the recovery process into several stages, first lowering coherence at the expense of sparsity to identify energy clusters which are then further treated through a second local recovery stage.

Of course, in Phases 2 and 3 of our experiments we can no longer expect to have *exact recovery*. In addition to treating non-constant bump functions, at the end of Phase 2 we have also tested the effect of aggregated measurements on high noise levels in the individual pixel intensities. Then, in Phase 3 we have also tested our method on aggregated measurements from spatially perturbed samples from a high definition STEM image. These first still idealized experiments indicate the principal feasibility of this concept towards repeated very low dose measurements.

The algorithms we have tested prove to be quite robust in our experiments, but they are not designed yet to fully exploit some of the special features of our target images, namely the fact that the intensity distributions around the atomic positions are well separated and the fact that the peak intensities are samples of a quantized function of the corresponding atomic numbers. Further improvements are the subject of ongoing research.

3.2 Electron Tomography

A quite different, very promising use of HAADF-STEM concerns *electron tomography*, see e.g. [36] for a detailed discussion. Here, specimens of considerably larger size and thickness are explored. Now the intensity values returned by the instrument are viewed

as integrals of the Z^2 distribution along the ray taken by the electron beam. The objective is to reconstruct the 3D-structure of the material under investigation.

It should be stressed that in such applications the target resolution is significantly lower than the one in the previous application. In particular, this means that the diameter of the electron beam is smaller than the size of the pixel to which it corresponds. Among the various possible questions to be asked in this context we shall focus in the sequel on the following scenario. Clumps of heavier atom clusters are embedded in some carrier material and one is interested in the distribution, size, and geometric shape of the clumps, see Figure 5. As pointed out before, the atomic structure of the carrier material is far from resolved. It therefore appears more like a gray soup similar to noise. While many of the clumps stand out clearly, some of them, depending on the projection direction, nearly merge with the soup.

Let us sketch the typical setup for this type of data acquisition in Figure 20:

- The electron gun scans the object in the usual STEM way rastering along a Cartesian grid which for simplicity of exposition is scaled here to discretize the unit square. The stepsize in the rastering process is denoted by H which is therefore also the resolution of the resulting 2D-images of density patterns received for each tilt angle. Recall that the diameter of the electron beam is (significantly) smaller than H which will be the main discretization parameter below.
- We adopt the convention that the scanning direction is parallel to the x -axis in the raster coordinate system. We are thinking of low dose measurements.
- The specimen is fixed to a holder handle that can be tilted around a (fixed) axis which is parallel to the y -axis in the raster plane and perpendicular to the optical axis (assumed as usual pointing along the z -direction).
- We are confined to a fixed tilt range ($\pm 60^\circ$, say), due to instrumental limitations for possible holder rotations and due to the fact that for longer paths through the specimen there will be ray deviations and interference. (In some cases it is possible to avoid both these limitations by preparing a specimen with a cone shape but the problem of unavailable observations at certain tilt angles cannot be dismissed, in general.) Let us suppose that

$$\theta_i, \quad i = 1, \dots, n_a,$$

are the (known equispaced) tilt angles.

Since the width of the electron beam is small compared to H it is justified to view the measurements as integrals over rays $R_{k,j}(\theta_i)$ corresponding to the gun position j in the k th scanning row and tilt angle θ_i . Thus, each fixed tilt angle θ_i and scanning row k , corresponding to a slice through the specimen at position k , produce the density integrals

$$d_{k,j}(\theta_i) := \int_{R_{k,j}(\theta_i)} f(\mathbf{x}) ds, \quad j = 1, \dots, M, \quad (3.19)$$

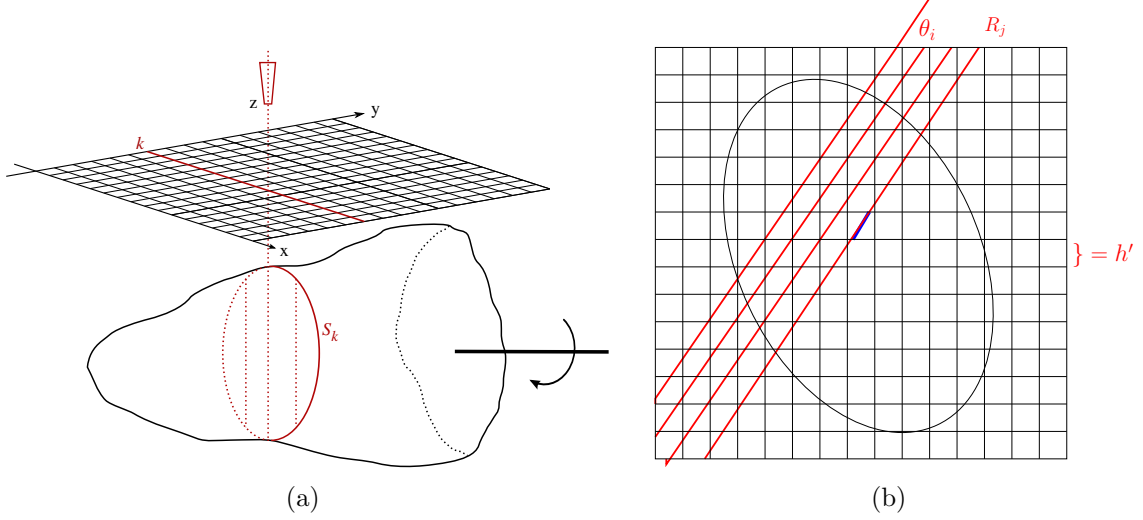


Figure 20: (a) The scanning plane (x, y) with the line k in red and the corresponding slice S_k from the specimen; (b) Family of parallel rays R_j for the tilt angle θ_i in the (ξ, η) -unit square for the slice S_k .

for the slice

$$S_k = \text{supp} f \cap \{\mathbf{x} : x = kH\}.$$

These slices thus correspond to a single family of parallel planes which, of course, would not provide enough information for a 3D-inversion (even if additional few directions of the tilt axis were added as in the “dual axis method”).

So, as a feasible principal strategy, the following commonly used two-stage process suggests itself:

- (1) For each slice S_k reconstruct the density distribution $f_k := f|_{S_k}$ from the ray data $d_{k,j}(\theta_i)$ on $R_{k,j}(\theta_i)$, $i = 1, \dots, n_a$, $j = 1, \dots, M$.
- (2) Then “stack” the slice distributions f_k together (e.g. by interpolation or more general fitting technique in a way to be discussed) to obtain an approximation to f .

In principle, this is a standard procedure and elaborate inversion schemes for various different types of tomographic applications are available. What hampers the straightforward use of such schemes in the present context, however, is first the “missing wedge” problem due to the restricted tilt range. Standard backprojection gives rise to severe artifacts in the reconstructed images.

The second issue concerns again dose. The coarser the angular resolution the less dose is applied to the specimen and the question arises how coarse the resolution can be kept while still reconstructing the structure of interest.

We shall now concentrate on stage (1). We fix the slice S_k and therefore suppress the index k in what follows. For convenience, we shall choose now a fixed (ξ, η) -coordinate system for the unknown density f in the plane of the slice S_k which is perpendicular to the tilt axis. The axis ξ is perpendicular to the electron beam at the initial tilt angle $\theta = 0$ oriented towards scanning direction and then the axis η is perpendicular to ξ , i.e.

parallel to the electron beam, and oriented towards the electron gun. For simplicity, we assume that the scanning area is the (ξ, η) -unit square $[0, 1]^2$ and the tilt axis projects to the point $(1/2, 1/2)$. In fact, the slice is thin enough so that, when rotated within the fixed angle range, the investigated area of the specimen falls always fully into the scanning unit square. Let us denote by $g(\xi, \eta)$ the unknown density of f restricted to the slice S under consideration in (ξ, η) coordinates. Due to our assumptions we may set $g(\xi, \eta) = 0$ if $(\xi, \eta) \notin [0, 1]^2$.

In this coordinate system the rays can be parameterized as

$$R_j^\theta := R_{k,j}(\theta) = \{(\xi, \eta) \in [0, 1]^2 : \xi = t_j + \eta \tan \theta\}, \quad j = 1, \dots, M \quad (3.20)$$

where t_j is the intersection of the ray R_j^θ with the ξ -axis. Therefore, one obtains

$$d_j^\theta := d_{k,j}(\theta) = \int_{R_j^\theta} g(\xi, \eta) ds = \frac{1}{\cos \theta} \int_0^1 g(t_j + \eta \tan \theta, \eta) d\eta, \quad j = 1, \dots, M. \quad (3.21)$$

see Figure 20(b). (In the last integral one could adjust the bounds for η to indicate the values at which the ray enters and leaves the unit square.)

To recover g from the data d_j^θ , one can, in principle, pursue two different strategies, namely

(I) applying the so called “algebraic” reconstruction technique (ART), or

(II) going through the Fourier-Slice-Theorem.

For a good distribution of rays (II) seems to be preferable since the FFT helps efficient computing. Recent relevant developments of reconstruction techniques based on a particularly adapted version of the Fast Fourier Transform in combination with regularization can be found in [26]. However, since in the given setting the Fourier data is incomplete due to the missing wedge, we shall concentrate here on (I). It basically does the following: the ray integrals are replaced by sums of weighted values of g on the given ray, where the weight reflects the contribution of the ray to the cell on which the unknown value is supposed to approximate g . Then, one formally obtains a linear system of equations in the unknown discrete approximations to g . Note that in our case this system (as in many other practical situations) will be underdetermined and most likely inconsistent.

The currently used discretizations all seem to fit into the following setting. Consider a Cartesian grid

$$\mathcal{G}_{h'} = \{\square_\alpha = [\alpha_1 h', (\alpha_1 + 1)h'] \times [\alpha_2 h', (\alpha_2 + 1)h'], \alpha \in \mathbb{Z}^2 : 0 \leq \alpha_1, \alpha_2 < N = 1/h'\}$$

where h' is the pixel width and a basis

$$\mathcal{B}_{h'} = \{B_\alpha : \square_\alpha \in \mathcal{G}_{h'}\}, \quad \|B_\alpha\|_{L_2} = 1, \quad \alpha \in \mathcal{G}_{h'},$$

where the basis functions that will be used to discretize g are normalized in $L_2([0, 1]^2)$. One cannot recover g from the finitely many measurements $d_j^{\theta_i}$, $i = 1, \dots, n_a$, $j = 1, \dots, M$. Instead, we can try to recover an approximation

$$g_{h'} = \sum_{\square_\alpha \in \mathcal{G}_{h'}} c_\alpha B_\alpha. \quad (3.22)$$

In fact, defining

$$w_{(\theta^i, j), \alpha} := \int_{R_j^{\theta^i}} B_\alpha(\mathbf{x}) ds = \frac{1}{\cos \theta_i} \int_0^1 B_\alpha(t_j + \eta \tan \theta_i, \eta) d\eta, \quad (3.23)$$

we arrive at the following system of linear equations

$$\sum_{\alpha \in \mathcal{G}_{h'}} w_{(\theta^i, j), \alpha} c_\alpha = d_j^{\theta^i}, \quad i = 1, \dots, n_a, \quad j = 1, \dots, M = 1/H, \quad (3.24)$$

in the unknown coefficients $c_\alpha, \alpha \in \mathcal{G}_{h'}$ which we'll abbreviate as

$$\mathbf{W}\mathbf{c} = \mathbf{d}. \quad (3.25)$$

Before discussing how to solve (3.25) we need to relate the approximation mesh size h' and the choice of the B_α to the given data, i.e. to the scanning step size H and the angular resolution. Since we aim at using a possibly small number n_a of tilt angles, although the scanning step size H is not too small, we expect that $n_a \ll M$. On the other hand, we would like to have a relatively good approximation by the ansatz (3.22), i.e. we would like to have h' as small as possible. Therefore, we will have $\#\mathcal{G}_{h'} = (N+1)^2 > n_a(M+1)$, even if $h' \sim H$ for reasonable constants. Thus, in any case of interest to us, the system (3.25) will be underdetermined.

The standard procedure to “invert” such systems, is the *Kaczmarz-Iteration* or one of its numerous variants. It views the solution (if it exists) as the intersection of hyperplanes (each given by one equation in the system) and projects current approximations successively onto these hyperplanes. Obviously, when the hyperplanes have big mutual angles this converges rapidly. If they are nearly parallel, the convergence becomes very slow. The reason for the popularity of this scheme is that it copes reasonably well with ill-conditioned systems and that as a “row-action” algorithm it exploits the sparsity of the system. At least from a theoretical point of view there is a problem, however, that in the inconsistent case (i.e., if there exists no solution of the system due to measurement noise) the iterations do not converge to a least squares solution of the system and one has to be content if the resulting image is visually satisfactory.

Here we want to propose an alternative to Kaczmarz' algorithm which is more in the spirit of the above mentioned treatment of the Logan-Shepp phantom (see Section 2.5). In order to be able to invert the corresponding ill-posed problem (3.25), we need a suitable regularization which, in turn, should be based on a proper sparsity model for g .

Trying to reconstruct the shape and position of heavy atom clumps corresponds to reconstructing the position and shape of higher “intensity islands”. Hence, for each slice one is looking for a piecewise constant of variable shape and height values but relatively small diameters. Nevertheless, for diameters of $10h'$ or more, it seems reasonable to take

$$B_\alpha = \frac{1}{h'} \chi_{\square_\alpha}. \quad (3.26)$$

In this case, the weights are proportional to the lengths of the intersections of the rays with the cells in the Cartesian grid:

$$w_{(\theta^i, j), \alpha} = \frac{|R_j^{\theta^i} \cap \square_\alpha|}{h'}. \quad (3.27)$$

Our first experiments will be concerned with this latter piecewise constant setting. Then, to recover a piecewise constant, a reasonable regularization should be

$$\min \{ \|\mathbf{c}\|_{TV} : \text{subject to } \|\mathbf{W}\mathbf{c} - \mathbf{d}\|_{\ell_2} \leq \epsilon \}, \quad (3.28)$$

where ϵ is the noise level and $\|\cdot\|_{TV}$ is a discrete total variation. Equivalently this problem can be formulated as (unconstrained) minimization problem

$$\min_{\mathbf{c}} \left\{ \frac{1}{2} \|\mathbf{W}\mathbf{c} - \mathbf{d}\|_{\ell_2}^2 + \mu \|\mathbf{c}\|_{TV} \right\}, \quad (3.29)$$

where the parameter μ is related to the noise level ϵ . These formulations relate to (2.19) and (2.20), just with the ℓ_1 -norm replaced by a TV-norm.

In the literature there are several ways to define the total variation. The definition most commonly used is

$$\|\mathbf{c}\|_{TV} := \sum_{\alpha \in \mathcal{G}_{h'}} \sqrt{(c_{\alpha+e^1} - c_{\alpha})^2 + (c_{\alpha+e^2} - c_{\alpha})^2}, \quad (3.30)$$

but one could also think of taking

$$\|\mathbf{c}\|_{TV} := \sum_{\alpha \in \mathcal{G}_{h'}} |c_{\alpha+e^1} - c_{\alpha}| + |c_{\alpha+e^2} - c_{\alpha}|. \quad (3.31)$$

In summary, the following questions will serve as guidelines for further investigations.

- In order to resolve the shape of the islands well, it would be desirable to have $h' \leq H$ as small as possible. How small can h' be chosen relative to H ? Of course, when h' gets too small compared with H , it could happen that some pixels are missed by all rays. Nevertheless, the TV-penalization would try to keep perimeters small, so that one may still be able to recover the correct shapes.
- It is not clear how the matrices $\mathbf{W}$ cooperate with TV -minimization. Many methods for the solution of the problem (3.28) (or (3.29)) have been developed under the assumption that the matrix W fulfills the restricted isometry property, but in the ART the matrix W is typically sparse (since each ray hits roughly the order of $1/h'$ pixels), ill-conditioned, and most often severely rank-deficient which even makes solving the standard least squares challenging, as explained above.

Since the efficient solution of problems like (3.29) for large data sets is the topic of current research, but satisfactory algorithms are not available yet, we shall address these questions first by some experiments with relatively small synthetic data. Nevertheless, the following studies should suffice at this point to provide a proof of concept. Specifically, we have used the matlab package NESTA, which, in principle, is able to solve the problem (3.29) using the total variation given by 3.30. However, since the computational performance of this method is relatively slow, we restricted ourselves to an image size of 64×64 pixels and only tried the Logan-Shepp phantom with a resolution of 128×128 .

3.2.1 Example 1

The first example is very simple and is used to demonstrate the difference between a linear reconstruction (based on Kaczmarz-iterations) and the nonlinear TV-reconstruction based on (3.29). As a “phantom” that we want to reconstruct from simulated data (i.e., from precomputed exact values of the ray integrals 3.19) we take a white square covering the region $[5/16, 5/8] \times [5/16, 5/8]$ on black ground covering the unit square $[0, 1]^2$. Here “white” means a grayscale value of 200 and “black” a grayscale value of 10. We discretize the image with 64×64 pixels (i.e., $h' = 1/64$) which allows an exact representation of the phantom.

We begin with a parameter study for the noise-free case and first examine the question how many measurements one needs to recover the phantom well with Kaczmarz iterations and with TV -regularization, see Figure 21. We find that the ℓ_2 -reconstructions computed with Kaczmarz’ algorithm show strong artifacts and become completely meaningless, if one diminishes the number of measurements. The TV-reconstruction is almost exact even for relatively small number of measurements and only starts to blur, if many pixels in the image are not hit by a ray any more.

We repeat the above experiment but add some noise to the right hand side $\mathbf{d}$. More precisely, we used the following model for the noise:

$$\tilde{d}_l^{\theta_i} = d_l^{\theta_i} + e_l^{\theta_i}, \quad e_l^{\theta_i} := l(R_l^{\theta_i}) \cdot N(0, \sigma), \quad (3.32)$$

where $l(R)$ is the length of the intersection of the ray R with the unit square, and $N(0, \sigma)$ is the Gaussian normal distribution with mean 0 and standard deviation σ . The idea behind this model is that we assume that most of the noise is caused by the ray crossing the soup surrounding the cluster.

Table 4 lists some information about the setup of these experiments and the numerical properties of the system matrix W . Clearly decreasing the number of measurements does not only make the problem smaller, but also improves the condition of the matrix. Note that the image has 4096 pixels, so that the system is severely underdetermined.

n_a	$1/H$	Θ_{max}	#Eq	Ran(W)	$e_{rel, \sigma=2}(\mathbf{d})$	$\mu_{opt, \sigma=2}$	$e_{rel, \sigma=4}(\mathbf{d})$	$\mu_{opt, \sigma=4}$
20	32	60	840	828	4.60e-02	0.3	9.18e-02	0.4
20	16	60	422	418	4.71e-02	0.7	9.51e-02	0.4
10	10	60	132	132	5.03e-02	0.03	10.1e-02	0.07
5	5	60	33	33	5.20e-02	0.0007	10.5e-02	0.01

Table 4: Parameters for the reconstructions in Figures 21 and 22. #Eq denotes the number of rows of the matrix W which equals the number of measurements; $e_{rel, \sigma}$ is the relative error $\|\tilde{\mathbf{d}} - \mathbf{d}\|/\|\mathbf{d}\|$ of the right hand side due to the added noise with standard deviation σ according to (3.32); μ_{opt} is the regularization parameter which delivered the optimal reconstruction for the given noise level.

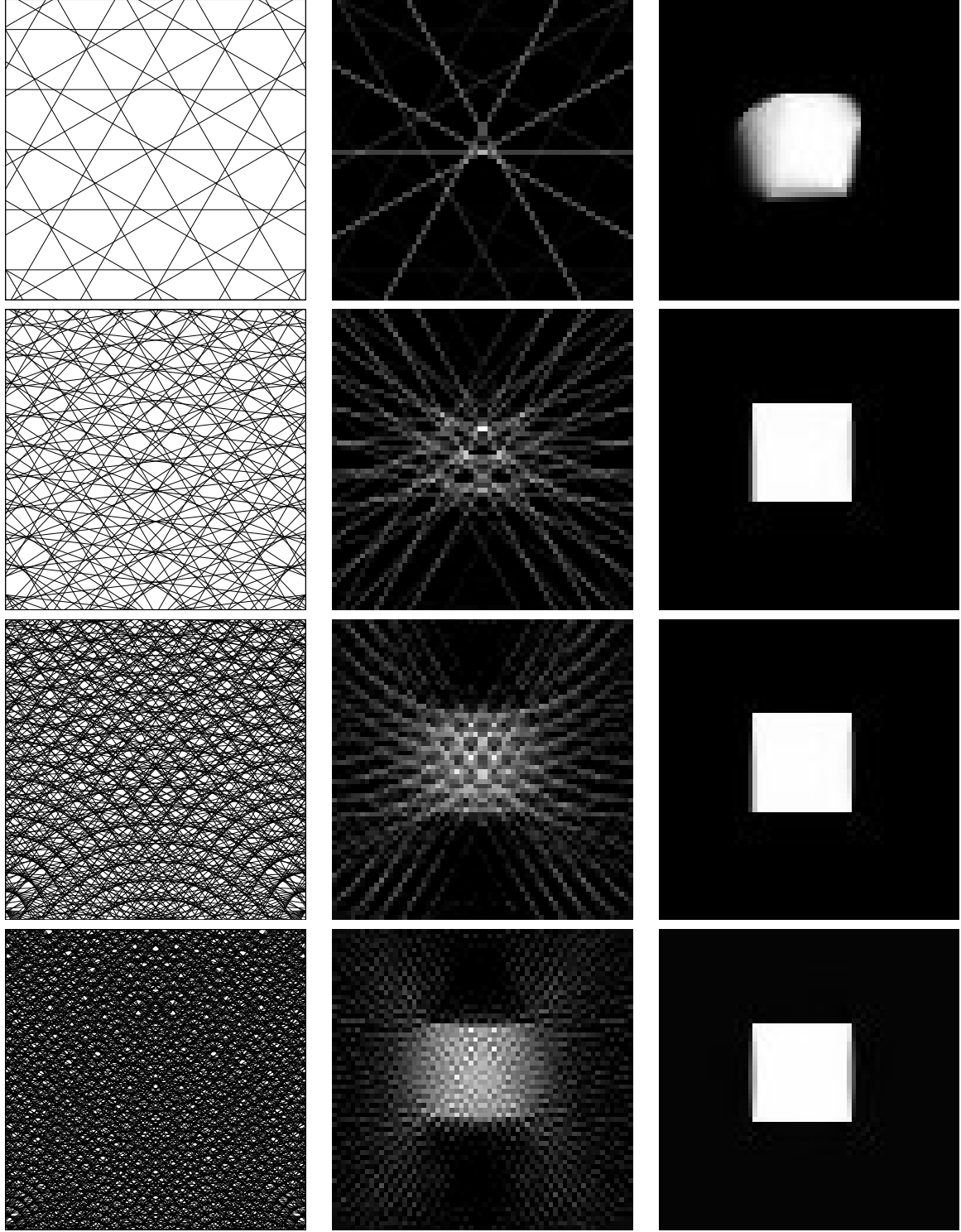


Figure 21: Reconstructions from noise free data with a maximum tilt angle of 60° . First column: ray trace diagrams of the measurements taken; Second column: least-squares reconstruction computed with Kaczmarz iterations; Third column: TV-regularized reconstruction; First row: $n_a = 5$, $H = 1/5$; Second row: $n_a = 10$, $H = 1/10$; Third row: $n_a = 20$, $H = 1/16$; Fourth row: $n_a = 20$, $H = 1/32$.

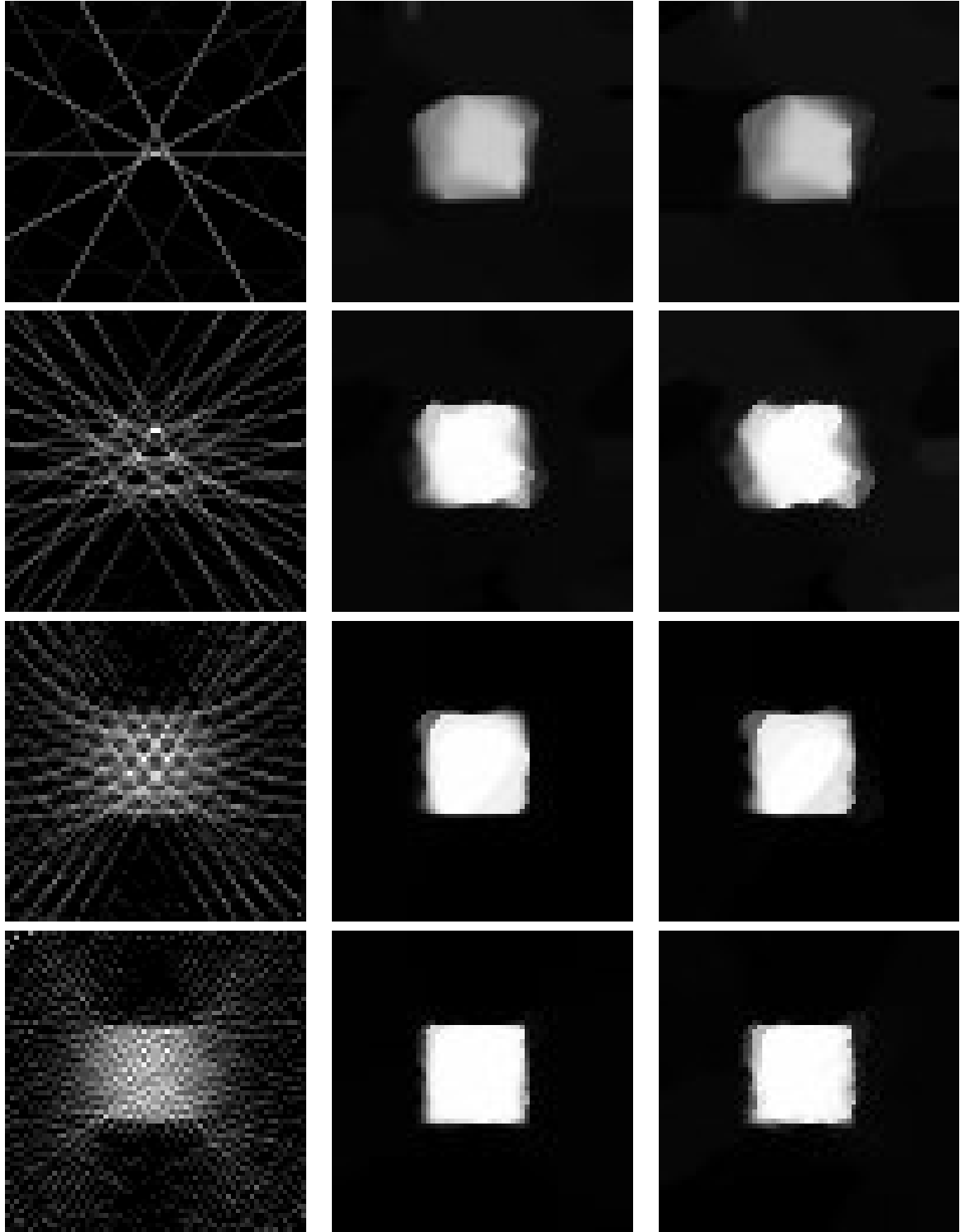


Figure 22: Reconstructions from noisy data with a maximum tilt angle of 60° . First column: least-squares reconstruction computed with Kaczmarz iterations ($\sigma = 4$); Second and third column: TV-regularized reconstruction for $\sigma = 2$ and $\sigma = 4$, respectively; First row: $n_a = 5$, $H = 1/5$; Second row: $n_a = 10$, $H = 1/10$; Third row: $n_a = 20$, $H = 1/16$; Fourth row: $n_a = 20$, $H = 1/32$.

3.2.2 Example 2

Now we turn to a more complicated example, where four clusters of different shape and size, two of them non-convex, have to be reconstructed. In this case the effect of the missing wedge can be well observed, as explained in Figures 23 and 24. As a general rule it seems, that the reconstruction of longer edges parallel to the axis of the missing wedge is a serious problem.

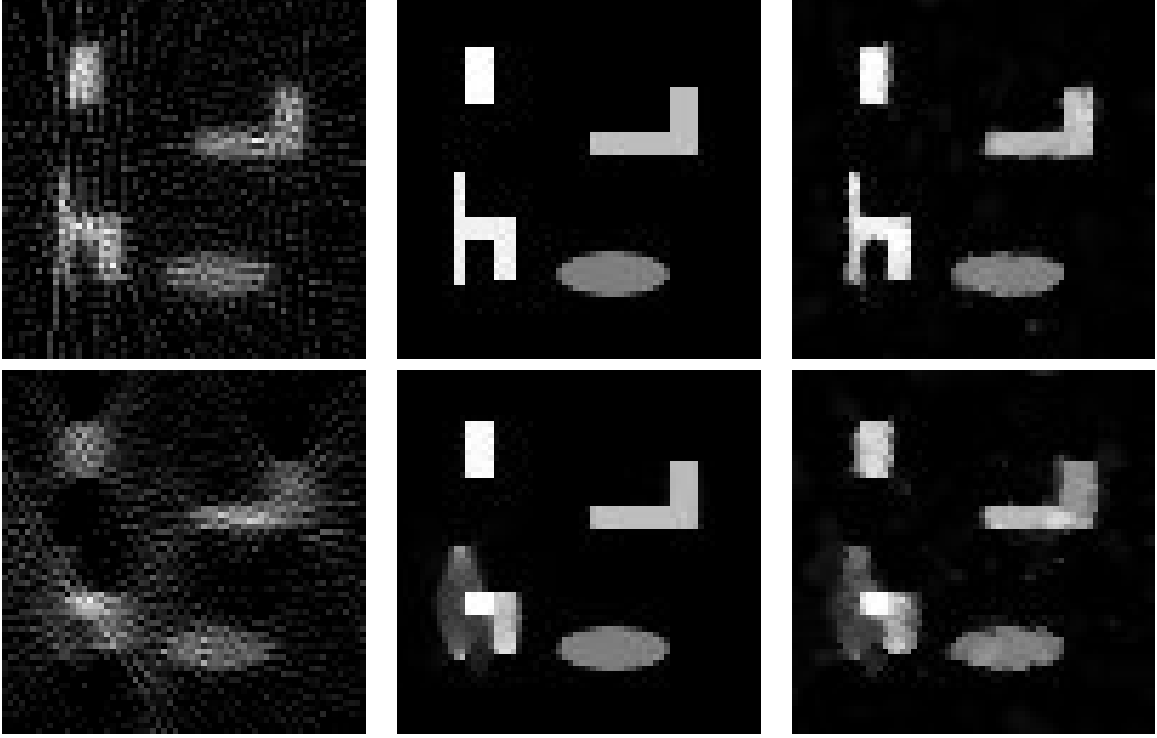


Figure 23: Reconstruction of Phantom 2 for $n_a = 20$, $H = 1/32$. Top row: Maximum tilt angle $\Theta_{max} = 85^\circ$; Bottom Row: $\Theta_{max} = 60^\circ$; First and second column: Kaczmarz and TV-reconstruction from noise-free data; Third column: TV-reconstruction from noisy data ($\sigma = 2$, $e_{rel}(\mathbf{d}) = 0.077$), where $\mu_{opt} = 0.03$ and 0.02 for $\Theta_{max} = 85^\circ$ and $\Theta_{max} = 60^\circ$, respectively.

In fact, as seen in Figure 23, whereas the reconstruction for the larger tilt angle is almost perfect, one observes that the vertical edges of the h-shaped object blur in case of a restricted tilt angle. These are the edges to which no parallel measurements are taken.

As we see in Figure 24, after a rotation, in case of a restricted tilt angle, the h-shaped object is reconstructed almost perfectly but the vertical edges of the l-shaped object blur, in particular in the presence of noise.

In the noisy case it is important to choose an appropriate value for μ . Unfortunately, this value seems to depend not only on the noise level and the discretization parameters, which are known beforehand or could at least be estimated. In the shown experiments we determine the optimal μ by minimizing the ℓ_2 -distance between the reconstruction and the original image. In this particular experiment we find different values for μ , although

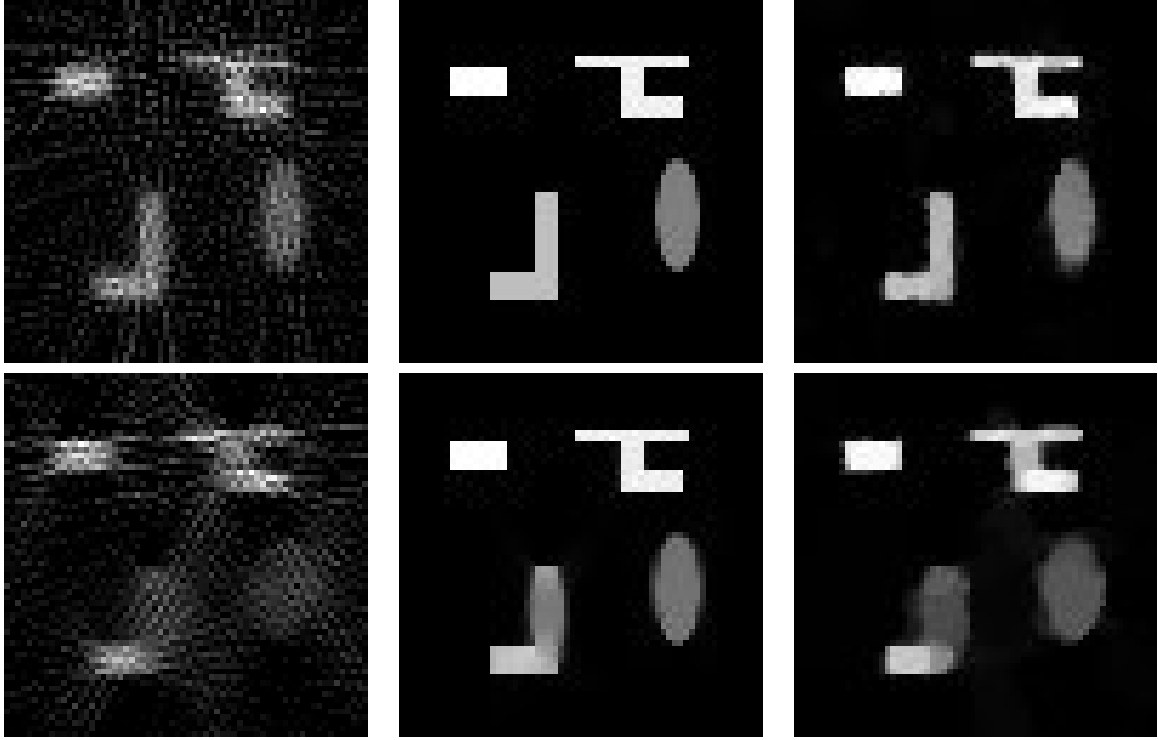


Figure 24: Reconstruction of a rotated Phantom 2. Same parameters as in Figure 23 for $n_a = 20$, $H = 1/32$. Top row: $\Theta_{max} = 85^\circ$; Bottom Row: $\Theta_{max} = 60^\circ$; First and second column: Kaczmarz and TV-reconstruction from noise-free data; Third column: TV-reconstruction from noisy data ($\sigma = 2$, $e_{rel}(\mathbf{d}) = 0.077$), where $\mu_{opt} = 0.07$ and 0.12 for $\Theta_{max} = 85^\circ$ and $\Theta_{max} = 60^\circ$, respectively.

the only difference between the images is the rotation. On the other hand the optimal μ usually needs to be determined only approximatively, because the reconstructions are very similar for a wide range of values.

3.2.3 Logan-Shepp Type Phantom

Finally, with reference to Figure 6, we have computed several 128×128 reconstructions for a Logan-Shepp type phantom (Logan-Shepp phantom modulo contrast change), although this is actually a little bit outside the scope of our application. In [8] a Fourier technique was used to reconstruct the image, however, in contrast to the present situation for the full angular range of 180° . The effect of the missing wedge is important in the present case though, because the reconstruction of the skull seriously deteriorates in the direction of the symmetry axis of the missing wedge. From Figures 25,26 it becomes clear that $n_a = 20$ projections are not enough, but $n_a = 40$ projections suffice to get an almost perfect reconstruction if the tilt angle is not restricted too much. If the tilt angle is restricted to 60° , then artifacts at the lower and upper part of the skull appear. This consequently also affects the interior reconstruction, such that the small ellipses in the lower part of the brain are hardly recognized any more. Furthermore, this phantom is

very sensitive to noise and the finer structures quickly become unrecognizable for both the ℓ_2 and TV reconstruction. Note that the image consists of 16384 pixels. For $n_a = 20$, $H = 1/32$ we perform 840 measurements, whereas in the case $n_a = 40$, $H = 1/64$ the matrix W has 3348 rows.

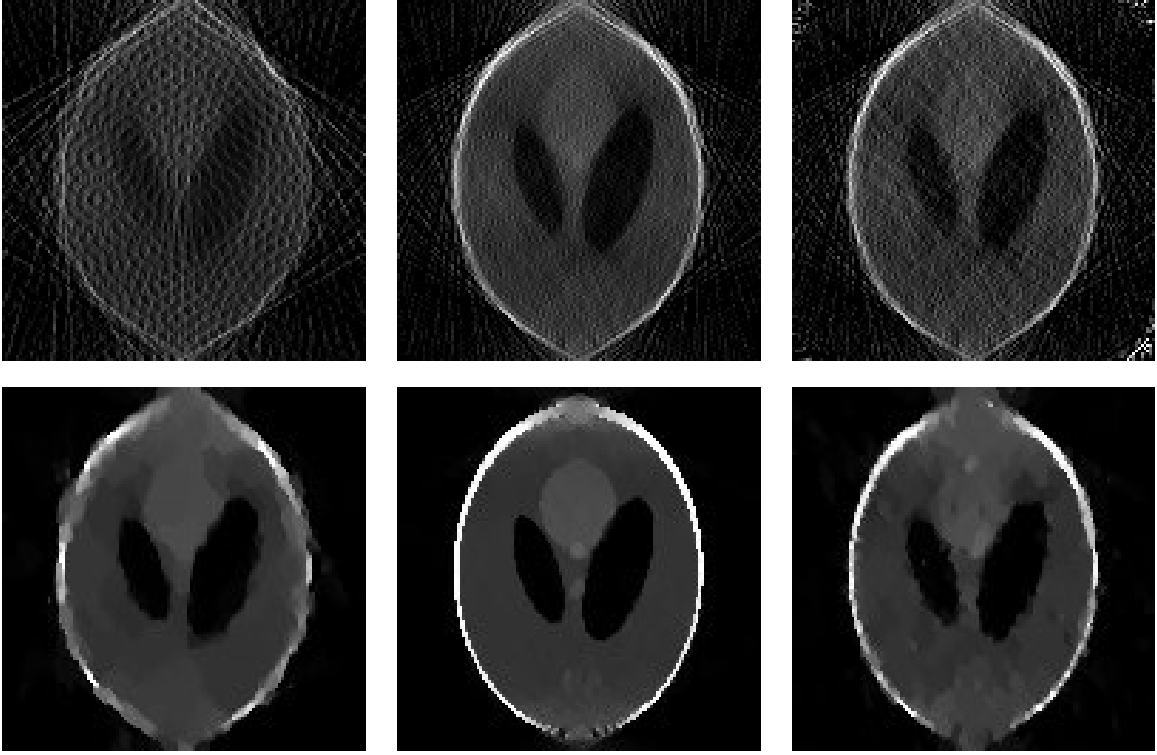


Figure 25: Reconstructions of the Logan-Shepp Phantom for a maximum tilt angle of $\Theta_{max} = 60^\circ$. First column: $n_a = 20$, $H = 1/32$, no noise added; Second and third column: $n_a = 40$, $H = 1/64$, without noise and with noise ($\sigma = 2$, $e_{rel}(\mathbf{d}) = 0.053$), respectively; First row: Kaczmarz reconstructions; Second row: TV -reconstructions (with $\mu = 0.0001$, $\mu = 0.001$, and $\mu = 0.05$, from left to right).

4 Conclusions

We have briefly summarized some mathematical foundations of Compressed Sensing from a perspective that, in our opinion, is relevant for developing new imaging concepts in the context of electron microscopy, with special emphasis on HAADF-STEM and electron tomography. To substantiate this claim we have discussed two application scenarios concerning STEM. The main objective in both cases is to argue the principal suitability of CS-concepts which requires identifying the notions of sparsity and measurements in this context. Moreover, we have outlined first steps towards possible solution strategies, identifying the arising key tasks and obstructions, illustrated by several experiments. More detailed presentations of corresponding findings are deferred to forthcoming papers.

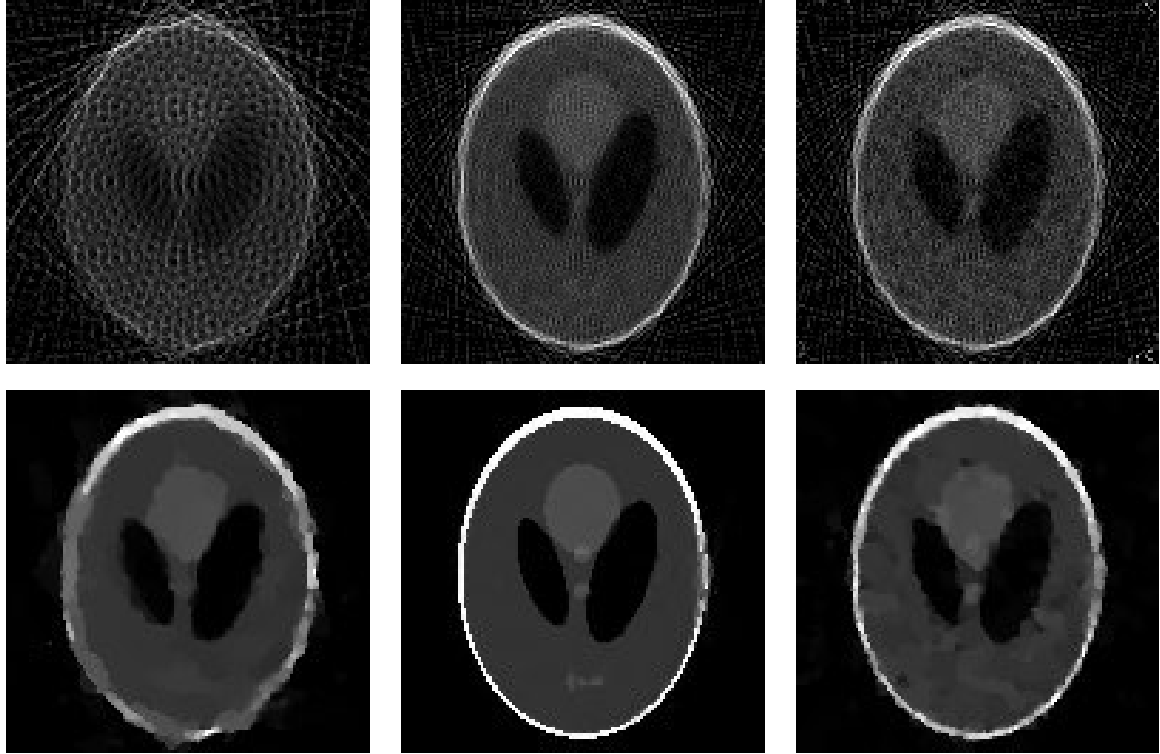


Figure 26: Reconstructions of the Logan-Shepp Phantom for a maximum tilt angle of $\Theta_{max} = 85^\circ$. First column: $n_a = 20$, $H = 1/32$, no noise added: Second and third column: $n_a = 40$, $H = 1/64$, without noise and with noise ($\sigma = 2$, $e_{rel}(\mathbf{d}) = 0.053$), respectively; First row: Kaczmarz reconstructions; Second row: TV-reconstructions (with $\mu = 0.0001$, $\mu = 0.001$, and $\mu = 0.05$, from left to right).

Acknowledgement: We are very indebted to Doug Blom and Sonali Mitra for providing us with several STEM simulations without which we would not have been able to validate the algorithmic concepts. Numerous discussions with Tom Vogt and Doug Blom have provided us with invaluable sources of information without which this research would not have been possible. We are also very grateful to Nigel Browning for providing tomography data. We would also like to thank Andreas Platen for his assistance in preparing the numerical experiments.

References

- [1] N. Baba, K. Terayama, T. Yoshimizu, N. Ichise, N. Tanaka, *An auto-tuning method for focusing and astigmatism correction in HAADF-STEM, based on the image contrast transfer function*, Journal of Electron Microscopy (3)**50**(2001), 163–176.

- [2] P.E. Batson, N. Dellby, O.L. Krivanek, *Sub-angstrom resolution using aberration corrected electron optics*, Nature **418**(2002), 617–620.
- [3] S. Becker, J. Bobin, and E. Candes, *NESTA: A Fast and Accurate First-order Method for Sparse Recovery*, CalTech Technical Report (2009).
- [4] P. Binev, F. Blanco-Silva, D. A. Blom, W. Dahmen, P. Lamby, R. Sharpley, T. Vogt, *Super-Resolution Image Reconstruction by Nonlocal-Means applied to HAADF-STEM*, chapter in this volume.
- [5] L. Bregman, *The relaxation method of finding the common points of convex sets and its application to the solution of problems in convex programming*, USSR Comput. Math. Math. Phys. **7**(1967), 200–217.
- [6] E. Candès and J. Romberg, ℓ_1 -MAGIC: *Recovery of Sparse Signals via Convex Programming*, available at <http://www.acm.caltech.edu/l1magic/> (2005).
- [7] E. Candès, J. Romberg, T. Tao, *Stable signal recovery from incomplete and inaccurate measurements*, Comm. Pure and Appl. Math. **59**(2006), 1207–1223.
- [8] E. Candès, J. Romberg, T. Tao, *Robust Uncertainty Principles: Exact Signal Reconstruction From Highly Incomplete Frequency Information*, IEEE Transactions on Information Theory **52**(2006), 489–509.
- [9] E. Candès and T. Tao, *Decoding by linear programming*, IEEE Trans. Inf. Theory **51**(2005), 4203–4215.
- [10] A. Chambolle, R. DeVore, B. Lucier, Y. Lee, *Nonlinear wavelet image processing: Variational problems, compression, and noise removal through wavelet shrinkage*, IEEE Image Processing **7**(1998), 319–335.
- [11] A. Cohen, W. Dahmen, I. Daubechies, R. DeVore, *Tree Approximation and Encoding*, ACHA **11**(2001), 192–226.
- [12] A. Cohen, W. Dahmen, R. DeVore, *Compressed sensing and best k -term approximation*, J. Amer. Math. Soc. **22**(2009), 211–231.
- [13] A. Cohen, W. Dahmen, R. DeVore, *Instance Optimal Decoding by Thresholding in Compressed Sensing*, in: Proceedings El Escorial 08, Contemporary Mathematics 2009.
- [14] R. DeVore, I. Daubechies, M. Fornasier, S. Güntürk, *Iterative re-weighted least squares*, Comm. Pure Appl. Math. (1) **63**(2010), 1–38.
- [15] R. DeVore, B. Jawerth, B. Lucier, *Image compression through transform coding*, IEEE Proceedings on Information Theory **38**(1992), 719–746.
- [16] R. DeVore, L.S. Johnson, C. Pan, R. Sharpley, *Optimal entropy encoders for mining multiply resolved data*, in: Data Mining II (N. Ebecken and C.A. Brebbia, Eds.), WIT Press, Boston, 2000, 73–82.

- [17] R. DeVore, G. Petrova and P. Wojtaszczyk, *Instance-optimality in probability with an l_1 -minimization decoder*, Appl. Comput. Harmon. Anal. **27**(2009), 275–288.
- [18] D. Donoho and Y. Tsaig, *Compressed Sensing*, IEEE Trans. Information Theory **52**(2006), 1289–1306.
- [19] M.F. Duarte, M.A. Davenport, D. Takhar, J.N. Laska, T. Sun, K.F. Kelly, R.G. Baraniuk, *Single Pixel Imaging via Compressive Sampling*, IEEE Signal Processing Magazine, **25**(2008), 83–91.
- [20] A. Garnaev and E.D. Gluskin, *The widths of a Euclidean ball* (Russian), Dokl. Akad. Nauk SSSR (5) **277**(1984), 1048–1052.
- [21] A.C. Gilbert, S. Mutukrishnan, M.J. Strauss, *Approximation of Functions over Redundant Dictionaries Using Coherence*, Proc. 14th Annu. ACM-SIAM Symp. Discrete Algorithms, Baltimore, MD (2003), 243–252.
- [22] M. Haider, S. Uhlemann, E. Schwan, H. Rose, B. Kabius, K. Urban, *Electron microscopy image enhanced*, Nature (6678) **392**(1998), 768–769.
- [23] E. M. James and N. D. Browning, *Practical aspects of atomic resolution imaging and analysis in STEM*, Ultramicroscopy (1-4)**78** (1999), 125–139.
- [24] B. Kashin, *The Widths of Certain Finite Dimensional Sets and Classes of Smooth Functions*, Izvestia **41**(1977), 334–351.
- [25] E.J. Kirkland, *Advanced Computing in Electron Microscopy*, Springer, NY (2010).
- [26] Y. Mao, B.P. Fahimian, S.J. Osher, J. Miao, *Development and Optimization of Regularized Tomographic Reconstruction Algorithms Utilizing Equally-Sloped Tomography*, IEEE Transactions on Image Processing, **19** (No 5)(2010), 1259–1268.
- [27] D. Needell and J.A. Tropp, *CoSaMP: Iterative signal recovery from incomplete and inaccurate samples*, Applied and Computational Harmonic Analysis **26**(No 3)(2008), 301–321.
- [28] Y. Nesterov, *A Method for Unconstrained Convex Minimization Problem with the Rate of Convergence $O(1/k^2)$* , Doklady AN USSR, **269** (1983).
- [29] Y. Nesterov, *Smooth Minimization of Non-Smooth Functions*, Mathematical Programming, **103** (2005), 127–152.
- [30] S. Osher, M. Burger, D. Goldfarb, J. Xu, W. Yin, *An Iterative Regularization Method for Total Variation Based Image Restoration*, Multi-scale Model. Simul., **4** (No. 2) (2005), 460–489.
- [31] S. Osher, Y. Mao, B. Dong, W. Yin, *Fast Linearized Bregman Iteration for Compressive Sensing and Sparse Denoising*, Commun. Math. Sci., **8** (2010), 93–111.

- [32] H. Sawada, Y. Tanishiro, N. Ohashi, T. Tomita, F. Hosokawa, T. Kaneyama, Y. Kondo, K. Takayanagi, *STEM imaging of 47-pm-separated atomic columns by a spherical aberration-corrected electron microscope with a 300-kV cold field emission gun*, Journal of Electron Microscopy (6) **58**(2009), 357–361.
- [33] J. Shapiro, *Embedded Image Coding Using Zero-Trees of Wavelet Coefficients*, IEEE Trans. Signal Process. **41**(1993), 3445–3462.
- [34] V.N. Temlyakov, *Greedy Approximation*, Acta Num. **10**(2008), 235–409.
- [35] J. Tropp, *Greedy is Good: Algorithmic Results for Sparse Approximation*, IEEE Trans. Inform. Theory **10**(2004), 2231–2242.
- [36] M. Weyland, P.A. Midgley, J. M. Thomas, *Electron Tomography of Nanoparticle Catalysts on Porous Supports: A New Technique Based on Rutherford Scattering*, J. Phys. Chem. B. **105**(2001), 7882–7886.
- [37] W. Yin, S. Osher, J. Darbon, and D. Goldfarb, *Bregman Iterative Algorithms for Compressed Sensing and Related Problems*, CAAM Technical Report, **TR07-13** (2007).
- [38] M. Zibulevsky and M. Elad, *L1-L2 Optimization in Signal and Image Processing: Iterative shrinkage and beyond*, IEEE Signal Processing Magazine, **10**, (2010), 76–88.

Peter Binev, Department of Mathematics, University of South Carolina, Columbia, SC 29208, USA, binev@math.sc.edu

Wolfgang Dahmen, Institut für Geometrie und Praktische Mathematik, RWTH Aachen, Templergraben 55, D-52056 Aachen, Germany, dahmen@igpm.rwth-aachen.de

Ronald DeVore, Department of Mathematics, Texas A&M University, College Station, TX 77840, USA, devore@math.tamu.edu

Philipp Lamby, Department of Mathematics, University of South Carolina, Columbia, SC 29208, USA, lamby@math.sc.edu

Daniel Savu, Department of Mathematics, University of South Carolina, Columbia, SC 29208, USA, savu@math.sc.edu

Robert C. Sharpley, Department of Mathematics, University of South Carolina, Columbia, SC 29208, USA, sharpley@math.sc.edu