

A Singular Value Thresholding Algorithm for Matrix Completion

Jian-Feng Cai[†] Emmanuel J. Candès[#] Zuowei Shen[§]

[†] Temasek Laboratories, National University of Singapore, Singapore 117543

[#] Applied and Computational Mathematics, Caltech, Pasadena, CA 91125

[§] Department of Mathematics, National University of Singapore, Singapore 117543

September 2008

Abstract

This paper introduces a novel algorithm to approximate the matrix with minimum nuclear norm among all matrices obeying a set of convex constraints. This problem may be understood as the convex relaxation of a rank minimization problem, and arises in many important applications as in the task of recovering a large matrix from a small subset of its entries (the famous Netflix problem). Off-the-shelf algorithms such as interior point methods are not directly amenable to large problems of this kind with over a million unknown entries.

This paper develops a simple first-order and easy-to-implement algorithm that is extremely efficient at addressing problems in which the optimal solution has low rank. The algorithm is iterative and produces a sequence of matrices $\{\mathbf{X}^k, \mathbf{Y}^k\}$ and at each step, mainly performs a soft-thresholding operation on the singular values of the matrix $\mathbf{Y}^k$. There are two remarkable features making this attractive for low-rank matrix completion problems. The first is that the soft-thresholding operation is applied to a sparse matrix; the second is that the rank of the iterates $\{\mathbf{X}^k\}$ is empirically nondecreasing. Both these facts allow the algorithm to make use of very minimal storage space and keep the computational cost of each iteration low. On the theoretical side, we provide a convergence analysis showing that the sequence of iterates converges. On the practical side, we provide numerical examples in which $1,000 \times 1,000$ matrices are recovered in less than a minute on a modest desktop computer. We also demonstrate that our approach is amenable to very large scale problems by recovering matrices of rank about 10 with nearly a billion unknowns from just about 0.4% of their sampled entries. Our methods are connected with the recent literature on linearized Bregman iterations for ℓ_1 minimization, and we develop a framework in which one can understand these algorithms in terms of well-known Lagrange multiplier algorithms.

Keywords. Nuclear norm minimization, matrix completion, singular value thresholding, Lagrange dual function, Uzawa's algorithm.

1 Introduction

1.1 Motivation

There is a rapidly growing interest in the recovery of an unknown low-rank or approximately low-rank matrix from very limited information. This problem occurs in many areas of engineering and

applied science such as machine learning [1, 3, 4], control [42] and computer vision, see [48]. As a motivating example, consider the problem of recovering a data matrix from a sampling of its entries. This routinely comes up whenever one collects partially filled out surveys, and one would like to infer the many missing entries. In the area of recommender systems, users submit ratings on a subset of entries in a database, and the vendor provides recommendations based on the user's preferences. Because users only rate a few items, one would like to infer their preference for unrated items; this is the famous Netflix problem [2]. Recovering a rectangular matrix from a sampling of its entries is known as the *matrix completion* problem. The issue is of course that this problem is extraordinarily ill posed since with fewer samples than entries, we have infinitely many completions. Therefore, it is apparently impossible to identify which of these candidate solutions is indeed the "correct" one without some additional information.

In many instances, however, the matrix we wish to recover has low rank or approximately low rank. For instance, the Netflix data matrix of all user-ratings may be approximately low-rank because it is commonly believed that only a few factors contribute to anyone's taste or preference. In computer vision, inferring scene geometry and camera motion from a sequence of images is a well-studied problem known as the structure-from-motion problem. This is an ill-conditioned problem for objects may be distant with respect to their size, or especially for "missing data" which occur because of occlusion or tracking failures. However, when properly stacked and indexed, these images form a matrix which has very low rank (e.g. rank 3 under orthography) [21, 48]. Other examples of low-rank matrix fitting abound; e.g. in control (system identification), machine learning (multi-class learning) and so on. Having said this, the premise that the unknown has (approximately) low rank radically changes the problem, making the search for solutions feasible since the lowest-rank solution now tends to be the right one.

In a recent paper [13], Candès and Recht showed that matrix completion is not as ill-posed as people thought. Indeed, they proved that most low-rank matrices can be recovered *exactly* from most sets of sampled entries even though these sets have surprisingly small cardinality, and more importantly, they proved that this can be done by solving a simple *convex* optimization problem. To state their results, suppose to simplify that the unknown matrix $\mathbf{M} \in \mathbb{R}^{n \times n}$ is square, and that one has available m sampled entries $\{\mathbf{M}_{ij} : (i, j) \in \Omega\}$ where Ω is a random subset of cardinality m . Then [13] proves that most matrices $\mathbf{M}$ of rank r can be perfectly recovered by solving the optimization problem

$$\begin{aligned} & \text{minimize} && \|\mathbf{X}\|_* \\ & \text{subject to} && X_{ij} = M_{ij}, \quad (i, j) \in \Omega, \end{aligned} \tag{1.1}$$

provided that the number of samples obeys

$$m \geq Cn^{6/5}r \log n \tag{1.2}$$

for some positive numerical constant C .¹ In (1.1), the functional $\|\mathbf{X}\|_*$ is the nuclear norm of the matrix $\mathbf{M}$, which is the sum of its singular values. The optimization problem (1.1) is convex and can be recast as a semidefinite program [29, 30]. In some sense, this is the tightest convex relaxation of the NP-hard rank minimization problem

$$\begin{aligned} & \text{minimize} && \text{rank}(\mathbf{X}) \\ & \text{subject to} && X_{ij} = M_{ij}, \quad (i, j) \in \Omega, \end{aligned} \tag{1.3}$$

¹Note that an $n \times n$ matrix of rank r depends upon $r(2n - r)$ degrees of freedom.

since the nuclear ball $\{\mathbf{X} : \|\mathbf{X}\|_* \leq 1\}$ is the convex hull of the set of rank-one matrices with spectral norm bounded by one. Another interpretation of Candès and Recht’s result is that under suitable conditions, the rank minimization program (1.3) and the convex program (1.1) are *formally equivalent* in the sense that they have exactly the same unique solution.

1.2 Algorithm outline

Because minimizing the nuclear norm both provably recovers the lowest-rank matrix subject to constraints (see [45] for related results) and gives generally good empirical results in a variety of situations, it is understandably of great interest to develop numerical methods for solving (1.1). In [13], this optimization problem was solved using one of the most advanced semidefinite programming solvers, namely, SDPT3 [47]. This solver and others like SeDuMi are based on interior-point methods, and are problematic when the size of the matrix is large because they need to solve huge systems of linear equations to compute the Newton direction. In fact, SDPT3 can only handle $n \times n$ matrices with $n \leq 100$. Presumably, one could resort to iterative solvers such as the method of conjugate gradients to solve for the Newton step but this is problematic as well since it is well known that the condition number of the Newton system increases rapidly as one gets closer to the solution. In addition, none of these general purpose solvers use the fact that the solution may have low rank. We refer the reader to [40] for some recent progress on interior-point methods concerning some special nuclear norm-minimization problems.

This paper develops the *singular value thresholding* algorithm for approximately solving the nuclear norm minimization problem (1.1) and by extension, problems of the form

$$\begin{aligned} & \text{minimize} && \|\mathbf{X}\|_* \\ & \text{subject to} && \mathcal{A}(\mathbf{X}) = \mathbf{b}, \end{aligned} \tag{1.4}$$

where $\mathcal{A}$ is a linear operator acting on the space of $n_1 \times n_2$ matrices and $\mathbf{b} \in \mathbb{R}^m$. This algorithm is a simple first-order method, and is especially well suited for problems of very large sizes in which the solution has low rank. We sketch this algorithm in the special matrix completion setting and let $\mathcal{P}_\Omega$ be the orthogonal projector onto the span of matrices vanishing outside of Ω so that the (i, j) th component of $\mathcal{P}_\Omega(\mathbf{X})$ is equal to X_{ij} if $(i, j) \in \Omega$ and zero otherwise. Our problem may be expressed as

$$\begin{aligned} & \text{minimize} && \|\mathbf{X}\|_* \\ & \text{subject to} && \mathcal{P}_\Omega(\mathbf{X}) = \mathcal{P}_\Omega(\mathbf{M}), \end{aligned} \tag{1.5}$$

with optimization variable $\mathbf{X} \in \mathbb{R}^{n_1 \times n_2}$. Fix $\tau > 0$ and a sequence $\{\delta_k\}_{k \geq 1}$ of scalar step sizes. Then starting with $\mathbf{Y}^0 = \mathbf{0} \in \mathbb{R}^{n_1 \times n_2}$, the algorithm inductively defines

$$\begin{cases} \mathbf{X}^k = \text{shrink}(\mathbf{Y}^{k-1}, \tau), \\ \mathbf{Y}^k = \mathbf{Y}^{k-1} + \delta_k \mathcal{P}_\Omega(\mathbf{M} - \mathbf{X}^k) \end{cases} \tag{1.6}$$

until a stopping criterion is reached. In (1.6), $\text{shrink}(\mathbf{Y}, \tau)$ is a nonlinear function which applies a soft-thresholding rule at level τ to the singular values of the input matrix, see Section 2 for details. The key property here is that for large values of τ , the sequence $\{\mathbf{X}^k\}$ converges to a solution which very nearly minimizes (1.5). Hence, at each step, one only needs to compute at most one singular value decomposition and perform a few elementary matrix additions. Two important remarks are in order:

1. *Sparsity.* For each $k \geq 0$, $\mathbf{Y}^k$ vanishes outside of Ω and is, therefore, sparse, a fact which can be used to evaluate the shrink function rapidly.
2. *Low-rank property.* The matrices $\mathbf{X}^k$ turn out to have low rank, and hence the algorithm has minimum storage requirement since we only need to keep principal factors in memory.

Our numerical experiments demonstrate that the proposed algorithm can solve problems, in Matlab, involving matrices of size $30,000 \times 30,000$ having close to a billion unknowns in 17 minutes on a standard desktop computer with a 1.86 GHz CPU (dual core with Matlab's multithreading option enabled) and 3 GB of memory. As a consequence, the singular value thresholding algorithm may become a rather powerful computational tool for large scale matrix completion.

1.3 General formulation

The singular value thresholding algorithm can be adapted to deal with other types of convex constraints. For instance, it may address problems of the form

$$\begin{aligned} & \text{minimize} && \|\mathbf{X}\|_* \\ & \text{subject to} && f_i(\mathbf{X}) \leq 0, \quad i = 1, \dots, m, \end{aligned} \tag{1.7}$$

where each f_i is a Lipschitz convex function (note that one can handle linear equality constraints by considering pairs of affine functionals). In the simpler case where the f_i 's are affine functionals, the general algorithm goes through a sequence of iterations which greatly resemble (1.6). This is useful because this enables the development of numerical algorithms which are effective for recovering matrices from a small subset of sampled entries possibly contaminated with noise.

1.4 Contents and notations

The rest of the paper is organized as follows. In Section 2, we derive the singular value thresholding (SVT) algorithm for the matrix completion problem, and recasts it in terms of a well-known Lagrange multiplier algorithm. In Section 3, we extend the SVT algorithm and formulate a general iteration which is applicable to general convex constraints. In Section 4, we establish the convergence results for the iterations given in Sections 2 and 3. We demonstrate the performance and effectiveness of the algorithm through numerical examples in Section 5, and review additional implementation details. Finally, we conclude the paper with a short discussion in Section 6.

Before continuing, we provide here a brief summary of the notations used throughout the paper. Matrices are bold capital, vectors are bold lowercase and scalars or entries are not bold. For instance, $\mathbf{X}$ is a matrix and X_{ij} its (i, j) th entry. Likewise, $\mathbf{x}$ is a vector and x_i its i th component. The nuclear norm of a matrix is denoted by $\|\mathbf{X}\|_*$, the Frobenius norm by $\|\mathbf{X}\|_F$ and the spectral norm by $\|\mathbf{X}\|_2$; note that these are respectively the 1-norm, the 2-norm and the sup-norm of the vector of singular values. The adjoint of a matrix $\mathbf{X}$ is $\mathbf{X}^*$ and similarly for vectors. The notation $\text{diag}(\mathbf{x})$, where $\mathbf{x}$ is a vector, stands for the diagonal matrix with $\{x_i\}$ as diagonal elements. We denote by $\langle \mathbf{X}, \mathbf{Y} \rangle = \text{trace}(\mathbf{X}^* \mathbf{Y})$ the standard inner product between two matrices ($\|\mathbf{X}\|_F^2 = \langle \mathbf{X}, \mathbf{X} \rangle$). The Cauchy-Schwarz inequality gives $\langle \mathbf{X}, \mathbf{Y} \rangle \leq \|\mathbf{X}\|_F \|\mathbf{Y}\|_F$ and it is well known that we also have $\langle \mathbf{X}, \mathbf{Y} \rangle \leq \|\mathbf{X}\|_* \|\mathbf{Y}\|_2$ (the spectral and nuclear norms are dual from one another), see e.g. [13, 45].

2 The Singular Value Thresholding Algorithm

This section introduces the singular value thresholding algorithm and discusses some of its basic properties. We begin with the definition of a key building block, namely, the singular value thresholding operator.

2.1 The singular value shrinkage operator

Consider the singular value decomposition (SVD) of a matrix $\mathbf{X} \in \mathbb{R}^{n_1 \times n_2}$ of rank r

$$\mathbf{X} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^*, \quad \mathbf{\Sigma} = \text{diag}(\{\sigma_i\}_{1 \leq i \leq r}), \quad (2.1)$$

where $\mathbf{U}$ and $\mathbf{V}$ are respectively $n_1 \times r$ and $n_2 \times r$ matrices with orthonormal columns, and the singular values σ_i are positive (unless specified otherwise, we will always assume that the SVD of a matrix is given in the reduced form above). For each $\tau \geq 0$, we introduce the soft-thresholding operator $\mathcal{D}_\tau$ defined as follows:

$$\mathcal{D}_\tau(\mathbf{X}) := \mathbf{U}\mathcal{D}_\tau(\mathbf{\Sigma})\mathbf{V}^*, \quad \mathcal{D}_\tau(\mathbf{\Sigma}) = \text{diag}(\{\sigma_i - \tau\}_+), \quad (2.2)$$

where t_+ is the positive part of t , namely, $t_+ = \max(0, t)$. In words, this operator simply applies a soft-thresholding rule to the singular values of $\mathbf{X}$, effectively shrinking these towards zero. This is the reason why we will also refer to this transformation as the *singular value shrinkage* operator. Even though the SVD may not be unique, it is easy to see that the singular value shrinkage operator is well defined and we do not elaborate further on this issue. In some sense, this shrinkage operator is a straightforward extension of the soft-thresholding rule for scalars and vectors. In particular, note that if many of the singular values of $\mathbf{X}$ are below the threshold τ , the rank of $\mathcal{D}_\tau(\mathbf{X})$ may be considerably lower than that of $\mathbf{X}$, just like the soft-thresholding rule applied to vectors leads to sparser outputs whenever some entries of the input are below threshold.

The singular value thresholding operator is the proximity operator associated with the nuclear norm. Details about the proximity operator can be found in e.g. [35].

Theorem 2.1 *For each $\tau \geq 0$ and $\mathbf{Y} \in \mathbb{R}^{n_1 \times n_2}$, the singular value shrinkage operator (2.2) obeys*

$$\mathcal{D}_\tau(\mathbf{Y}) = \arg \min_{\mathbf{X}} \left\{ \frac{1}{2} \|\mathbf{X} - \mathbf{Y}\|_F^2 + \tau \|\mathbf{X}\|_* \right\}. \quad (2.3)$$

Proof. Since the function $h_0(\mathbf{X}) := \tau \|\mathbf{X}\|_* + \frac{1}{2} \|\mathbf{X} - \mathbf{Y}\|_F^2$ is strictly convex, it is easy to see that there exists a unique minimizer, and we thus need to prove that it is equal to $\mathcal{D}_\tau(\mathbf{Y})$. To do this, recall the definition of a subgradient of a convex function $f : \mathbb{R}^{n_1 \times n_2} \rightarrow \mathbb{R}$. We say that $\mathbf{Z}$ is a subgradient of f at $\mathbf{X}_0$, denoted $\mathbf{Z} \in \partial f(\mathbf{X}_0)$, if

$$f(\mathbf{X}) \geq f(\mathbf{X}_0) + \langle \mathbf{Z}, \mathbf{X} - \mathbf{X}_0 \rangle \quad (2.4)$$

for all $\mathbf{X}$. Now $\hat{\mathbf{X}}$ minimizes h_0 if and only if $\mathbf{0}$ is a subgradient of the functional h_0 at the point $\hat{\mathbf{X}}$, i.e.

$$\mathbf{0} \in \hat{\mathbf{X}} - \mathbf{Y} + \tau \partial \|\hat{\mathbf{X}}\|_*, \quad (2.5)$$

where $\partial \|\hat{\mathbf{X}}\|_*$ is the set of subgradients of the nuclear norm. Let $\mathbf{X} \in \mathbb{R}^{n_1 \times n_2}$ be an arbitrary matrix and $\mathbf{U}\mathbf{\Sigma}\mathbf{V}^*$ be its SVD. It is known [13, 37, 49] that

$$\partial \|\mathbf{X}\|_* = \{ \mathbf{U}\mathbf{V}^* + \mathbf{W} : \mathbf{W} \in \mathbb{R}^{n_1 \times n_2}, \mathbf{U}^*\mathbf{W} = \mathbf{0}, \mathbf{W}\mathbf{V} = \mathbf{0}, \|\mathbf{W}\|_2 \leq 1 \}. \quad (2.6)$$

Set $\hat{\mathbf{X}} := \mathcal{D}_\tau(\mathbf{Y})$ for short. In order to show that $\hat{\mathbf{X}}$ obeys (2.5), decompose the SVD of $\mathbf{Y}$ as

$$\mathbf{Y} = \mathbf{U}_0 \mathbf{\Sigma}_0 \mathbf{V}_0^* + \mathbf{U}_1 \mathbf{\Sigma}_1 \mathbf{V}_1^*,$$

where $\mathbf{U}_0, \mathbf{V}_0$ (resp. $\mathbf{U}_1, \mathbf{V}_1$) are the singular vectors associated with singular values greater than τ (resp. smaller than or equal to τ). With these notations, we have

$$\hat{\mathbf{X}} = \mathbf{U}_0 (\mathbf{\Sigma}_0 - \tau \mathbf{I}) \mathbf{V}_0^*$$

and, therefore,

$$\mathbf{Y} - \hat{\mathbf{X}} = \tau (\mathbf{U}_0 \mathbf{V}_0^* + \mathbf{W}), \quad \mathbf{W} = \tau^{-1} \mathbf{U}_1 \mathbf{\Sigma}_1 \mathbf{V}_1^*.$$

By definition, $\mathbf{U}_0^* \mathbf{W} = 0$, $\mathbf{W} \mathbf{V}_0 = 0$ and since the diagonal elements of $\mathbf{\Sigma}_1$ have magnitudes bounded by τ , we also have $\|\mathbf{W}\|_2 \leq 1$. Hence $\mathbf{Y} - \hat{\mathbf{X}} \in \tau \partial \|\hat{\mathbf{X}}\|_*$, which concludes the proof. $\square$

2.2 Shrinkage iterations

We are now in the position to introduce the singular value thresholding algorithm. Fix $\tau > 0$ and a sequence $\{\delta_k\}$ of positive step sizes. Starting with $\mathbf{Y}_0$, inductively define for $k = 1, 2, \dots$,

$$\begin{cases} \mathbf{X}^k = \mathcal{D}_\tau(\mathbf{Y}^{k-1}), \\ \mathbf{Y}^k = \mathbf{Y}^{k-1} + \delta_k \mathcal{P}_\Omega(\mathbf{M} - \mathbf{X}^k) \end{cases} \quad (2.7)$$

until a stopping criterion is reached (we postpone the discussion this stopping criterion and of the choice of step sizes). This shrinkage iteration is very simple to implement. At each step, we only need to compute an SVD and perform elementary matrix operations. With the help of a standard numerical linear algebra package, the whole algorithm can be coded in just a few lines.

Before addressing further computational issues, we would like to make explicit the relationship between this iteration and the original problem (1.1). In Section 4, we will show that the sequence $\{\mathbf{X}^k\}$ converges to the unique solution of an optimization problem closely related to (1.1), namely,

$$\begin{aligned} & \text{minimize} && \tau \|\mathbf{X}\|_* + \frac{1}{2} \|\mathbf{X}\|_F^2 \\ & \text{subject to} && \mathcal{P}_\Omega(\mathbf{X}) = \mathcal{P}_\Omega(\mathbf{M}). \end{aligned} \quad (2.8)$$

Furthermore, it is intuitive that the solution to this modified problem converges to that of (1.5) as $\tau \rightarrow \infty$ as shown in Section 3. Thus by selecting a large value of the parameter τ , the sequence of iterates converges to a matrix which nearly minimizes (1.1).

As mentioned earlier, there are two crucial properties which make this algorithm ideally suited for matrix completion.

- *Low-rank property.* A remarkable empirical fact is that the matrices in the sequence $\{\mathbf{X}^k\}$ have low rank (provided, of course, that the solution to (2.8) has low rank). We use the word “empirical” because all of our numerical experiments have produced low-rank sequences but we cannot rigorously prove that this is true in general. The reason for this phenomenon is, however, simple: because we are interested in large values of τ (as to better approximate the solution to (1.1)), the thresholding step happens to ‘kill’ most of the small singular values and produces a low-rank output. In fact, our numerical results show that the rank of $\mathbf{X}^k$ is nondecreasing with k , and the maximum rank is reached in the last steps of the algorithm, see Section 5.

Thus, when the rank of the solution is substantially smaller than either dimension of the matrix, the storage requirement is low since we could store each $\mathbf{X}^k$ in its SVD form (note that we only need to keep the current iterate and may discard earlier values).

- *Sparsity.* Another important property of the SVT algorithm is that the iteration matrix $\mathbf{Y}^k$ is sparse. Since $\mathbf{Y}^0 = \mathbf{0}$, we have by induction that $\mathbf{Y}^k$ vanishes outside of Ω . The fewer entries available, the sparser $\mathbf{Y}^k$. Because the sparsity pattern Ω is fixed throughout, one can then apply sparse matrix techniques to save storage. Also, if $|\Omega| = m$, the computational cost of updating $\mathbf{Y}^k$ is of order m . Moreover, we can call subroutines supporting sparse matrix computations, which can further reduce computational costs.

One such subroutine is the SVD. However, note that we do not need to compute the entire SVD of $\mathbf{Y}^k$ to apply the singular value thresholding operator. Only the part corresponding to singular values greater than τ is needed. Hence, a good strategy is to apply the iterative Lanczos algorithm to compute the first few singular values and singular vectors. Because $\mathbf{Y}^k$ is sparse, $\mathbf{Y}^k$ can be applied to arbitrary vectors rapidly, and this procedure offers a considerable speedup over naive methods.

2.3 Relation with other works

Our algorithm is inspired by recent work in the area of ℓ_1 minimization, and especially by the work on linearized Bregman iterations for compressed sensing, see [9–11, 23, 44, 51] for linearized Bregman iterations and [14–17, 26] for some information about the field of compressed sensing. In this line of work, linearized Bregman iterations are used to find the solution to an underdetermined system of linear equations with minimum ℓ_1 norm. In fact, Theorem 2.1 asserts that the singular value thresholding algorithm can be formulated as a linearized Bregman iteration. Bregman iterations were first introduced in [43] as a convenient tool for solving computational problems in the imaging sciences, and a later paper [51] showed that they were useful for solving ℓ_1 -norm minimization problems in the area of compressed sensing. Linearized Bregman iterations were proposed in [23] to improve performance of plain Bregman iterations, see also [51]. Additional details together with a technique for improving the speed of convergence called *kicking* are described in [44]. On the practical side, the paper [11] applied Bregman iterations to solve a deblurring problem while on the theoretical side, the references [9, 10] gave a rigorous analysis of the convergence of such iterations. New developments keep on coming out at a rapid pace and recently, [32] introduced a new iteration, the *split Bregman iteration*, to extend Bregman-type iterations (such as linearized Bregman iterations) to problems involving the minimization of ℓ_1 -like functionals such as total-variation norms, Besov norms, and so forth.

When applied to ℓ_1 -minimization problems, linearized Bregman iterations are sequences of soft-thresholding rules operating on vectors. Iterative soft-thresholding algorithms in connection with ℓ_1 or total-variation minimization have quite a bit of history in signal and image processing and we would like to mention the works [12, 39] for total-variation minimization, [24, 25, 31] for ℓ_1 minimization, and [5, 7, 8, 19, 20, 27, 28, 46] for some recent applications in the area of image inpainting and image restoration. Just as iterative soft-thresholding methods are designed to find sparse solutions, our iterative singular value thresholding scheme is designed to find a sparse vector of singular values. In classical problems arising in the areas of compressed sensing, and signal or image processing, the sparsity is expressed in a known transformed domain and soft-thresholding is applied to transformed coefficients. In contrast, the shrinkage operator $\mathcal{D}_\tau$ is adaptive. The SVT not

only discovers a sparse singular vector but also the bases in which we have a sparse representation. In this sense, the SVT algorithm is an extension of earlier iterative soft-thresholding schemes.

Finally, we would like to contrast the SVT iteration (2.7) with the popular iterative soft-thresholding algorithm used in many papers in imaging processing and perhaps best known under the name of Proximal Forward-Backward Splitting method (PFBS), see [8,22,24,31,33] for example. The constrained minimization problem (1.5) may be relaxed into

$$\text{minimize} \quad \lambda \|\mathbf{X}\|_* + \frac{1}{2} \|\mathcal{P}_\Omega(\mathbf{X}) - \mathcal{P}_\Omega(\mathbf{M})\|_F^2 \quad (2.9)$$

for some $\lambda > 0$. Theorem 2.1 asserts that $\mathcal{D}_\lambda$ is the proximity operator of $\lambda \|\mathbf{X}\|_*$ and Proposition 3.1(iii) in [22] gives that the solution to this unconstrained problem is characterized by the fixed point equation $\mathbf{X} = \mathcal{D}_{\lambda\delta}(\mathbf{X} + \delta \mathcal{P}_\Omega(\mathbf{M} - \mathbf{X}))$ for each $\delta > 0$. One can then apply a simplified version of the PFBS method (see (3.6) in [22]) to obtain iterations of the form

$$\mathbf{X}^k = \mathcal{D}_{\lambda\delta_{k-1}}(\mathbf{X}^{k-1} + \delta_{k-1} \mathcal{P}_\Omega(\mathbf{M} - \mathbf{X}^{k-1})).$$

Introducing an intermediate matrix $\mathbf{Y}^k$, this algorithm may be expressed as

$$\begin{cases} \mathbf{X}^k = \mathcal{D}_{\lambda\delta_{k-1}}(\mathbf{Y}^{k-1}), \\ \mathbf{Y}^k = \mathbf{X}^k + \delta_k \mathcal{P}_\Omega(\mathbf{M} - \mathbf{X}^k). \end{cases} \quad (2.10)$$

The difference with (2.7) may seem subtle at first—replacing $\mathbf{X}^k$ in (2.10) with $\mathbf{Y}^{k-1}$ and setting $\delta_k = \delta$ gives (2.7) with $\tau = \lambda\delta$ —but has enormous consequences as this gives entirely different algorithms. First, they have different limits: while (2.7) converges to the solution of the constrained minimization (2.8), (2.10) converges to the solution of (2.9) provided that the sequence of step sizes is appropriately selected. Second, selecting a large λ (or a large value of $\tau = \lambda\delta$) in (2.10) gives a low-rank sequence of iterates and a limit with small nuclear norm. The limit, however, does not fit the data and this is why one has to choose a small or moderate value of λ (or of $\tau = \lambda\delta$). However, when λ is not sufficiently large, the $\mathbf{X}^k$ may not have low rank even though the solution has low rank (and one may need to compute many singular vectors), and $\mathbf{Y}^k$ is not sufficiently sparse to make the algorithm computationally attractive. Moreover, the limit does not necessary have a small nuclear norm. These are reasons why (2.10) is not suitable for matrix completion.

2.4 Interpretation as a Lagrange multiplier method

In this section, we recast the SVT algorithm as a type of Lagrange multiplier algorithm known as Uzawa’s algorithm. An important consequence is that this will allow us to extend the SVT algorithm to other problems involving the minimization of the nuclear norm under convex constraints, see Section 3. Further, another contribution of this paper is that this framework actually recasts linear Bregman iterations as a very special form of Uzawa’s algorithm, hence providing fresh and clear insights about these iterations.

In what follows, we set $f_\tau(\mathbf{X}) = \tau \|\mathbf{X}\|_* + \frac{1}{2} \|\mathbf{X}\|_F^2$ for some fixed $\tau > 0$ and recall that we wish to solve (2.8)

$$\begin{aligned} &\text{minimize} && f_\tau(\mathbf{X}) \\ &\text{subject to} && \mathcal{P}_\Omega(\mathbf{X}) = \mathcal{P}_\Omega(\mathbf{M}). \end{aligned}$$

The Lagrangian for this problem is given by

$$\mathcal{L}(\mathbf{X}, \mathbf{Y}) = f_\tau(\mathbf{X}) + \langle \mathbf{Y}, \mathcal{P}_\Omega(\mathbf{M} - \mathbf{X}) \rangle,$$

where $\mathbf{Y} \in \mathbb{R}^{n_1 \times n_2}$. Strong duality holds and $\mathbf{X}^*$ and $\mathbf{Y}^*$ are primal-dual optimal if $(\mathbf{X}^*, \mathbf{Y}^*)$ is a saddlepoint of the Lagrangian $\mathcal{L}(\mathbf{X}, \mathbf{Y})$, i.e. a pair obeying

$$\sup_{\mathbf{Y}} \inf_{\mathbf{X}} \mathcal{L}(\mathbf{X}, \mathbf{Y}) = \mathcal{L}(\mathbf{X}^*, \mathbf{Y}^*) = \inf_{\mathbf{X}} \sup_{\mathbf{Y}} \mathcal{L}(\mathbf{X}, \mathbf{Y}). \quad (2.11)$$

(The function $g_0(\mathbf{Y}) = \inf_{\mathbf{X}} \mathcal{L}(\mathbf{X}, \mathbf{Y})$ is called the dual function.) Uzawa's algorithm approaches the problem of finding a saddlepoint with an iterative procedure. From $\mathbf{Y}_0 = \mathbf{0}$, say, inductively define

$$\begin{cases} \mathcal{L}(\mathbf{X}^k, \mathbf{Y}^{k-1}) = \min_{\mathbf{X}} \mathcal{L}(\mathbf{X}, \mathbf{Y}^{k-1}) \\ \mathbf{Y}^k = \mathbf{Y}^{k-1} + \delta_k \mathcal{P}_{\Omega}(\mathbf{M} - \mathbf{X}^k), \end{cases} \quad (2.12)$$

where $\{\delta_k\}_{k \geq 1}$ is a sequence of positive step sizes. Uzawa's algorithm is, in fact, a subgradient method applied to the dual problem, where each step moves the current iterate in the direction of the gradient or of a subgradient. Indeed, observe that

$$\partial_{\mathbf{Y}} g_0(\mathbf{Y}) = \partial_{\mathbf{Y}} \mathcal{L}(\tilde{\mathbf{X}}, \mathbf{Y}) = \mathcal{P}_{\Omega}(\mathbf{M} - \tilde{\mathbf{X}}), \quad (2.13)$$

where $\tilde{\mathbf{X}}$ is the minimizer of the Lagrangian for that value of $\mathbf{Y}$ so that a gradient descent update for $\mathbf{Y}$ is of the form

$$\mathbf{Y}^k = \mathbf{Y}^{k-1} + \delta_k \partial_{\mathbf{Y}} g_0(\mathbf{Y}^{k-1}) = \mathbf{Y}^{k-1} + \delta_k \mathcal{P}_{\Omega}(\mathbf{M} - \mathbf{X}^k).$$

It remains to compute the minimizer of the Lagrangian (2.12), and note that

$$\arg \min f_{\tau}(\mathbf{X}) + \langle \mathbf{Y}, \mathcal{P}_{\Omega}(\mathbf{M} - \mathbf{X}) \rangle = \arg \min \tau \|\mathbf{X}\|_* + \frac{1}{2} \|\mathbf{X} - \mathcal{P}_{\Omega} \mathbf{Y}\|_F^2. \quad (2.14)$$

However, we know that the minimizer is given by $\mathcal{D}_{\tau}(\mathcal{P}_{\Omega}(\mathbf{Y}))$ and since $\mathbf{Y}^k = \mathcal{P}_{\Omega}(\mathbf{Y}^k)$ for all $k \geq 0$, Uzawa's algorithm takes the form

$$\begin{cases} \mathbf{X}^k = \mathcal{D}_{\tau}(\mathbf{Y}^{k-1}) \\ \mathbf{Y}^k = \mathbf{Y}^{k-1} + \delta_k \mathcal{P}_{\Omega}(\mathbf{M} - \mathbf{X}^k), \end{cases}$$

which is exactly the update (2.7). This point of view brings to bear many different mathematical tools for proving the convergence of the singular value thresholding iterations. For an early use of Uzawa's algorithm minimizing an ℓ_1 -like functional, the total-variation norm, under linear inequality constraints, see [12].

3 General Formulation

This section presents a general formulation of the SVT algorithm for approximately minimizing the nuclear norm of a matrix under convex constraints.

3.1 Linear equality constraints

Set the objective functional $f_{\tau}(\mathbf{X}) = \tau \|\mathbf{X}\|_* + \frac{1}{2} \|\mathbf{X}\|_F^2$ for some fixed $\tau > 0$, and consider the following optimization problem:

$$\begin{aligned} & \text{minimize} && f_{\tau}(\mathbf{X}) \\ & \text{subject to} && \mathcal{A}(\mathbf{X}) = \mathbf{b}, \end{aligned} \quad (3.1)$$

where $\mathcal{A}$ is a linear transformation mapping $n_1 \times n_2$ matrices into $\mathbb{R}^m$ ($\mathcal{A}^*$ is the adjoint of $\mathcal{A}$). This more general formulation is considered in [13] and [45] as an extension of the matrix completion problem. Then the Lagrangian for this problem is of the form

$$\mathcal{L}(\mathbf{X}, \mathbf{y}) = f_\tau(\mathbf{X}) + \langle \mathbf{y}, \mathbf{b} - \mathcal{A}(\mathbf{X}) \rangle, \quad (3.2)$$

where $\mathbf{X} \in \mathbb{R}^{n_1 \times n_2}$ and $\mathbf{y} \in \mathbb{R}^m$, and starting with $\mathbf{y}^0 = \mathbf{0}$, Uzawa's iteration is given by

$$\begin{cases} \mathbf{X}^k = \mathcal{D}_\tau(\mathcal{A}^*(\mathbf{y}^{k-1})), \\ \mathbf{y}^k = \mathbf{y}^{k-1} + \delta_k(\mathbf{b} - \mathcal{A}(\mathbf{X}^k)). \end{cases} \quad (3.3)$$

The iteration (3.3) is of course the same as (2.7) in the case where $\mathcal{A}$ is a sampling operator extracting m entries with indices in Ω out of an $n_1 \times n_2$ matrix. To verify this claim, observe that in this situation, $\mathcal{A}^*\mathcal{A} = \mathcal{P}_\Omega$, and let $\mathbf{M}$ be any matrix obeying $\mathcal{A}(\mathbf{M}) = \mathbf{b}$. Then defining $\mathbf{Y}^k = \mathcal{A}^*(\mathbf{y}^k)$ and substituting this expression in (3.3) gives (2.7).

3.2 General convex constraints

One can also adapt the algorithm to handle general convex constraints. Suppose we wish to minimize $f_\tau(\mathbf{X})$ defined as before over a convex set $\mathbf{X} \in \mathcal{C}$. To simplify, we will assume that this convex set is given by

$$\mathcal{C} = \{\mathbf{X} : f_i(\mathbf{X}) \leq 0, \forall i = 1, \dots, m\},$$

where the f_i 's are convex functionals (note that one can handle linear equality constraints by considering pairs of affine functionals). The problem of interest is then of the form

$$\begin{aligned} & \text{minimize} && f_\tau(\mathbf{X}) \\ & \text{subject to} && f_i(\mathbf{X}) \leq 0, \quad i = 1, \dots, m. \end{aligned} \quad (3.4)$$

Just as before, it is intuitive that as $\tau \rightarrow \infty$, the solution to this problem converges to a minimizer of the nuclear norm under the same constraints (1.7) as shown in Theorem 3.1 at the end of this section.

Put $\mathcal{F}(\mathbf{X}) := (f_1(\mathbf{X}), \dots, f_m(\mathbf{X}))$ for short. Then the Lagrangian for (3.4) is equal to

$$\mathcal{L}(\mathbf{X}, \mathbf{y}) = f_\tau(\mathbf{X}) + \langle \mathbf{y}, \mathcal{F}(\mathbf{X}) \rangle,$$

where $\mathbf{X} \in \mathbb{R}^{n_1 \times n_2}$ and $\mathbf{y} \in \mathbb{R}^m$ is now a vector with nonnegative components denoted, as usual, by $\mathbf{y} \geq \mathbf{0}$. One can apply Uzawa's method just as before with the only modification that we will use a subgradient method with projection to maximize the dual function since we need to make sure that the successive updates $\mathbf{y}^k$ belong to the nonnegative orthant. This gives

$$\begin{cases} \mathbf{X}^k = \arg \min \{f_\tau(\mathbf{X}) + \langle \mathbf{y}^{k-1}, \mathcal{F}(\mathbf{X}) \rangle\}, \\ \mathbf{y}^k = [\mathbf{y}^{k-1} + \delta_k \mathcal{F}(\mathbf{X}^k)]_+. \end{cases} \quad (3.5)$$

Above, $\mathbf{x}_+$ is of course the vector with entries equal to $\max(x_i, 0)$. When $\mathcal{F}$ is an affine mapping of the form $\mathbf{b} - \mathcal{A}(\mathbf{X})$ so that one solves

$$\begin{aligned} & \text{minimize} && f_\tau(\mathbf{X}) \\ & \text{subject to} && \mathcal{A}(\mathbf{X}) \geq \mathbf{b}, \end{aligned}$$

this simplifies to

$$\begin{cases} \mathbf{X}^k = \mathcal{D}_\tau(\mathcal{A}^*(\mathbf{y}^{k-1})), \\ \mathbf{y}^k = [\mathbf{y}^{k-1} + \delta_k(\mathbf{b} - \mathcal{A}(\mathbf{X}^k))]_+, \end{cases} \quad (3.6)$$

and thus the extension to linear inequality constraints is straightforward.

3.3 Example

An interesting example concerns the extension of the Dantzig selector [18] to matrix problems. Suppose we have available linear measurements about a matrix $\mathbf{M}$ of interest

$$\mathbf{b} = \mathcal{A}(\mathbf{M}) + \mathbf{z}, \quad (3.7)$$

where $\mathbf{z} \in \mathbb{R}^m$ is a noise vector. Then under these circumstances, one might want to find the matrix which minimizes the nuclear norm among all matrices which are consistent with the data $\mathbf{b}$. Inspired by the work on the Dantzig selector which was originally developed for estimating sparse parameter vectors from noisy data, one could approach this problem by solving

$$\begin{aligned} & \text{minimize} && \|\mathbf{X}\|_* \\ & \text{subject to} && |\text{vec}(\mathcal{A}^*(\mathbf{r}))| \leq \text{vec}(\mathbf{E}), \quad \mathbf{r} := \mathbf{b} - \mathcal{A}(\mathbf{X}), \end{aligned} \quad (3.8)$$

where $\mathbf{E}$ is an array of tolerances, which is adjusted to fit the noise statistics [18]. Above, $\text{vec}(\mathbf{A}) \leq \text{vec}(\mathbf{B})$, for any two matrices $\mathbf{A}$ and $\mathbf{B}$, means componentwise inequalities; that is, $A_{ij} \leq B_{ij}$ for all indices i, j . We use this notation as not to confuse the reader with the positive semidefinite ordering. In the case of the matrix completion problem where $\mathcal{A}$ extracts sampled entries indexed by Ω , one can always see the data vector as the sampled entries of some matrix $\mathbf{B}$ obeying $\mathcal{P}_\Omega(\mathbf{B}) = \mathcal{A}^*(\mathbf{b})$; the constraint is then natural for it may be expressed as

$$|B_{ij} - X_{ij}| \leq E_{ij}, \quad (i, j) \in \Omega,$$

If $\mathbf{z}$ is white noise with standard deviation σ , one may want to use a multiple of σ for E_{ij} . In words, we are looking for a matrix with minimum nuclear norm under the constraint that all of its sampled entries do not deviate too much from what has been observed.

Let $\mathbf{Y}_+ \in \mathbb{R}^{n_1 \times n_2}$ (resp. $\mathbf{Y}_- \in \mathbb{R}^{n_1 \times n_2}$) be the Lagrange multiplier associated with the componentwise linear inequality constraints $\text{vec}(\mathcal{A}^*(\mathbf{r})) \leq \text{vec}(\mathbf{E})$ (resp. $-\text{vec}(\mathcal{A}^*(\mathbf{r})) \leq \text{vec}(\mathbf{E})$). Then starting with $\mathbf{Y}_\pm^0 = \mathbf{0}$, the SVT iteration for this problem is of the form

$$\begin{cases} \mathbf{X}^k = \mathcal{D}_\tau(\mathcal{A}^*(\mathbf{Y}_+^{k-1} - \mathbf{Y}_-^{k-1})), \\ \mathbf{Y}_\pm^k = [\mathbf{Y}_\pm^{k-1} + \delta_k(\pm \mathcal{A}^*(\mathbf{r}^k) - \mathbf{E})]_+, \quad \mathbf{r}^k = \mathbf{b} - \mathcal{A}(\mathbf{X}^k), \end{cases} \quad (3.9)$$

where again $[\cdot]_+$ is applied componentwise.

We conclude by noting that in the matrix completion problem where $\mathcal{A}^*\mathcal{A} = \mathcal{P}_\Omega$ and one observes $\mathcal{P}_\Omega(\mathbf{B})$, one can check that this iteration simplifies to

$$\begin{cases} \mathbf{X}^k = \mathcal{D}_\tau(\mathbf{Y}_+^{k-1} - \mathbf{Y}_-^{k-1}), \\ \mathbf{Y}_\pm^k = [\mathbf{Y}_\pm^{k-1} + \delta_k \mathcal{P}_\Omega(\pm(\mathbf{B} - \mathbf{X}^k) - \mathbf{E})]_+. \end{cases} \quad (3.10)$$

Again, this is easy to implement and whenever the solution has low rank, the iterates $\mathbf{X}^k$ have low rank as well.

3.4 When the proximal problem gets close

We now show that minimizing the proximal objective $f_\tau(\mathbf{X}) = \tau\|\mathbf{X}\|_* + \frac{1}{2}\|\mathbf{X}\|_F^2$ is the same as minimizing the nuclear norm in the limit of large τ 's. The theorem below is general and covers the special case of linear equality constraints as in (2.8).

Theorem 3.1 *Let $\mathbf{X}_\tau^*$ be the solution to (3.4) and $\mathbf{X}_\infty$ be the minimum Frobenius-norm solution to (1.7) defined as*

$$\mathbf{X}_\infty := \arg \min_{\mathbf{X}} \{\|\mathbf{X}\|_F^2 : \mathbf{X} \text{ is a solution of (1.7)}\}. \quad (3.11)$$

Assume that the $f_i(\mathbf{X})$'s, $1 \leq i \leq m$, are convex and lower semi-continuous. Then

$$\lim_{\tau \rightarrow \infty} \|\mathbf{X}_\tau^* - \mathbf{X}_\infty\|_F = 0. \quad (3.12)$$

Proof. It follows from the definition of $\mathbf{X}_\tau^*$ and $\mathbf{X}_\infty$ that

$$\|\mathbf{X}_\tau^*\|_* + \frac{1}{2\tau}\|\mathbf{X}_\tau^*\|_F^2 \leq \|\mathbf{X}_\infty\|_* + \frac{1}{2\tau}\|\mathbf{X}_\infty\|_F^2, \quad \text{and} \quad \|\mathbf{X}_\infty\|_* \leq \|\mathbf{X}_\tau^*\|_*. \quad (3.13)$$

Summing these two inequalities gives

$$\|\mathbf{X}_\tau^*\|_F^2 \leq \|\mathbf{X}_\infty\|_F^2, \quad (3.14)$$

which implies that $\|\mathbf{X}_\tau^*\|_F^2$ is bounded uniformly in τ . Thus, we would prove the theorem if we could establish that any convergent subsequence $\{\mathbf{X}_{\tau_k}^*\}_{k \geq 1}$ must converge to $\mathbf{X}_\infty$.

Consider an arbitrary converging subsequence $\{\mathbf{X}_{\tau_k}^*\}$ and set $\mathbf{X}_c := \lim_{k \rightarrow \infty} \mathbf{X}_{\tau_k}^*$. Since for each $1 \leq i \leq m$, $f_i(\mathbf{X}_{\tau_k}^*) \leq 0$ and f_i is lower semi-continuous, $\mathbf{X}_c$ obeys

$$f_i(\mathbf{X}_c) \leq 0, \quad i = 1, \dots, m. \quad (3.15)$$

Furthermore, since $\|\mathbf{X}_\tau^*\|_F^2$ is bounded, (3.13) yields

$$\limsup_{\tau \rightarrow \infty} \|\mathbf{X}_\tau^*\|_* \leq \|\mathbf{X}_\infty\|_*, \quad \|\mathbf{X}_\infty\|_* \leq \liminf_{\tau \rightarrow \infty} \|\mathbf{X}_\tau^*\|_*.$$

An immediate consequence is $\lim_{\tau \rightarrow \infty} \|\mathbf{X}_\tau^*\|_* = \|\mathbf{X}_\infty\|_*$ and, therefore, $\|\mathbf{X}_c\|_* = \|\mathbf{X}_\infty\|_*$. This shows that $\mathbf{X}_c$ is a solution to (1.1). Now it follows from the definition of $\mathbf{X}_\infty$ that $\|\mathbf{X}_c\|_F \geq \|\mathbf{X}_\infty\|_F$, while we also have $\|\mathbf{X}_c\|_F \leq \|\mathbf{X}_\infty\|_F$ because of (3.14). We conclude that $\|\mathbf{X}_c\|_F = \|\mathbf{X}_\infty\|_F$ and thus $\mathbf{X}_c = \mathbf{X}_\infty$ since $\mathbf{X}_\infty$ is unique. $\square$

4 Convergence Analysis

This section establishes the convergence of the SVT iterations. We begin with the simpler proof of the convergence of (2.7) in the special case of the matrix completion problem, and then present the argument for the more general constraints (3.5). We hope that this progression will make the second and more general proof more transparent.

4.1 Convergence for matrix completion

We begin by recording a lemma which establishes the strong convexity of the objective f_τ .

Lemma 4.1 *Let $\mathbf{Z} \in \partial f_\tau(\mathbf{X})$ and $\mathbf{Z}' \in \partial f_\tau(\mathbf{X}')$. Then*

$$\langle \mathbf{Z} - \mathbf{Z}', \mathbf{X} - \mathbf{X}' \rangle \geq \|\mathbf{X} - \mathbf{X}'\|_F^2. \quad (4.1)$$

Proof. An element $\mathbf{Z}$ of $\partial f_\tau(\mathbf{X})$ is of the form $\mathbf{Z} = \tau \mathbf{Z}_0 + \mathbf{X}$, where $\mathbf{Z}_0 \in \partial \|\mathbf{X}\|_*$, and similarly for $\mathbf{Z}'$. This gives

$$\langle \mathbf{Z} - \mathbf{Z}', \mathbf{X} - \mathbf{X}' \rangle = \tau \langle \mathbf{Z}_0 - \mathbf{Z}'_0, \mathbf{X} - \mathbf{X}' \rangle + \|\mathbf{X} - \mathbf{X}'\|_F^2$$

and it thus suffices to show that the first term of the right-hand side is nonnegative. From (2.6), we have that any subgradient of the nuclear norm at $\mathbf{X}$ obeys $\|\mathbf{Z}_0\|_2 \leq 1$ and $\langle \mathbf{Z}_0, \mathbf{X} \rangle = \|\mathbf{X}\|_*$. In particular, this gives

$$|\langle \mathbf{Z}_0, \mathbf{X}' \rangle| \leq \|\mathbf{Z}_0\|_2 \|\mathbf{X}'\|_* \leq \|\mathbf{X}'\|_*, \quad |\langle \mathbf{Z}'_0, \mathbf{X} \rangle| \leq \|\mathbf{Z}'_0\|_2 \|\mathbf{X}\|_* \leq \|\mathbf{X}\|_*.$$

Whence,

$$\begin{aligned} \langle \mathbf{Z}_0 - \mathbf{Z}'_0, \mathbf{X} - \mathbf{X}' \rangle &= \langle \mathbf{Z}_0, \mathbf{X} \rangle + \langle \mathbf{Z}'_0, \mathbf{X}' \rangle - \langle \mathbf{Z}_0, \mathbf{X}' \rangle - \langle \mathbf{Z}'_0, \mathbf{X} \rangle \\ &= \|\mathbf{X}\|_* + \|\mathbf{X}'\|_* - \langle \mathbf{Z}_0, \mathbf{X}' \rangle - \langle \mathbf{Z}'_0, \mathbf{X} \rangle \geq 0, \end{aligned}$$

which proves the lemma. $\square$

This lemma is key in showing that the SVT algorithm (2.7) converges.

Theorem 4.2 *Suppose that the sequence of step sizes obeys $0 < \inf \delta_k \leq \sup \delta_k < 2$. Then the sequence $\{\mathbf{X}^k\}$ obtained via (2.7) converges to the unique solution of (2.8).*

Proof. Let $(\mathbf{X}^*, \mathbf{Y}^*)$ be primal-dual optimal for the problem (2.8). The optimality conditions give

$$\begin{aligned} \mathbf{0} &= \mathbf{Z}^k - \mathcal{P}_\Omega(\mathbf{Y}^{k-1}) \\ \mathbf{0} &= \mathbf{Z}^* - \mathcal{P}_\Omega(\mathbf{Y}^*), \end{aligned}$$

for some $\mathbf{Z}^k \in \partial f_\tau(\mathbf{X}^k)$ and some $\mathbf{Z}^* \in \partial f_\tau(\mathbf{X}^*)$. We then deduce that

$$(\mathbf{Z}^k - \mathbf{Z}^*) - \mathcal{P}_\Omega(\mathbf{Y}^{k-1} - \mathbf{Y}^*) = \mathbf{0}$$

and, therefore, it follows from Lemma 4.1 that

$$\langle \mathbf{X}^k - \mathbf{X}^*, \mathcal{P}_\Omega(\mathbf{Y}^{k-1} - \mathbf{Y}^*) \rangle = \langle \mathbf{Z}^k - \mathbf{Z}^*, \mathbf{X}^k - \mathbf{X}^* \rangle \geq \|\mathbf{X}^k - \mathbf{X}^*\|_F^2. \quad (4.2)$$

We continue and observe that because $\mathcal{P}_\Omega \mathbf{X}^* = \mathcal{P}_\Omega \mathbf{M}$,

$$\|\mathcal{P}_\Omega(\mathbf{Y}^k - \mathbf{Y}^*)\|_F = \|\mathcal{P}_\Omega(\mathbf{Y}^{k-1} - \mathbf{Y}^*) + \delta_k \mathcal{P}_\Omega(\mathbf{X}^* - \mathbf{X}^k)\|_F.$$

Therefore, setting $r_k = \|\mathcal{P}_\Omega(\mathbf{Y}^k - \mathbf{Y}^*)\|_F$,

$$\begin{aligned} r_k^2 &= r_{k-1}^2 - 2\delta_k \langle \mathcal{P}_\Omega(\mathbf{Y}^{k-1} - \mathbf{Y}^*), \mathbf{X}^k - \mathbf{X}^* \rangle + \delta_k^2 \|\mathcal{P}_\Omega(\mathbf{X}^* - \mathbf{X}^k)\|_F^2 \\ &\leq r_{k-1}^2 - 2\delta_k \|\mathbf{X}^k - \mathbf{X}^*\|_F^2 + \delta_k^2 \|\mathbf{X}^k - \mathbf{X}^*\|_F^2 \end{aligned} \quad (4.3)$$

since for any matrix $\mathbf{X}$, $\|\mathcal{P}_\Omega(\mathbf{X})\|_F \leq \|\mathbf{X}\|_F$. Under our assumptions about the size of δ_k , we have $2\delta_k - \delta_k^2 \geq \beta$ for all $k \geq 1$ and some $\beta > 0$ and thus

$$r_k^2 \leq r_{k-1}^2 - \beta \|\mathbf{X}^k - \mathbf{X}^*\|_F^2. \quad (4.4)$$

Two properties follow from this:

1. The sequence $\{\|\mathcal{P}_\Omega(\mathbf{Y}^k - \mathbf{Y}^*)\|_F\}$ is nonincreasing and, therefore, converges to a limit.
2. As a consequence, $\|\mathbf{X}^k - \mathbf{X}^*\|_F^2 \rightarrow 0$ as $k \rightarrow \infty$.

The theorem is established. $\square$

4.2 General convergence theorem

Our second result is more general and establishes the convergence of the SVT iterations to the solution of (3.4) under general convex constraints. From now on, we will only assume that the function $\mathcal{F}(\mathbf{X})$ is Lipschitz in the sense that

$$\|\mathcal{F}(\mathbf{X}) - \mathcal{F}(\mathbf{Y})\| \leq L(\mathcal{F})\|\mathbf{X} - \mathbf{Y}\|_F, \quad (4.5)$$

for some nonnegative constant $L(\mathcal{F})$. Note that if $\mathcal{F}$ is affine, $\mathcal{F}(\mathbf{X}) = \mathbf{b} - \mathcal{A}(\mathbf{X})$, we have $L(\mathcal{F}) = \|\mathcal{A}\|_2$ where $\|\mathcal{A}\|_2$ is the spectrum norm of the linear transformation $\mathcal{A}$ defined as $\|\mathcal{A}\|_2 := \sup\{\|\mathcal{A}(\mathbf{X})\|_{\ell_2} : \|\mathbf{X}\|_F = 1\}$. We also recall that $\mathcal{F}(\mathbf{X}) = (f_1(\mathbf{X}), \dots, f_m(\mathbf{X}))$ where each f_i is convex, and that the Lagrangian for the problem (3.4) is given by

$$\mathcal{L}(\mathbf{X}, \mathbf{y}) = f_\tau(\mathbf{X}) + \langle \mathbf{y}, \mathcal{F}(\mathbf{X}) \rangle, \quad \mathbf{y} \geq \mathbf{0}.$$

We will assume to simplify that strong duality holds which is automatically true if the constraints obey constraint qualifications such as Slater's condition [6].

We first establish the following preparatory lemma.

Lemma 4.3 *Let $(\mathbf{X}^*, \mathbf{y}^*)$ be a primal-dual optimal pair for (3.4). Then for each $\delta > 0$, $\mathbf{y}^*$ obeys*

$$\mathbf{y}^* = [\mathbf{y}^* + \delta \mathcal{F}(\mathbf{X}^*)]_+. \quad (4.6)$$

Proof. Recall that the projection $\mathbf{x}_0$ of a point $\mathbf{x}$ onto a convex set $\mathcal{C}$ is characterized by

$$\begin{cases} \mathbf{x}_0 \in \mathcal{C}, \\ \langle \mathbf{y} - \mathbf{x}_0, \mathbf{x} - \mathbf{x}_0 \rangle \leq 0, \quad \forall \mathbf{y} \in \mathcal{C}. \end{cases}$$

In the case where $\mathcal{C} = \mathbb{R}_+^m = \{\mathbf{x} \in \mathbb{R}^m : \mathbf{x} \geq \mathbf{0}\}$, this condition becomes $\mathbf{x}_0 \geq \mathbf{0}$ and

$$\langle \mathbf{y} - \mathbf{x}_0, \mathbf{x} - \mathbf{x}_0 \rangle \leq 0, \quad \forall \mathbf{y} \geq \mathbf{0}.$$

Now because $\mathbf{y}^*$ is dual optimal we have

$$\mathcal{L}(\mathbf{X}^*, \mathbf{y}^*) \geq \mathcal{L}(\mathbf{X}^*, \mathbf{y}), \quad \forall \mathbf{y} \geq \mathbf{0}.$$

Substituting the expression for the Lagrangian, this is equivalent to

$$\langle \mathbf{y} - \mathbf{y}^*, \mathcal{F}(\mathbf{X}^*) \rangle \leq 0, \quad \forall \mathbf{y} \geq \mathbf{0},$$

which is the same as

$$\langle \mathbf{y} - \mathbf{y}^*, \mathbf{y}^* + \rho \mathcal{F}(\mathbf{X}^*) - \mathbf{y}^* \rangle \leq 0, \quad \forall \mathbf{y} \geq \mathbf{0}, \quad \forall \rho \geq 0.$$

Hence it follows that $\mathbf{y}^*$ must be the projection of $\mathbf{y}^* + \rho \mathcal{F}(\mathbf{X}^*)$ onto the nonnegative orthant $\mathbb{R}_+^m$. Since the projection of an arbitrary vector $\mathbf{x}$ onto $\mathbb{R}_+^m$ is given by $\mathbf{x}_+$, our claim follows. $\square$

We are now in the position to state our general convergence result.

Theorem 4.4 *Suppose that the sequence of step sizes obeys $0 < \inf \delta_k \leq \sup \delta_k < 2/\|L(\mathcal{F})\|^2$, where $L(\mathcal{F})$ is the Lipschitz constant in (4.5). Then assuming strong duality, the sequence $\{\mathbf{X}^k\}$ obtained via (3.5) converges to the unique solution of (3.4).*

Proof. Let $(\mathbf{X}^*, \mathbf{y}^*)$ be primal-dual optimal for the problem (3.4). We claim that the optimality conditions give that for all $\mathbf{X}$

$$\begin{aligned}\langle \mathbf{Z}^k, \mathbf{X} - \mathbf{X}^k \rangle + \langle \mathbf{y}^{k-1}, \mathcal{F}(\mathbf{X}) - \mathcal{F}(\mathbf{X}^k) \rangle &\geq 0, \\ \langle \mathbf{Z}^*, \mathbf{X} - \mathbf{X}^* \rangle + \langle \mathbf{y}^*, \mathcal{F}(\mathbf{X}) - \mathcal{F}(\mathbf{X}^*) \rangle &\geq 0,\end{aligned}\tag{4.7}$$

for some $\mathbf{Z}^k \in \partial f_\tau(\mathbf{X}^k)$ and some $\mathbf{Z}^* \in \partial f_\tau(\mathbf{X}^*)$. We justify this assertion by proving one of the two inequalities since the other is exactly similar. For the first, $\mathbf{X}^k$ minimizes $\mathcal{L}(\mathbf{X}, \mathbf{y}^{k-1})$ over all $\mathbf{X}$ and, therefore, there exist $\mathbf{Z}^k \in \partial f_\tau(\mathbf{X}^k)$ and $\mathbf{Z}_i^k \in \partial f_i(\mathbf{X}^k)$, $1 \leq i \leq m$, such that

$$\mathbf{Z}^k + \sum_{i=1}^m y_i^{k-1} \mathbf{Z}_i^k = 0.$$

Now because each f_i is convex,

$$f_i(\mathbf{X}) - f_i(\mathbf{X}^k) \geq \langle \mathbf{Z}_i^k, \mathbf{X} - \mathbf{X}^k \rangle$$

and, therefore,

$$\langle \mathbf{Z}^k, \mathbf{X} - \mathbf{X}^k \rangle + \sum_{i=1}^m y_i^{k-1} (f_i(\mathbf{X}) - f_i(\mathbf{X}^k)) \geq \langle \mathbf{Z}^k + \sum_{i=1}^m y_i^{k-1} \mathbf{Z}_i^k, \mathbf{X} - \mathbf{X}^k \rangle = 0.$$

This is (4.7).

Now write the first inequality in (4.7) for $\mathbf{X}^*$, the second for $\mathbf{X}^k$ and sum the two inequalities. This gives

$$\langle \mathbf{Z}^k - \mathbf{Z}^*, \mathbf{X}^k - \mathbf{X}^* \rangle + \langle \mathbf{y}^{k-1} - \mathbf{y}^*, \mathcal{F}(\mathbf{X}^k) - \mathcal{F}(\mathbf{X}^*) \rangle \leq 0.$$

The rest of the proof is essentially the same as that of Theorem 4.5. It follows from Lemma 4.1 that

$$\langle \mathbf{y}^{k-1} - \mathbf{y}^*, \mathcal{F}(\mathbf{X}^k) - \mathcal{F}(\mathbf{X}^*) \rangle \leq -\langle \mathbf{Z}^k - \mathbf{Z}^*, \mathbf{X}^k - \mathbf{X}^* \rangle \leq -\|\mathbf{X}^k - \mathbf{X}^*\|_F^2.\tag{4.8}$$

We continue and observe that because $\mathbf{y}^* = [\mathbf{y}^* + \delta_k \mathcal{F}(\mathbf{X})]_+$ by Lemma 4.3, we have

$$\begin{aligned}\|\mathbf{y}^k - \mathbf{y}^*\| &= \|[\mathbf{y}^{k-1} + \delta_k \mathcal{F}(\mathbf{X}^k)]_+ - [\mathbf{y}^* + \delta_k \mathcal{F}(\mathbf{X}^*)]_+\| \\ &\leq \|\mathbf{y}^{k-1} - \mathbf{y}^* + \delta_k (\mathcal{F}(\mathbf{X}^k) - \mathcal{F}(\mathbf{X}^*))\|\end{aligned}$$

since the projection onto the convex set $\mathbb{R}_+^m$ is a contraction. Therefore,

$$\begin{aligned}\|\mathbf{y}^k - \mathbf{y}^*\|^2 &= \|\mathbf{y}^{k-1} - \mathbf{y}^*\|^2 + 2\delta_k \langle \mathbf{y}^{k-1} - \mathbf{y}^*, \mathcal{F}(\mathbf{X}^k) - \mathcal{F}(\mathbf{X}^*) \rangle + \delta_k^2 \|\mathcal{F}(\mathbf{X}^k) - \mathcal{F}(\mathbf{X}^*)\|^2 \\ &\leq \|\mathbf{y}^{k-1} - \mathbf{y}^*\|^2 - 2\delta_k \|\mathbf{X}^k - \mathbf{X}^*\|_F^2 + \delta_k^2 L^2 \|\mathbf{X}^k - \mathbf{X}^*\|_F^2,\end{aligned}$$

where we have put L instead of $L(\mathcal{F})$ for short. Under our assumptions about the size of δ_k , we have $2\delta_k - \delta_k^2 L^2 \geq \beta$ for all $k \geq 1$ and some $\beta > 0$. Then

$$\|\mathbf{y}^k - \mathbf{y}^*\|^2 \leq \|\mathbf{y}^{k-1} - \mathbf{y}^*\|^2 - \beta \|\mathbf{X}^k - \mathbf{X}^*\|_F^2,\tag{4.9}$$

and the conclusion is as before. □

The problem (3.1) with linear constraints can be reduced to (3.4) by choosing

$$\mathcal{F}(\mathbf{X}) = \begin{bmatrix} \mathbf{b} \\ -\mathbf{b} \end{bmatrix} - \begin{bmatrix} \mathcal{A} \\ -\mathcal{A} \end{bmatrix} \mathbf{X},$$

and we have the following corollary:

Corollary 4.5 *Suppose that the sequence of step sizes obeys $0 < \inf \delta_k \leq \sup \delta_k < 2/\|\mathcal{A}\|_2^2$. Then the sequence $\{\mathbf{X}^k\}$ obtained via (3.3) converges to the unique solution of (3.1).*

Let $\|\mathcal{A}\|_2 := \sup\{\|\mathcal{A}(\mathbf{X})\|_F : \|\mathbf{X}\|_F = 1\}$. With $\mathcal{F}(\mathbf{X})$ given as above, we have $|L(\mathcal{F})|^2 = 2\|\mathcal{A}\|_2^2$ and thus, Theorem 4.4 guarantees convergence as long as $0 < \inf \delta_k \leq \sup \delta_k < 1/\|\mathcal{A}\|_2^2$. However, an argument identical to the proof of Theorem 4.2 would remove the extra factor of two. We omit the details.

5 Implementation and Numerical Results

This section provides implementation details of the SVT algorithm—as to make it practically effective for matrix completion—such as the numerical evaluation of the singular value thresholding operator, the selection of the step size δ_k , the selection of a stopping criterion, and so on. This section also introduces several numerical simulation results which demonstrate the performance and effectiveness of the SVT algorithm. We show that $30,000 \times 30,000$ matrices of rank 10 are recovered from just about 0.4% of their sampled entries in a matter of a few minutes on a modest desktop computer with a 1.86 GHz CPU (dual core with Matlab’s multithreading option enabled) and 3 GB of memory.

5.1 Implementation details

5.1.1 Evaluation of the singular value thresholding operator

To apply the singular value thresholding operator at level τ to an input matrix, it suffices to know those singular values and corresponding singular vectors above the threshold τ . In the matrix completion problem, the singular value thresholding operator is applied to sparse matrices $\{\mathbf{Y}^k\}$ since the number of sampled entries is typically much lower than the number of entries in the unknown matrix $\mathbf{M}$, and we are hence interested in numerical methods for computing the dominant singular values and singular vectors of large sparse matrices. The development of such methods is a relatively mature area in scientific computing and numerical linear algebra in particular. In fact, many high-quality packages are readily available. Our implementation uses PROPACK, see [36] for documentation and availability. One reason for this choice is convenience: PROPACK comes in a Matlab and a Fortran version, and we find it convenient to use the well-documented Matlab version. More importantly, PROPACK uses the iterative Lanczos algorithm to compute the singular values and singular vectors directly, by using the Lanczos bidiagonalization algorithm with partial reorthogonalization. In particular, PROPACK does not compute the eigenvalues and eigenvectors of $(\mathbf{Y}^k)^* \mathbf{Y}^k$ and $\mathbf{Y}^k (\mathbf{Y}^k)^*$, or of an augmented matrix as in the Matlab built-in function ‘svds’ for example. Consequently, PROPACK is an efficient—both in terms of number of flops and storage requirement—and stable package for computing the dominant singular values and singular vectors

of a large sparse matrix. For information, the available documentation [36] reports a speedup factor of about ten over Matlab’s ‘svds’. Furthermore, the Fortran version of PROPACK is about 3–4 times faster than the Matlab version. Despite this significant speedup, we have only used the Matlab version but since the singular value shrinkage operator is by-and-large the dominant cost in the SVT algorithm, we expect that a Fortran implementation would run about 3 to 4 times faster.

As for most SVD packages, though one can specify the number of singular values to compute, PROPACK can not automatically compute only those singular values exceeding the threshold τ . One must instead specify the number s of singular values ahead of time, and the software will compute the s largest singular values and corresponding singular vectors. To use this package, we must then determine the number s_k of singular values of $\mathbf{Y}^{k-1}$ to be computed at the k th iteration. We use the following simple method. Let $r_{k-1} = \text{rank}(\mathbf{X}^{k-1})$ be the number of nonzero singular values of $\mathbf{X}^{k-1}$ at the previous iteration. Set $s_k = r_{k-1} + 1$ and compute the first s_k singular values of $\mathbf{Y}^{k-1}$. If some of the computed singular values are already smaller than τ , then s_k is a right choice. Otherwise, increment s_k by a predefined integer ℓ repeatedly until some of the singular values fall below τ . In the experiments, we choose $\ell = 5$. Another rule might be to repeatedly multiply s_k by a positive number—e.g. 2—until our criterion is met. Incrementing s_k by a fixed integer works very well in practice; in our experiments, we very rarely need more than one update.

We note that it is not necessary to rerun the Lanczos iterations for the first s_k vectors since they have been already computed; only a few new singular values (ℓ of them) need to be numerically evaluated. This can be done by modifying the PROPACK routines. We have not yet modified PROPACK, however. Had we done so, our run times would be decreased.

5.1.2 Step sizes

There is a large literature on ways of selecting a step size but for simplicity, we shall use step sizes that are independent of the iteration count; that is $\delta_k = \delta$ for $k = 1, 2, \dots$. From Theorem 4.2, convergence for the completion problem is guaranteed (2.7) provided that $0 < \delta < 2$. This choice is, however, too conservative and the convergence is typically slow. In our experiments, we use instead

$$\delta = 1.2 \frac{n_1 n_2}{m}, \quad (5.1)$$

i.e. 1.2 times the undersampling ratio. We give a heuristic justification below.

Consider a fixed matrix $\mathbf{A} \in \mathbb{R}^{n_1 \times n_2}$. Under the assumption that the column and row spaces of $\mathbf{A}$ are not well aligned with the vectors taken from the canonical basis of $\mathbb{R}^{n_1}$ and $\mathbb{R}^{n_2}$ respectively—the *incoherence assumption* in [13]—then with very large probability over the choices of Ω , we have

$$(1 - \epsilon)p \|\mathbf{A}\|_F^2 \leq \|\mathcal{P}_\Omega(\mathbf{A})\|_F^2 \leq (1 + \epsilon)p \|\mathbf{A}\|_F^2, \quad p := m/(n_1 n_2), \quad (5.2)$$

provided that the rank of $\mathbf{A}$ is not too large. The probability model is that Ω is a set of sampled entries of cardinality m sampled uniformly at random so that all the choices are equally likely. In (5.2), we want to think of ϵ as a small constant, e.g. smaller than 1/2. In other words, the ‘energy’ of $\mathbf{A}$ on Ω (the set of sampled entries) is just about proportional to the size of Ω . The near isometry (5.2) is a consequence of Theorem 4.1 in [13], and we omit the details.

Now returning to the proof of Theorem 4.2, we see that a sufficient condition for the convergence of (2.7) is

$$\exists \beta > 0, \quad -2\delta \|\mathbf{X}^* - \mathbf{X}^k\|_F^2 + \delta^2 \|\mathcal{P}_\Omega(\mathbf{X}^* - \mathbf{X}^k)\|_F^2 \leq -\beta \|\mathbf{X}^* - \mathbf{X}^k\|_F^2,$$

compare (4.4), which is equivalent to

$$0 < \delta < 2 \frac{\|\mathbf{X}^\star - \mathbf{X}^k\|_F^2}{\|\mathcal{P}_\Omega(\mathbf{X}^\star - \mathbf{X}^k)\|_F^2}.$$

Since $\|\mathcal{P}_\Omega(\mathbf{X})\|_F \leq \|\mathbf{X}\|_F$ for any matrix $\mathbf{X} \in \mathbb{R}^{n_1 \times n_2}$, it is safe to select $\delta < 2$. But suppose that we could apply (5.2) to the matrix $\mathbf{A} = \mathbf{X}^\star - \mathbf{X}^k$. Then we could take δ inversely proportional to p ; e.g. with $\epsilon = 1/4$, we could take $\delta \leq 1.6p^{-1}$. Below, we shall use the value $\delta = 1.2p^{-1}$ which allows us to take large steps and still provides convergence, at least empirically.

The reason why this is not a rigorous argument is that (5.2) cannot be applied to $\mathbf{A} = \mathbf{X}^\star - \mathbf{X}^k$ even though this matrix difference may obey the incoherence assumption. The issue here is that $\mathbf{X}^\star - \mathbf{X}^k$ is not a fixed matrix, but rather depends on Ω since the iterates $\{\mathbf{X}^k\}$ are computed with the knowledge of the sampled set.

5.1.3 Initial steps

The SVT algorithm starts with $\mathbf{Y}^0 = \mathbf{0}$, and we want to choose a large τ to make sure that the solution of (2.8) is close enough to a solution of (1.1). Define k_0 as that integer obeying

$$\frac{\tau}{\delta \|\mathcal{P}_\Omega(\mathbf{M})\|_2} \in (k_0 - 1, k_0]. \quad (5.3)$$

Since $\mathbf{Y}^0 = \mathbf{0}$, it is not difficult to see that

$$\mathbf{X}^k = \mathbf{0}, \quad \mathbf{Y}^k = k\delta \mathcal{P}_\Omega(\mathbf{M}), \quad k = 1, \dots, k_0.$$

To save work, we may simply skip the computations of $\mathbf{X}^1, \dots, \mathbf{X}^{k_0}$, and start the iteration by computing $\mathbf{X}^{k_0+1}$ from $\mathbf{Y}^{k_0}$.

This strategy is a special case of a *kicking device* introduced in [44]; the main idea of such a kicking scheme is that one can ‘jump over’ a few steps whenever possible. Just like in the aforementioned reference, we can develop similar kicking strategies here as well. Because in our numerical experiments the kicking is rarely triggered, we forgo the description of such strategies.

5.1.4 Stopping criteria

Here, we discuss stopping criteria for the sequence of SVT iterations (2.7), and present two possibilities.

The first is motivated by the first-order optimality conditions or KKT conditions tailored to the minimization problem (2.8). By (2.14) and letting $\partial_{\mathbf{Y}} g_0(\mathbf{Y}) = \mathbf{0}$ in (2.13), we see that the solution $\mathbf{X}_\tau^\star$ to (2.8) must also verify

$$\begin{cases} \mathbf{X} = \mathcal{D}_\tau(\mathbf{Y}), \\ \mathcal{P}_\Omega(\mathbf{X} - \mathbf{M}) = \mathbf{0}, \end{cases} \quad (5.4)$$

where $\mathbf{Y}$ is a matrix vanishing outside of Ω^c . Therefore, to make sure that $\mathbf{X}^k$ is close to $\mathbf{X}_\tau^\star$, it is sufficient to check how close $(\mathbf{X}^k, \mathbf{Y}^{k-1})$ is to obeying (5.4). By definition, the first equation in (5.4) is always true. Therefore, it is natural to stop (2.7) when the error in the second equation is below a specified tolerance. We suggest stopping the algorithm when

$$\frac{\|\mathcal{P}_\Omega(\mathbf{X}^k - \mathbf{M})\|_F}{\|\mathcal{P}_\Omega(\mathbf{M})\|_F} \leq \epsilon, \quad (5.5)$$

where ϵ is a fixed tolerance, e.g. 10^{-4} . We provide a short heuristic argument justifying this choice below.

In the matrix completion problem, we know that under suitable assumptions

$$\|\mathcal{P}_\Omega(\mathbf{M})\|_F^2 \asymp p \|\mathbf{M}\|_F^2,$$

which is just (5.2) applied to the fixed matrix $\mathbf{M}$ (the symbol $\asymp$ here means that there is a constant ϵ as in (5.2)). Suppose we could also apply (5.2) to the matrix $\mathbf{X}^k - \mathbf{M}$ (which we rigorously cannot since $\mathbf{X}^k$ depends on Ω), then we would have

$$\|\mathcal{P}_\Omega(\mathbf{X}^k - \mathbf{M})\|_F^2 \asymp p \|\mathbf{X}^k - \mathbf{M}\|_F^2, \quad (5.6)$$

and thus

$$\frac{\|\mathcal{P}_\Omega(\mathbf{X}^k - \mathbf{M})\|_F}{\|\mathcal{P}_\Omega(\mathbf{M})\|_F} \asymp \frac{\|\mathbf{X}^k - \mathbf{M}\|_F}{\|\mathbf{M}\|_F}.$$

In words, one would control the relative reconstruction error by controlling the relative error on the set of sampled locations.

A second stopping criterion comes from duality theory. Firstly, the iterates $\mathbf{X}^k$ are generally not feasible for (2.8) although they become asymptotically feasible. One can construct a feasible point from $\mathbf{X}^k$ by projecting it onto the affine space $\{\mathbf{X} : \mathcal{P}_\Omega(\mathbf{X}) = \mathcal{P}_\Omega(\mathbf{M})\}$ as follows:

$$\tilde{\mathbf{X}}^k = \mathbf{X}^k + \mathcal{P}_\Omega(\mathbf{M} - \mathbf{X}^k).$$

As usual let $f_\tau(\mathbf{X}) = \tau \|\mathbf{X}\|_* + \frac{1}{2} \|\mathbf{X}\|_F^2$ and denote by p^* the optimal value of (2.8). Since $\tilde{\mathbf{X}}^k$ is feasible, we have

$$p^* \leq f_\tau(\tilde{\mathbf{X}}^k) := b_k.$$

Secondly, using the notations of Section 2.4, duality theory gives that

$$a_k := g_0(\mathbf{Y}^{k-1}) = \mathcal{L}(\mathbf{X}^k, \mathbf{Y}^{k-1}) \leq p^*.$$

Therefore, $b_k - a_k$ is an upper bound on the duality gap and one can stop the algorithm when this quantity falls below a given tolerance.

For very large problems in which one holds $\mathbf{X}^k$ in reduced SVD form, one may not want to compute the projection $\tilde{\mathbf{X}}^k$ since this matrix would not have low rank and would require significant storage space (presumably, one would not want to spend much time computing this projection either). Hence, the second method only makes practical sense when the dimensions are not prohibitively large, or when the iterates do not have low rank.

5.1.5 Algorithm

We conclude this section by summarizing the implementation details and give the SVT algorithm for matrix completion below (Algorithm 1). Of course, one would obtain a very similar structure for the more general problems of the form (3.1) and (3.4) with linear inequality constraints. For convenience, define for each nonnegative integer $s \leq \min\{n_1, n_2\}$,

$$[\mathbf{U}^k, \mathbf{\Sigma}^k, \mathbf{V}^k]_s, \quad k = 1, 2, \dots,$$

where $\mathbf{U}^k = [\mathbf{u}_1^k, \dots, \mathbf{u}_s^k]$ and $\mathbf{V}^k = [\mathbf{v}_1^k, \dots, \mathbf{v}_s^k]$ are the first s singular vectors of the matrix $\mathbf{Y}^k$, and $\mathbf{\Sigma}^k$ is a diagonal matrix with the first s singular values $\sigma_1^k, \dots, \sigma_s^k$ on the diagonal.

Algorithm 1: Singular Value Thresholding (SVT) Algorithm

Input: sampled set Ω and sampled entries $\mathcal{P}_\Omega(\mathbf{M})$, step size δ , tolerance ϵ , parameter τ , increment ℓ , and maximum iteration count $k_{\max}$

Output: $\mathbf{X}^{\text{opt}}$

Description: Recover a low-rank matrix $\mathbf{M}$ from a subset of sampled entries

```

1  Set  $\mathbf{Y}^0 = k_0 \delta \mathcal{P}_\Omega(\mathbf{M})$  ( $k_0$  is defined in (5.3))
2  Set  $r_0 = 0$ 
3  for  $k = 1$  to  $k_{\max}$ 
4      Set  $s_k = r_{k-1} + 1$ 
5      repeat
6          Compute  $[\mathbf{U}^{k-1}, \mathbf{\Sigma}^{k-1}, \mathbf{V}^{k-1}]_{s_k}$ 
7          Set  $s_k = s_k + \ell$ 
8      until  $\sigma_{s_k-\ell}^{k-1} \leq \tau$ 
9      Set  $r_k = \max\{j : \sigma_j^{k-1} > \tau\}$ 
10     Set  $\mathbf{X}^k = \sum_{j=1}^{r_k} (\sigma_j^{k-1} - \tau) \mathbf{u}_j^{k-1} \mathbf{v}_j^{k-1}$ 
11
12     if  $\|\mathcal{P}_\Omega(\mathbf{X}^k - \mathbf{M})\|_F / \|\mathcal{P}_\Omega \mathbf{M}\|_F \leq \epsilon$  then break
13
14     Set  $Y_{ij}^k = \begin{cases} 0 & \text{if } (i, j) \notin \Omega, \\ Y_{ij}^{k-1} + \delta(M_{ij} - X_{ij}^k) & \text{if } (i, j) \in \Omega \end{cases}$ 
15 end for  $k$ 
16 Set  $\mathbf{X}^{\text{opt}} = \mathbf{X}^k$ 

```

5.2 Numerical results

5.2.1 Linear equality constraints

Our implementation is in Matlab and all the computational results we are about to report were obtained on a desktop computer with a 1.86 GHz CPU (dual core with Matlab's multithreading option enabled) and 3 GB of memory. In our simulations, we generate $n \times n$ matrices of rank r by sampling two $n \times r$ factors $\mathbf{M}_L$ and $\mathbf{M}_R$ independently, each having i.i.d. Gaussian entries, and setting $\mathbf{M} = \mathbf{M}_L \mathbf{M}_R^*$ as it is suggested in [13]. The set of observed entries Ω is sampled uniformly at random among all sets of cardinality m .

The recovery is performed via the SVT algorithm (Algorithm 1), and we use

$$\|\mathcal{P}_\Omega(\mathbf{X}^k - \mathbf{M})\|_F / \|\mathcal{P}_\Omega \mathbf{M}\|_F < 10^{-4} \quad (5.7)$$

as a stopping criterion. As discussed earlier, the step sizes are constant and we set $\delta = 1.2p^{-1}$. Throughout this section, we denote the output of the SVT algorithm by $\mathbf{X}^{\text{opt}}$. The parameter τ is chosen empirically and set to $\tau = 5n$. A heuristic argument is as follows. Clearly, we would like the term $\tau \|\mathbf{M}\|_*$ to dominate the other, namely, $\frac{1}{2} \|\mathbf{M}\|_F^2$. For products of Gaussian matrices as above, standard random matrix theory asserts that the Frobenius norm of $\mathbf{M}$ concentrates around $n\sqrt{r}$, and that the nuclear norm concentrates around about nr (this should be clear in the simple case where $r = 1$ and is generally valid). The value $\tau = 5n$ makes sure that on the average, the

Unknown $\mathbf{M}$				Computational results		
size ($n \times n$)	rank (r)	m/d_r	m/n^2	time(s)	# iters	relative error
$1,000 \times 1,000$	10	6	0.12	23	117	1.64×10^{-4}
	50	4	0.39	196	114	1.59×10^{-4}
	100	3	0.57	501	129	1.68×10^{-4}
$5,000 \times 5,000$	10	6	0.024	147	123	1.73×10^{-4}
	50	5	0.10	950	108	1.61×10^{-4}
	100	4	0.158	3,339	123	1.72×10^{-4}
$10,000 \times 10,000$	10	6	0.012	281	123	1.73×10^{-4}
	50	5	0.050	2,096	110	1.65×10^{-4}
	100	4	0.080	7,059	127	1.79×10^{-4}
$20,000 \times 20,000$	10	6	0.006	588	124	1.73×10^{-4}
	50	5	0.025	4,581	111	1.66×10^{-4}
$30,000 \times 30,000$	10	6	0.004	1,030	125	1.73×10^{-4}

Table 1: Experimental results for matrix completion. The rank r is the rank of the unknown matrix $\mathbf{M}$, m/d_r is the ratio between the number of sampled entries and the number of degrees of freedom in an $n \times n$ matrix of rank r (oversampling ratio), and m/n^2 is the fraction of observed entries. All the computational results on the right are averaged over five runs.

value of $\tau\|\mathbf{M}\|_*$ is about 10 times that of $\frac{1}{2}\|\mathbf{M}\|_F^2$ as long as the rank is bounded away from the dimension n .

Our computational results are displayed in Table 1. There, we report the run time in seconds, the number of iterations it takes to reach convergence (5.7), and the relative error of the reconstruction

$$\text{relative error} = \|\mathbf{X}^{\text{opt}} - \mathbf{M}\|_F / \|\mathbf{M}\|_F, \quad (5.8)$$

where $\mathbf{M}$ is the real unknown matrix. All of these quantities are averaged over five runs. The table also gives the percentage of entries that are observed, namely, m/n^2 together with a quantity that we may want to think as the information oversampling ratio. Recall that an $n \times n$ matrix of rank r depends upon $d_r := r(2n - r)$ degrees of freedom. Then m/d_r is the ratio between the number of sampled entries and the ‘true dimensionality’ of an $n \times n$ matrix of rank r .

The first observation is that the SVT algorithm performs extremely well in these experiments. In all of our experiments, it takes fewer than 200 SVT iterations to reach convergence. As a consequence, the run times are short. As indicated in the table, we note that one recovers a $1,000 \times 1,000$ matrix of rank 10 in less than a minute. The algorithm also recovers $30,000 \times 30,000$ matrices of rank 10 from about 0.4% of their sampled entries in just about 17 minutes. In addition, higher-rank matrices are also efficiently completed: for example, it takes between one and two hours to recover $10,000 \times 10,000$ matrices of rank 100 and $20,000 \times 20,000$ matrices of rank 50. We would like to stress that these numbers were obtained on a modest CPU (1.86GHz). Furthermore, a Fortran implementation is likely to cut down on these numbers by a multiplicative factor typically between three and four.

We also check the validity of the stopping criterion (5.7) by inspecting the relative error defined in (5.8). The table shows that the heuristic and nonrigorous analysis of Section 5.1 holds in practice since the relative reconstruction error is of the same order as $\|\mathcal{P}_\Omega(\mathbf{X}^{\text{opt}} - \mathbf{M})\|_F / \|\mathcal{P}_\Omega \mathbf{M}\|_F \sim 10^{-4}$. Indeed, the overall relative errors reported in Table 1 are all less than 2×10^{-4} .

We emphasized all along an important feature of the SVT algorithm, which is that the matrices $\mathbf{X}^k$ have low rank. We demonstrate this fact empirically in Figure 1, which plots the rank of $\mathbf{X}^k$ versus the iteration count k , and does this for unknown matrices of size $5,000 \times 5,000$ with different ranks. The plots reveal an interesting phenomenon: in our experiments, the rank of $\mathbf{X}^k$ is nondecreasing so that the maximum rank is reached in the final steps of the algorithm. In fact, the rank of the iterates quickly reaches the value r of the true rank. After these few initial steps, the SVT iterations search for that matrix with rank r minimizing the objective functional. As mentioned earlier, the low-rank property is crucial for making the algorithm run fast.

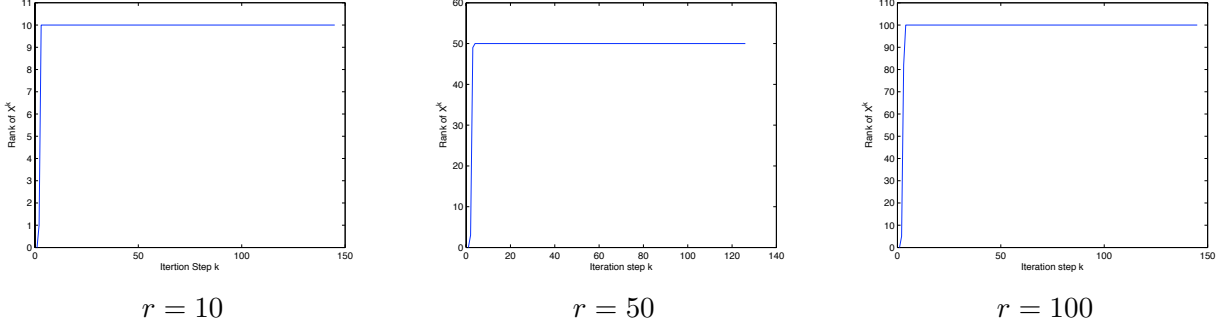


Figure 1: Rank of $\mathbf{X}^k$ as a function k when the unknown matrix $\mathbf{M}$ is of size $5,000 \times 5,000$ and of rank r .

Finally, we demonstrate the results of the SVT algorithm for matrix completion from noisy sampled entries. Suppose we observe data from the model

$$B_{ij} = M_{ij} + Z_{ij}, \quad (i, j) \in \Omega, \quad (5.9)$$

where $\mathbf{Z}$ is a zero-mean Gaussian white noise with standard deviation σ . We run the SVT algorithm but stop early, as soon as $\mathbf{X}^k$ is consistent with the data and obeys

$$\|\mathcal{P}_\Omega(\mathbf{X}^k - \mathbf{B})\|_F^2 \leq (1 + \epsilon) m \sigma^2, \quad (5.10)$$

where ϵ is a small parameter. Our reconstruction $\hat{\mathbf{M}}$ is the first $\mathbf{X}^k$ obeying (5.10). The results are shown in Table 2 (the quantities are averages of 5 runs). Define the noise ratio as

$$\|\mathcal{P}_\Omega(\mathbf{Z})\|_F / \|\mathcal{P}_\Omega(\mathbf{M})\|_F,$$

and the relative error by (5.8). From Table 2, we see that the SVT algorithm works well as the relative error between the recovered and the true data matrix is just about equal to the noise ratio.

The theory of low-rank matrix recovery from noisy data is nonexistent at the moment, and is obviously beyond the scope of this paper. Having said this, we would like to conclude this section with an intuitive and nonrigorous discussion, which may explain why the observed recovery error is within the noise level. Suppose again that $\hat{\mathbf{M}}$ obeys (5.6), namely,

$$\|\mathcal{P}_\Omega(\hat{\mathbf{M}} - \mathbf{M})\|_F^2 \asymp p \|\hat{\mathbf{M}} - \mathbf{M}\|_F^2. \quad (5.11)$$

noise ratio	Unknown matrix $\mathbf{M}$				Computational results		
	size ($n \times n$)	rank (r)	m/d_r	m/n^2	time(s)	# iters	relative error
10^{-2}	$1,000 \times 1,000$	10	6	0.12	10.8	51	0.78×10^{-2}
		50	4	0.39	87.7	48	0.95×10^{-2}
		100	3	0.57	216	50	1.13×10^{-2}
10^{-1}	$1,000 \times 1,000$	10	6	0.12	4.0	19	0.72×10^{-1}
		50	4	0.39	33.2	17	0.89×10^{-1}
		100	3	0.57	85.2	17	1.01×10^{-1}
1	$1,000 \times 1,000$	10	6	0.12	0.9	3	0.52
		50	4	0.39	7.8	3	0.63
		100	3	0.57	34.8	3	0.69

Table 2: Simulation results for noisy data. The computational results are averaged over five runs.

As mentioned earlier, one condition for this to happen is that $\mathbf{M}$ and $\hat{\mathbf{M}}$ have low rank. This is the reason why it is important to stop the algorithm early as we hope to obtain a solution which is both consistent with the data and has low rank (the limit of the SVT iterations, $\lim_{k \rightarrow \infty} \mathbf{X}^k$, will not generally have low rank since there may be no low-rank matrix matching the noisy data). From

$$\|\mathcal{P}_\Omega(\hat{\mathbf{M}} - \mathbf{M})\|_F \leq \|\mathcal{P}_\Omega(\hat{\mathbf{M}} - \mathbf{B})\|_F + \|\mathcal{P}_\Omega(\mathbf{B} - \mathbf{M})\|_F,$$

and the fact that both terms on the right-hand side are on the order of $\sqrt{m\sigma^2}$, we would have $p\|\hat{\mathbf{M}} - \mathbf{M}\|_F^2 = O(m\sigma^2)$ by (5.11). In particular, this would give that the relative reconstruction error is on the order of the noise ratio since $\|\mathcal{P}_\Omega(\mathbf{M})\|_F^2 \asymp p\|\mathbf{M}\|_F^2$ —as observed experimentally.

5.2.2 Linear inequality constraints

We now examine the speed at which one can solve similar problems with linear inequality constraints instead of linear equality constraints. We assume the model (5.9), where the matrix $\mathbf{M}$ of rank r is sampled as before, and solve the problem (3.8) by using (3.10). We formulate the inequality constraints in (3.8) with $E_{ij} = \sigma$ so that one searches for a solution $\hat{\mathbf{M}}$ with minimum nuclear norm among all those matrices whose sampled entries deviate from the observed ones by at most the noise level σ .² In this experiment, we adjust σ to be one tenth of a typical absolute entry of $\mathbf{M}$, i.e. $\sigma = 0.1 \sum_{ij \in \Omega} |M_{ij}|/m$, and the noise ratio as defined earlier is 0.780. We set $n = 1,000$, $r = 10$, and the number m of sampled entries is five times the number of degrees of freedom, i.e. $m = 5d_r$. Just as before, we set $\tau = 5n$, and choose a constant step size $\delta = 1.2p^{-1}$.

The results, reported in Figure 2, show that the algorithm behaves just as well with linear inequality constraints. To make this point, we compare our results with those obtained from noiseless data (same unknown matrix and sampled locations). In the noiseless case, it takes about 150 iterations to reach the tolerance $\epsilon = 10^{-4}$ whereas in the noisy case, convergence occurs in about 200 iterations (Figure 2(a)). In addition, just as in the noiseless problem, the rank of the iterates is nondecreasing and quickly reaches the true value r of the rank of the unknown matrix

²This may not be conservative enough from a statistical viewpoint but this works well in this case, and our emphasis here is on computational rather than statistical issues.

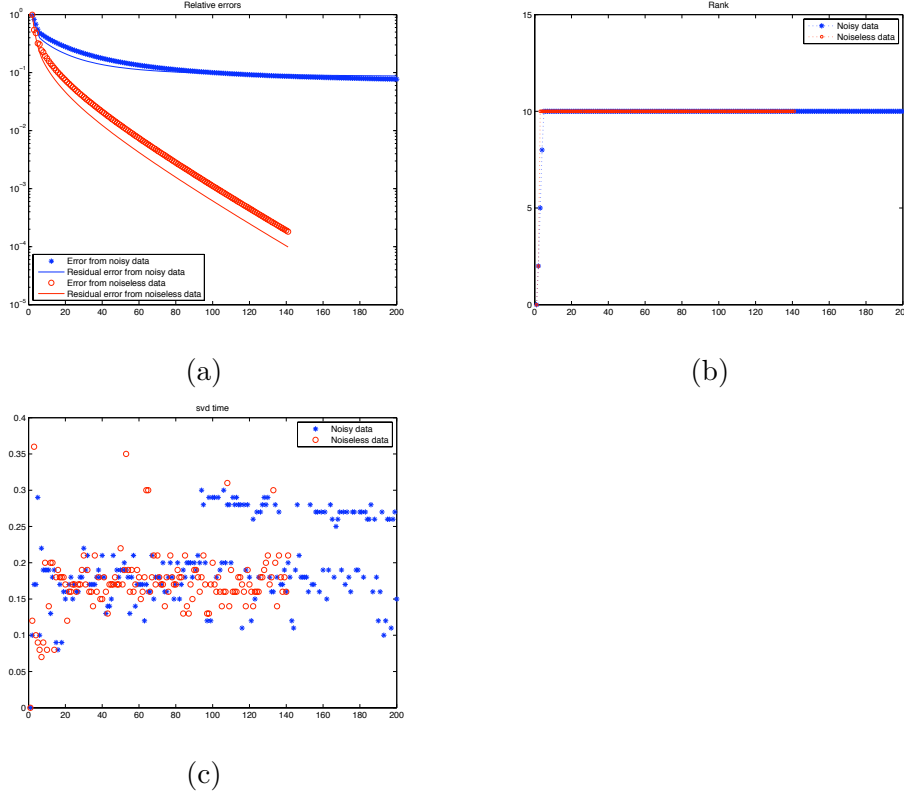


Figure 2: Computational results of the algorithm applied to noisy (linear inequality constraints as in (3.8)) and noiseless data (equality constraints). The blue (resp. red) color is used for the noisy (resp. noiseless) experiment. (a) Plot of the reconstruction errors from noisy and noiseless data as a function of the iteration count. The thin line is the residual relative error $\|\mathcal{P}_\Omega(\mathbf{X}^k - \mathbf{M})\|_F / \|\mathcal{P}_\Omega(\mathbf{M})\|_F$ and the thick line is the overall relative error $\|\mathbf{X}^k - \mathbf{M}\|_F / \|\mathbf{M}\|_F$. (b) Rank of the iterates as a function of the iteration count. (c) Time it takes to compute the singular value thresholding operation as a function of the iteration count. The computer here is a single-core 3.00GHz Pentium 4 running Matlab 7.2.0.

$\mathbf{M}$ we wish to recover (Figure 2(b)). As a consequence the SVT iterations take about the same amount of time as in the noiseless case (Figure 2(c)) so that the total running time of the algorithm does not appear to be substantially different from that in the noiseless case.

We close by pointing out that from a statistical point of view, the recovery of the matrix $\mathbf{M}$ from undersampled and noisy entries by the matrix equivalent of the Dantzig selector appears to be accurate since the relative error obeys $\|\hat{\mathbf{M}} - \mathbf{M}\|_F / \|\mathbf{M}\|_F = 0.0769$ (recall that the noise ratio is about 0.08).

6 Discussion

This paper introduced a novel algorithm, namely, the singular value thresholding algorithm for matrix completion and related nuclear norm minimization problems. This algorithm is easy to implement and surprisingly effective both in terms of computational cost and storage requirement

when the minimum nuclear-norm solution is also the lowest-rank solution. We would like to close this paper by discussing a few open problems and research directions related to this work.

Our algorithm exploits the fact that the sequence of iterates $\{\mathbf{X}^k\}$ have low rank when the minimum nuclear solution has low rank. An interesting question is whether one can prove (or disprove) that in a majority of the cases, this is indeed the case.

It would be interesting to explore other ways of computing $\mathcal{D}_\tau(\mathbf{Y})$ —in words, the action of the singular value shrinkage operator. Our approach uses the Lanczos bidiagonalization algorithm with partial reorthogonalization which takes advantages of sparse inputs but other approaches are possible. We mention two of them.

1. A series of papers have proposed the use of randomized procedures for the approximation of a matrix $\mathbf{Y}$ with a matrix $\mathbf{Z}$ of rank r [38, 41]. When this approximation consists of the truncated SVD retaining the part of the expansion corresponding to singular values greater than τ , this can be used to evaluate $\mathcal{D}_\tau(\mathbf{Y})$. Some of these algorithms are efficient when the input $\mathbf{Y}$ is sparse [41], and it would be interesting to know whether these methods are fast and accurate enough to be used in the SVT iteration (2.7).
2. A wide range of iterative methods for computing matrix functions of the general form $f(\mathbf{Y})$ are available today, see [34] for a survey. A valuable research direction is to investigate whether some of these iterative methods, or other to be developed, would provide powerful ways for computing $\mathcal{D}_\tau(\mathbf{Y})$.

In practice, one would like to solve (2.8) for large values of τ . However, a larger value of τ generally means a slower rate of convergence. A good strategy might be to start with a value of τ , which is large enough so that (2.8) admits a low-rank solution, and at the same time for which the algorithm converges rapidly. One could then use a continuation method as in [50] to increase the value of τ sequentially according to a schedule $\tau_0, \tau_1, \dots$, and use the solution to the previous problem with $\tau = \tau_{i-1}$ as an initial guess for the solution to the current problem with $\tau = \tau_i$ (warm starting). We hope to report on this in a separate paper.

Acknowledgments

J-F. C. is supported by the Wavelets and Information Processing Programme under a grant from DSTA, Singapore. E. C. is partially supported by the Waterman Award from the National Science Foundation and by an ONR grant N00014-08-1-0749. Z. S. is supported in part by Grant R-146-000-113-112 from the National University of Singapore. E. C. would like to thank Benjamin Recht and Joel Tropp for fruitful conversations related to this project, and Stephen Becker for his help in preparing the computational results of Section 5.2.2.

References

- [1] J. Abernethy, F. Bach, T. Evgeniou, and J.-P. Vert. Low-rank matrix factorization with attributes. Technical Report N24/06/MM, Ecole des Mines de Paris, 2006.
- [2] ACM SIGKDD and Netflix. *Proceedings of KDD Cup and Workshop*, 2007. Proceedings available online at <http://www.cs.uic.edu/~liub/KDD-cup-2007/proceedings.html>.
- [3] Y. Amit, M. Fink, N. Srebro, and S. Ullman. Uncovering shared structures in multiclass classification. In *Proceedings of the Twenty-fourth International Conference on Machine Learning*, 2007.

- [4] A. Argyriou, T. Evgeniou, and M. Pontil. Multi-task feature learning. In *Neural Information Processing Systems*, 2007.
- [5] J. Bect, L. Blanc-Féraud, G. Aubert, and A. Chambolle. A ℓ_1 unified variational framework for image restoration, in *Proc. Eighth Europ. Conf. Comput. Vision*, 2004.
- [6] S. Boyd, and L. Vandenberghe. *Convex Optimization*. Cambridge University Press, 2004.
- [7] J.-F. Cai, R. Chan, L. Shen, and Z. Shen. Restoration of chopped and noddled images by framelets. *SIAM J. Sci. Comput.*, 30(3):1205–1227, 2008.
- [8] J.-F. Cai, R. H. Chan, and Z. Shen. A framelet-based image inpainting algorithm. *Appl. Comput. Harmon. Anal.*, 24(2):131–149, 2008.
- [9] J.-F. Cai, S. Osher, and Z. Shen. *Convergence of the Linearized Bregman Iteration for ℓ_1 -norm Minimization*, 2008. UCLA CAM Report (08-52).
- [10] J.-F. Cai, S. Osher, and Z. Shen. *Linearized Bregman Iterations for Compressed Sensing*, 2008. Math. Comp., to appear; see also UCLA CAM Report (08-06).
- [11] J.-F. Cai, S. Osher, and Z. Shen. *Linearized Bregman Iterations for Frame-Based Image Deblurring*, 2008. preprint.
- [12] E. J. Candès, and F. Guo. New multiscale transforms, minimum total variation synthesis: Applications to edge-preserving image reconstruction. *Signal Processing*, 82:1519–1543, 2002.
- [13] E. J. Candès and B. Recht. *Exact Matrix Completion via Convex Optimization*, 2008.
- [14] E. J. Candès and J. Romberg. Sparsity and incoherence in compressive sampling. *Inverse Problems*, 23(3):969–985, 2007.
- [15] E. J. Candès, J. Romberg, and T. Tao. Robust uncertainty principles: exact signal reconstruction from highly incomplete frequency information. *IEEE Trans. Inform. Theory*, 52(2):489–509, 2006.
- [16] E. J. Candès and T. Tao. Decoding by linear programming. *IEEE Trans. Inform. Theory*, 51(12):4203–4215, 2005.
- [17] E. J. Candès and T. Tao. Near-optimal signal recovery from random projections: universal encoding strategies? *IEEE Trans. Inform. Theory*, 52(12):5406–5425, 2006.
- [18] E. J. Candès and T. Tao. The Dantzig selector: statistical estimation when p is much larger than n . *Annals of Statistics* 35:2313–2351, 2007.
- [19] A. Chai and Z. Shen. Deconvolution: A wavelet frame approach. *Numer. Math.*, 106(4):529–587, 2007.
- [20] R. H. Chan, T. F. Chan, L. Shen, and Z. Shen. Wavelet algorithms for high-resolution image reconstruction. *SIAM J. Sci. Comput.*, 24(4):1408–1432 (electronic), 2003.
- [21] P. Chen, and D. Suter. Recovering the missing components in a large noisy low-rank matrix: application to SFM source. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 26(8):1051–1063, 2004.
- [22] P. L. Combettes and V. R. Wajs. Signal recovery by proximal forward-backward splitting. *Multiscale Model. Simul.*, 4(4):1168–1200 (electronic), 2005.
- [23] J. Darbon and S. Osher. *Fast discrete optimization for sparse approximations and deconvolutions*, 2007. preprint.
- [24] I. Daubechies, M. Defrise, and C. De Mol. An iterative thresholding algorithm for linear inverse problems with a sparsity constraint. *Comm. Pure Appl. Math.*, 57(11):1413–1457, 2004.
- [25] I. Daubechies, G. Teschke, and L. Vese. Iteratively solving linear inverse problems under general convex constraints. *Inverse Probl. Imaging*, 1(1):29–46, 2007.

- [26] D. L. Donoho. Compressed sensing. *IEEE Trans. Inform. Theory*, 52(4):1289–1306, 2006.
- [27] M. Elad, J.-L. Starck, P. Querre, and D. L. Donoho. Simultaneous cartoon and texture image inpainting using morphological component analysis (MCA). *Appl. Comput. Harmon. Anal.*, 19(3):340–358, 2005.
- [28] M. J. Fadili, J.-L. Starck, and F. Murtagh. Inpainting and zooming using sparse representations. *The Computer Journal*, to appear.
- [29] M. Fazel. *Matrix Rank Minimization with Applications*. PhD thesis, Stanford University, 2002.
- [30] M. Fazel, H. Hindi, and S. Boyd. Log-det heuristic for matrix rank minimization with applications to Hankel and Euclidean distance matrices. in *Proc. Am. Control Conf.*, June 2003.
- [31] M. Figueiredo, and R. Nowak. An EM algorithm for wavelet-based image restoration. *IEEE Transactions on Image Processing*, 12(8):906–916, 2003.
- [32] T. Goldstein and S. Osher. *The Split Bregman Algorithm for L1 Regularized Problems*, 2008. UCLA CAM Reprots (08-29).
- [33] E. T. Hale, W. Yin, and Y. Zhang. *Fixed-point continuation for l1-minimization: methodology and convergence*. 2008. preprint.
- [34] N. J. Higham. *Functions of Matrices: Theory and Computation*. Society for Industrial and Applied Mathematics, Philadelphia, PA, USA, 2008.
- [35] J.-B. Hiriart-Urruty and C. Lemaréchal. *Convex analysis and minimization algorithms. I*, volume 305 of *Grundlehren der Mathematischen Wissenschaften [Fundamental Principles of Mathematical Sciences]*. Springer-Verlag, Berlin, 1993. Fundamentals.
- [36] R. M. Larsen, *PROPACK – Software for large and sparse SVD calculations*, Available from <http://sun.stanford.edu/~rmunk/PROPACK/>.
- [37] A. S. Lewis. The mathematics of eigenvalue optimization. *Math. Program.*, 97(1-2, Ser. B):155–176, 2003. ISMP, 2003 (Copenhagen).
- [38] E. Liberty, F. Woolfe, P.-G. Martinsson, V. Rokhlin, and M. Tygert. Randomized algorithms for the low-rank approximation of matrices. *Proc. Natl. Acad. Sci. USA*, 104(51): 20167–20172, 2007.
- [39] S. Lintner, and F. Malgouyres. Solving a variational image restoration model which involves ℓ_∞ constraints. *Inverse Problems*, 20:815–831, 2004.
- [40] Z. Liu, and L. Vandenbergh. Interior-point method for nuclear norm approximation with application to system identification. submitted to *Mathematical Programming*, 2008.
- [41] P.-G. Martinsson, V. Rokhlin, and M. Tygert. A randomized algorithm for the approximation of matrices Department of Computer Science, Yale University, New Haven, CT, Technical Report 1361, 2006.
- [42] M. Mesbahi and G. P. Papavassilopoulos. On the rank minimization problem over a positive semidefinite linear matrix inequality. *IEEE Transactions on Automatic Control*, 42(2):239–243, 1997.
- [43] S. Osher, M. Burger, D. Goldfarb, J. Xu, and W. Yin. An iterative regularization method for total variation-based image restoration. *Multiscale Model. Simul.*, 4(2):460–489 (electronic), 2005.
- [44] S. Osher, Y. Mao, B. Dong, and W. Yin. *Fast Linearized Bregman Iteration for Compressed Sensing and Sparse Denoising*, 2008. UCLA CAM Reprots (08-37).
- [45] B. Recht, M. Fazel, and P. Parrilo. Guaranteed minimum rank solutions of matrix equations via nuclear norm minimization. 2007. Submitted to *SIAM Review*.
- [46] J.-L. Starck, D. L. Donoho, and E. J. Candès. Astronomical image representation by the curvelet transform. *Astronom. and Astrophys.*, 398:785–800, 2003.

- [47] K. C. Toh, M. J. Todd, and R. H. Tütüncü. *SDPT3 – a MATLAB software package for semidefinite-quadratic-linear programming*, Available from <http://www.math.nus.edu.sg/~mattohc/sdpt3.html>.
- [48] C. Tomasi and T. Kanade. Shape and motion from image streams under orthography: a factorization method. *International Journal of Computer Vision*, 9(2):137–154, 1992.
- [49] G. A. Watson. Characterization of the subdifferential of some matrix norms. *Linear Algebra Appl.*, 170:33–45, 1992.
- [50] S. J. Wright, R. Nowak, and M. Figueiredo. Sparse reconstruction by separable approximation. Submitted for publication, 2007.
- [51] W. Yin, S. Osher, D. Goldfarb, and J. Darbon. Bregman iterative algorithms for ℓ_1 -minimization with applications to compressed sensing. *SIAM J. Imaging Sci.*, 1(1):143–168, 2008.