



Efficient Solvers for Sparse Subspace Clustering

Farhad Pourkamali-Anaraki^a, James Folberth^b, Stephen Becker^b

^a Department of Computer Science, University of Massachusetts Lowell, MA, USA

^b Department of Applied Mathematics, University of Colorado Boulder, CO, USA

ARTICLE INFO

Article history:

Received 26 June 2019

Revised 18 February 2020

Accepted 19 February 2020

Available online 21 February 2020

Keywords:

Unsupervised learning

Subspace clustering

Large-scale optimization

Sparsity-driven methods

ABSTRACT

Sparse subspace clustering (SSC) clusters n points that lie near a union of low-dimensional subspaces. The SSC model expresses each point as a linear or affine combination of the other points, using either ℓ_1 or ℓ_0 regularization. Using ℓ_1 regularization results in a convex problem but requires $\mathcal{O}(n^2)$ storage, and is typically solved by the alternating direction method of multipliers which takes $\mathcal{O}(n^3)$ flops. The ℓ_0 model is non-convex but only needs memory linear in n , and is solved via orthogonal matching pursuit and cannot handle the case of affine subspaces. This paper shows that a proximal gradient framework can solve SSC, covering both ℓ_1 and ℓ_0 models, and both linear and affine constraints. For both ℓ_1 and ℓ_0 , algorithms to compute the proximity operator in the presence of affine constraints have not been presented in the SSC literature, so we derive an exact and efficient algorithm that solves the ℓ_1 case with just $\mathcal{O}(n^2)$ flops. In the ℓ_0 case, our algorithm retains the low-memory overhead, and is the first algorithm to solve the SSC- ℓ_0 model with affine constraints. Experiments show our algorithms do not rely on sensitive regularization parameters, and they are less sensitive to sparsity misspecification and high noise.

© 2020 Elsevier B.V. All rights reserved.

1. Introduction

In modern data analysis, clustering is an important tool for extracting information from large-scale data sets by identifying groups of similar data points without the presence of ground-truth labels. Therefore, there has been growing interest in developing accurate and efficient clustering algorithms by taking account of the intrinsic structure of large high-dimensional data sets. For instance, the popular K-means algorithm and its kernel-based variants are based on the assumption that (mapped) data points are evenly distributed within linearly separable clusters [1,2].

There has been much work on approaches for more complicated clustering models, such as data that comes from a mixture of manifolds. For some problems, a reasonable assumption is that of data points lying near a union of low-dimensional subspaces [3,4]. The dimensions and orientations of the subspaces are unknown and there are possibly non-trivial intersections between every pair of subspaces. The main task is to partition a given data set such that each group contains only data points from the same subspace. This problem is referred to as “subspace clustering” and has numerous applications in machine learning and computer vision such as motion segmentation and face clustering [5,6].

Among existing subspace clustering techniques, a popular line of work is focused on applying spectral clustering to an affinity matrix obtained by solving a global optimization problem, which represents each data point as a linear or affine combination of other points [7]. Given $\mathbf{x}_1, \dots, \mathbf{x}_n$ that lie near a union of subspaces in $\mathbb{R}^p$, let $\mathbf{X} \in \mathbb{R}^{p \times n}$ be the matrix whose columns are the points. Then, each $\mathbf{x}_j$, $j = 1, \dots, n$, can be expressed as:

$$\mathbf{x}_j = \mathbf{X}\mathbf{c}_j + \mathbf{e}_j, \text{ s.t. } [\mathbf{c}_j]_j = 0, \mathbf{c}_j^T \mathbf{1} = 1, \quad (1)$$

where $\mathbf{c}_j \in \mathbb{R}^n$ is the coefficient vector and $\mathbf{e}_j \in \mathbb{R}^p$ is the representation error. The constraint $[\mathbf{c}_j]_j = 0$ eliminates the trivial solution of expressing a point as a linear combination of itself. Also, the constraint $\mathbf{c}_j^T \mathbf{1} = 1$ allows us to represent data points that lie near a union of affine rather than linear subspaces [5,8,9].

When representing each data point in a low-dimensional subspace in terms of other points in the same subspace, the vector $\mathbf{c}_j$ in Eq. (1) is not unique. However, the main goal is to find a “subspace-preserving” solution such that there are no connections between points from different subspaces. Thus, $[\mathbf{c}_j]_{li} \neq 0$ should indicate that $\mathbf{x}_i$ is in the same subspace as $\mathbf{x}_j$. Given subspace-preserving representations $\mathbf{C} = [\mathbf{c}_1, \dots, \mathbf{c}_n] \in \mathbb{R}^{n \times n}$, a graph with n vertices corresponding to data points is constructed where its affinity matrix is given by the symmetric matrix $\mathbf{W} = |\mathbf{C}| + |\mathbf{C}^T|$. Then, spectral clustering [10] is applied to $\mathbf{W}$ to cluster the data. Note that the subspace clustering problem we consider in this work can be viewed as a particular case of spectral clustering in

E-mail address: farhad_pourkamali@uml.edu (F. Pourkamali-Anaraki).

which the affinity matrix is formed using the self-expressiveness property. More generally, spectral clustering uses manifold structures of data points and a distance metric for constructing graphs that represent such relationships [11–13].

Sparse subspace clustering (SSC) approaches the problem of finding subspace-preserving coefficients by enforcing a sparsity prior on the columns of the matrix $\mathbf{C}$. To do so, a popular technique is centered on solving the following convex optimization program [5,7] (referred to as SSC- ℓ_1 in this paper):

$$\min_{\mathbf{C}} \|\mathbf{C}\|_1 + \frac{\lambda_e}{2} \|\mathbf{X} - \mathbf{XC}\|_F^2 \text{ s.t. } \text{diag}(\mathbf{C}) = \mathbf{0}, \mathbf{C}^T \mathbf{1} = \mathbf{1}, \quad (2)$$

where the ℓ_1 norm promotes the sparsity of $\mathbf{C}$ and $\lambda_e > 0$ is the regularization parameter. Prior work has shown that the solution of (2) is guaranteed to be subspace-preserving under broad conditions on the subspaces as well as under the presence of noise and outliers [14–16]. Although SSC- ℓ_1 is supported by a rich body of theory, the computational complexity associated with solving (2) using the alternating direction method of multipliers (ADMM, cf. [17]) scales cubically with the number of data points. In addition, the process of optimal parameter selection for ADMM requires a significantly increased amount of computational time [18]. In fact, as we will corroborate later, a poor parameter selection for ADMM leads to low accuracy clustering results. Moreover, variants of ADMM with adaptive schemes for updating the solver parameter do not seem to be effective.

Therefore, despite the existence of strong theoretical guarantees, finding subspace-preserving coefficients based on ℓ_1 norm regularization is computationally prohibitive for large-scale data sets [19–23]. One solution to this problem has been to use ℓ_0 instead of ℓ_1 regularization on the columns of $\mathbf{C}$ [24,25]. The resulting model is the following non-convex optimization program (referred to as SSC- ℓ_0 in this paper): for all $j = 1, \dots, n$, solve:

$$\min_{\mathbf{c}_j} \frac{1}{2} \|\mathbf{x}_j - \mathbf{Xc}_j\|_2^2 \text{ s.t. } \|\mathbf{c}_j\|_0 \leq k, [\mathbf{c}_j]_j = 0, \mathbf{c}_j^T \mathbf{1} = 1. \quad (3)$$

If we remove the linear equality constraint $\mathbf{c}_j^T \mathbf{1} = 1$ associated with affine subspaces, then the k -sparse coefficient vector $\mathbf{c}_j$ can be estimated using the orthogonal matching pursuit (OMP) algorithm. However, OMP cannot directly deal with the more general class of affine subspaces, as OMP is a specialized greedy algorithm that enforces the sparsity constraint by only taking k steps and cannot enforce any other kind of constraint. It is also worth pointing out that OMP is only known to solve the problem accurately under certain assumptions that do not hold in the subspace clustering problem. In particular, the data matrix does not satisfy mutual incoherence or restricted isometry properties under the union of subspaces model. The work of [26] presents a theoretical analysis of the sparse subspace clustering problem using ℓ_0 norm regularization for the noiseless case. The work [27] proposes to use the elastic net regularizer (mixture of two norms) to address the scalability issue.

In this paper, we present two first-order methods that can efficiently solve SSC- ℓ_1 and SSC- ℓ_0 optimization problems for the more general case of affine subspaces. Specifically, motivated by theoretical guarantees and empirical success of SSC- ℓ_1 , an efficient proximal gradient method is proposed that requires $\mathcal{O}(n^2)$ time and $\mathcal{O}(n^2)$ memory to find the representation matrix $\mathbf{C}$ for a fixed $p < n$. Another noticeable advantage of the introduced method over ADMM is the lack of additional parameter tuning for a given λ_e . In the case of SSC- ℓ_0 , the main advantage of our proposed solver, compared to other sparse approximation techniques such as OMP, is the ability to handle the more general case of affine subspaces. Recent work [28] showed that the affine constraint might be discarded when the ambient dimension p is large enough compared to the sum of subspace dimensions. However, this assumption

Table 1

Summary of complexity of algorithms discussed, showing the leading order terms assuming $p < n$ where $\mathbf{X} \in \mathbb{R}^{p \times n}$ and nnz is the number of non-zero entries in $\mathbf{X}$, and k is the sparsity in (3).

[T = # iterations]			computation	memory
	linear	ADMM [5]	$n^3 + Tn^3$	n^2
		ADMM (Remark 1)	$pn^2 + Tpn^2$	n^2
SSC- ℓ_1 (Eq. (2))	affine	Proposed	Tpn^2	n^2
		ADMM [5]	$n^3 + Tn^3$	n^2
		ADMM (Remark 1)	$pn^2 + Tpn^2$	n^2
		Proposed	$T(p + \log n)n^2$	n^2
SSC- ℓ_0 (Eq. (3))	linear	OMP [24]	$k(\text{nnz} \cdot n + pkn)$	$\text{nnz} + kn$
	affine	Proposed	$T(\text{nnz} \cdot n + kn)$	$\text{nnz} + kn$
		Proposed	$T(\text{nnz} \cdot n + kn)$	$\text{nnz} + kn$

tion is not realistic in many large-scale problems consisting of multiple subspaces. Our proposed solvers perform SSC efficiently on large data sets regardless of their ambient dimensions.

There are two prior works on using proximal gradient methods in the context of subspace clustering. For example, the authors in [29] used a proximal gradient method for the low-rank subspace clustering problem, where nuclear norm regularization enforces the coefficient matrix $\mathbf{C}$ to be low-rank instead of being sparse. Although using the nuclear norm simplifies the problem, such a regularizer may lead to performance degradation [4], and theoretical guarantees are very limited. Another work used a proximal gradient method to solve the SSC- ℓ_0 problem without the affine constraint [30]. Therefore, this work advances the previous research in this direction by presenting efficient solvers for the sparse subspace clustering problem even with the more general case of affine subspaces.

Additionally, we present an efficient implementation of ADMM for SSC- ℓ_1 using the matrix-inversion lemma. The improved implementation in Remark 1 reduces the computational cost of ADMM for SSC- ℓ_1 [5] from $\mathcal{O}(n^3)$ down to $\mathcal{O}(n^2)$. Such observations have been made for ADMM in general before, but not for SSC in particular, and many popular codes for SSC via ADMM do not use the efficient implementation. A summary of complexity of our proposed solvers is presented in Table 1.

The rest of the paper is organized as follows. In Section 2, we provide a brief review of two main existing solvers for SSC: ADMM and OMP. Section 3 introduces the proposed proximal gradient framework along with detailed instructions on finding proximal operators for the case of affine subspaces. In Section 4, we present various numerical experiments to compare our methods with the existing solvers in terms of computational savings, robustness to solver parameters, and superior performance. Concluding remarks and future research directions are given in Section 5.

Notation Lower-case and upper-case bold letters represent column vectors and matrices, respectively. For a vector $\mathbf{c} \in \mathbb{R}^n$ and $q \geq 1$, let $\|\mathbf{c}\|_q = (\sum_{i=1}^n |\mathbf{c}_i|^q)^{1/q}$ denote the ℓ_q norm, where $[\mathbf{c}]_i$ is the i -th element of $\mathbf{c}$. Also, $\|\mathbf{c}\|_0$ represents the ℓ_0 pseudo-norm which counts the number of non-zero entries in $\mathbf{c}$. Let $\|\mathbf{C}\| = \max_{\mathbf{x}: \|\mathbf{x}\|_2=1} \mathbf{x}^T \mathbf{C} \mathbf{x}$ stand for the spectral norm and let $\|\mathbf{C}\|_F = \sqrt{\sum_{i,j} |\mathbf{C}_{ij}|^2}$ represent the Frobenius norm with the (i, j) -th entry denoted by $[\mathbf{C}]_{ij}$. We use the standard matrix norm $\|\mathbf{C}\|_1 = \sum_{ij} |[\mathbf{C}]_{ij}|$. Finally, $\text{diag}(\mathbf{C})$ returns a column vector of the main diagonal elements of $\mathbf{C}$ and $\mathbf{1}$ denotes the all-ones vector of matching dimensions.

2. Review of Sparse Subspace Clustering

The SSC- ℓ_1 optimization problem can be solved using generic convex solvers such as interior point methods (IPM). However,

even an IPM that is customized to take advantage of problem structure would still require $\mathcal{O}(n^3)$ flops per iteration, and generally 15 to 30 iterations. To reduce the computational cost, Elhamifar and Vidal [5] proposed to use the alternating direction method of multipliers (ADMM). Here, we briefly explain the procedure to solve SSC- ℓ_1 via ADMM to compare with our proposed method in the next section. In our experiments, we also compare our solvers with a variant of ADMM known as Adaptive ADMM (AADMM) [31], which adaptively tunes a penalty parameter to achieve fast convergence.

Let us first introduce an auxiliary matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ and consider the following program whose solution coincides with the solution of the original program:

$$\begin{aligned} \min_{\mathbf{C}, \mathbf{A}} \|\mathbf{C}\|_1 + \frac{\lambda_e}{2} \|\mathbf{X} - \mathbf{X}\mathbf{A}\|_F^2 \\ \text{s.t. } \mathbf{A}^T \mathbf{1} = \mathbf{1}, \quad \mathbf{A} = \mathbf{C} - \text{diag}(\mathbf{C}). \end{aligned} \quad (4)$$

With an abuse of notation in this discussion, $\text{diag}(\mathbf{C})$ also denotes the matrix formed by zeroing all but the diagonal entries of $\mathbf{C}$. Next, the augmented Lagrangian with $\rho > 0$ is formed,

$$\begin{aligned} \mathcal{L}(\mathbf{C}, \mathbf{A}, \boldsymbol{\delta}, \boldsymbol{\Delta}) = \|\mathbf{C}\|_1 + \frac{\lambda_e}{2} \|\mathbf{X} - \mathbf{X}\mathbf{A}\|_F^2 + \frac{\rho}{2} h(\mathbf{C}, \mathbf{A}) \dots \\ \dots + \boldsymbol{\delta}^T (\mathbf{A}^T \mathbf{1} - \mathbf{1}) + \text{trace}(\boldsymbol{\Delta}^T (\mathbf{C} - \text{diag}(\mathbf{C}))) \quad (5) \\ h(\mathbf{C}, \mathbf{A}) \stackrel{\text{def}}{=} \|\mathbf{A}^T \mathbf{1} - \mathbf{1}\|_2^2 + \|\mathbf{A} - (\mathbf{C} - \text{diag}(\mathbf{C}))\|_F^2 \end{aligned}$$

where the Lagrange multipliers are $\boldsymbol{\delta} \in \mathbb{R}^n$ and a matrix $\boldsymbol{\Delta} \in \mathbb{R}^{n \times n}$.

In the i -th iteration of ADMM, the two matrices $\mathbf{A}$ and $\mathbf{C}$ are updated sequentially (à la Gauss-Seidel) by minimizing the Lagrangian with respect to the primal variables. Specifically, $\mathbf{A}^{(i+1)} = \arg \min_{\mathbf{A}} \mathcal{L}(\mathbf{C}^{(i)}, \mathbf{A}, \boldsymbol{\delta}^{(i)}, \boldsymbol{\Delta}^{(i)})$ which can be found by solving the normal equations

$$(\lambda_e \mathbf{X}^T \mathbf{X} + \rho \mathbf{I} + \rho \mathbf{1} \mathbf{1}^T) \mathbf{A}^{(i+1)} = \lambda_e \mathbf{X}^T \mathbf{X} + \rho (\mathbf{1} \mathbf{1}^T + \mathbf{C}^{(i)}) \dots \dots - \mathbf{1} \boldsymbol{\delta}^{(i)T} - \boldsymbol{\Delta}^{(i)}, \quad (6)$$

and $\mathbf{C}^{(i+1)} = \arg \min_{\mathbf{C}} \mathcal{L}(\mathbf{C}, \mathbf{A}^{(i+1)}, \boldsymbol{\delta}^{(i)}, \boldsymbol{\Delta}^{(i)})$ which can be solved as $\mathbf{C}^{(i+1)} = \mathbf{J} - \text{diag}(\mathbf{J})$, where $\mathbf{J} = \text{prox}_{\rho^{-1} \|\cdot\|_1}(\mathbf{A}^{(i+1)} + \rho^{-1} \boldsymbol{\Delta}^{(i)})$ and $\text{prox}_{\eta \|\cdot\|_1}$ applies to each element of the matrix and is defined as $\text{prox}_{\eta \|\cdot\|_1}(v) = \text{sign}(v) \cdot [|v| - \eta]_+$, with $[\tau]_+ \stackrel{\text{def}}{=} \max\{0, \tau\}$, cf. Eq. (13). At the same iteration, $\boldsymbol{\delta}$ and $\boldsymbol{\Delta}$ are updated by a gradient descent step on the dual function: $\boldsymbol{\Delta}^{(i+1)} = \boldsymbol{\Delta}^{(i)} + \rho(\mathbf{A}^{(i+1)} - \mathbf{C}^{(i+1)})$ and $\boldsymbol{\delta}^{(i+1)} = \boldsymbol{\delta}^{(i)} + \rho(\mathbf{A}^{(i+1)T} \mathbf{1} - \mathbf{1})$.

The ADMM solver for SSC- ℓ_1 incurs complexity $\mathcal{O}(n^3 + n^2 p)$ to form $\mathbf{X}^T \mathbf{X}$ and compute the matrix inversion for updating $\mathbf{A}$ in Eq. (6). If it is possible to store the resulting $n \times n$ matrix, one can apply that to the right-hand side of Eq. (6), which incurs complexity $\mathcal{O}(n^3)$ per iteration. Since the overall complexity of ADMM scales cubically with the number of data points n , finding subspace-preserving coefficients based on ℓ_1 norm regularization is computationally prohibitive for large data sets. Hence, there is a need for SSC- ℓ_1 solvers that are computationally efficient.

Remark 1. The implementation of ADMM in [5] has $\mathcal{O}(n^3)$ upfront complexity cost and also $\mathcal{O}(n^3)$ complexity per iteration¹ (for both linear and affine subspace clustering). However, by using the matrix inversion lemma (aka Sherman-Morrison-Woodbury identity), one can reduce the up-front cost to $\mathcal{O}(pn^2 + p^3)$ and the per-iteration cost to $\mathcal{O}(pn^2)$. Our numerical experiments use code from [5] with this modification. Specifically, consider a simplified version of (6) as $(\mathbf{X}^T \mathbf{X} + \rho \mathbf{I}) \mathbf{A}^{(i+1)} = \tilde{\mathbf{C}}$ where $\tilde{\mathbf{C}}$ represents

the right-hand side of (6) and $\mathbf{X}$ has absorbed $\sqrt{\lambda_e}$ and appended the row $\sqrt{\rho} \mathbf{1}^T$ (to account for $\rho \mathbf{1} \mathbf{1}^T$). To initialize, compute $\mathbf{M} = (\mathbf{I}_{p+1} + \rho^{-1} \mathbf{X} \mathbf{X}^T)^{-1}$ (directly or implicitly via a Cholesky factorization) which costs $\mathcal{O}(p^2 n)$ for $\mathbf{X} \mathbf{X}^T$ and $\mathcal{O}(p^3)$ for the inversion/factorization, then use the matrix inversion lemma

$$(\mathbf{X}^T \mathbf{X} + \rho \mathbf{I})^{-1} = \rho^{-1} \mathbf{I} - \rho^{-2} \mathbf{X}^T \mathbf{M} \mathbf{X},$$

and never explicitly form this matrix but rather apply it to $\tilde{\mathbf{C}}$ in $\mathcal{O}(pn^2 + p^2 n)$ time to get

$$\mathbf{A}^{(i+1)} = \rho^{-1} \tilde{\mathbf{C}} - \rho^{-2} \mathbf{X}^T (\mathbf{M} \tilde{\mathbf{C}}).$$

A further disadvantage of ADMM is that tuning the parameter ρ that was introduced in Eq. (5) substantially increases the computational complexity of the ADMM solver. In the implementation of SSC- ℓ_1 solver, the regularization parameter λ_e and the parameter ρ for ADMM are controlled by a parameter α [5, Prop. 1], where $\lambda_e = \alpha/\mu$ for some $\alpha > 1$, $\rho = \alpha$, and

$$\mu \stackrel{\text{def}}{=} \min_i \max_{j \neq i} |\mathbf{x}_i^T \mathbf{x}_j| \quad (7)$$

depends on the data set. In Section 4, we show that the choice of ρ can greatly impact the performance of SSC, and that $\rho = \alpha$ is not a good choice for some data sets. Furthermore, adaptive techniques for updating the parameter ρ do not address this issue.

An alternative method to reduce the memory and computational costs of SSC- ℓ_1 is based on using ℓ_0 norm regularization on the columns of the coefficient matrix $\mathbf{C}$ [24]. Let k be a pre-defined parameter that is proportional to the intrinsic dimensions of subspaces; in practice, it is a parameter that must be estimated. For each point $\mathbf{x}_j$ in the data set, a k -sparse coefficient vector $\mathbf{c}_j \in \mathbb{R}^n$ is obtained by solving the non-convex optimization problem in Eq. (3). Without the linear equality constraint for affine subspaces, the orthogonal matching pursuit (OMP) algorithm can be used to approximately solve this problem. To do so, the j -th column of the data matrix $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_n]$ should be removed and one column of the reduced matrix is selected at a time until k columns are chosen. A simple implementation of OMP requires storing $\mathbf{X}$ and $\mathcal{O}(kn)$ additional storage (for the nonzero entries of $\mathbf{C}$), and incurs complexity $\mathcal{O}(\text{nnz} \cdot k + k^2 p)$ per column j , where $\text{nnz} \leq np$ is the number of non-zero entries in $\mathbf{X}$. Thus, the overall complexity of solving SSC- ℓ_0 via OMP is a quadratic function of n when the sparsity parameter k is small enough compared to n .

3. The Proposed Methods

Section 3.1 reviews the generic proximal gradient descent framework, and then in §3.2 we show how the SSC- ℓ_1 and SSC- ℓ_0 problems can be solved within that framework. The methods consist of a gradient step, which is straightforward and the same for all the variants, and a proximal step. The nature of the proximal step depends on which variant of the problem we solve, and details on all four variants are in §3.3. Because we are able to fit the problems in an existing framework, we can apply standard convergence results, as discussed in §3.4.

3.1. Proximal Gradient Descent Framework

Our methods to solve SSC- ℓ_1 and SSC- ℓ_0 derive from the proximal gradient framework, which we briefly explain. For background on the convex proximal gradient algorithm see [32] or the book [33]; for background on the non-convex version, see [34]. The generic framework is:

$$\min_{\mathbf{y}} f(\mathbf{y}) + g(\mathbf{y}) \quad (8)$$

where f and g are both proper and lower semi-continuous (lsc) extended valued functions, and f has full domain and a Lipschitz

¹ <http://vision.jhu.edu/code/>.

continuous gradient with Lipschitz constant L , and $\mathbf{y}$ is in a finite-dimensional Euclidean space. The function g can be an indicator function $\delta_{\mathcal{Y}}$ of a closed non-empty set $\mathcal{Y}$ meaning that $g(\mathbf{y}) = 0$ if $\mathbf{y} \in \mathcal{Y}$ and $+\infty$ otherwise.

Taking $g \equiv 0$ for the moment, observe that the basic gradient descent iteration $\mathbf{y}^{t+1} = \mathbf{y}^t - \frac{1}{L} \nabla f(\mathbf{y}^t)$ can be equivalently written as:

$$\mathbf{y}^{t+1} = \arg \min_{\mathbf{y}} \underbrace{f(\mathbf{y}^t) + \nabla f(\mathbf{y}^t)^T (\mathbf{y} - \mathbf{y}^t) + \frac{L}{2} \|\mathbf{y} - \mathbf{y}^t\|_2^2}_{Q_f(\mathbf{y}; \mathbf{y}^t)}$$

and that due to the smoothness assumption on f , $Q_f(\mathbf{y}; \mathbf{y}^t) \geq f(\mathbf{y}) \forall \mathbf{y}$ (cf., e.g., [35]), so gradient descent can be viewed as minimizing a majorizing function.

Now allowing a general g , it immediately follows that $Q_f(\mathbf{y}; \mathbf{y}^t) + g(\mathbf{y}) \geq f(\mathbf{y}) + g(\mathbf{y})$, $\forall \mathbf{y}$, and this motivates the update:

$$\mathbf{y}^{t+1} \in \arg \min_{\mathbf{y}} Q_f(\mathbf{y}; \mathbf{y}^t) + g(\mathbf{y}). \quad (9)$$

For any $\gamma > 0$, define the proximity operator (or “prox” for short) to be:

$$\text{prox}_{\gamma g}(\bar{\mathbf{y}}) \in \arg \min_{\mathbf{y}} \gamma \cdot g(\mathbf{y}) + \frac{1}{2} \|\mathbf{y} - \bar{\mathbf{y}}\|_2^2$$

The minimizer may not be unique if g is not convex, in which case the prox is defined as any minimizer. The prox is a natural extension of the Euclidean projection onto a closed nonempty set $\mathcal{Y}$, and indeed if g is the indicator function of $\mathcal{Y}$ then the proximity operator is just the projection onto $\mathcal{Y}$.

By completing the square, the update (9) can be cast as:

$$\mathbf{y}^{t+1} = \text{prox}_{L^{-1}g}(\mathbf{y}^t - L^{-1} \nabla f(\mathbf{y}^t)) \quad (10)$$

which defines the generic proximal gradient algorithm.

3.2. Algorithms for SSC- ℓ_1 and SSC- ℓ_0

The proximal gradient framework applies to SSC- ℓ_1 by identifying f as: $f(\mathbf{C}) = \frac{\lambda}{2} \|\mathbf{X} - \mathbf{X}\mathbf{C}\|_F^2$, $\mathcal{Y}_0 = \{\mathbf{C} \mid \text{diag}(\mathbf{C}) = 0\}$, $\mathcal{Y}_1 = \{\mathbf{C} \mid \mathbf{C}^T \mathbf{1} = \mathbf{1}\}$, and

$$g(\mathbf{C}) = \|\mathbf{C}\|_1 + \delta_{\mathcal{Y}_0}(\mathbf{C}) + \delta_{\mathcal{Y}_1}(\mathbf{C}) = \sum_{j=1}^n g_j(\mathbf{c}_j). \quad (11)$$

Both f and g are separable in the columns $\mathbf{c}_j$ of $\mathbf{C}$ in the sense that $g(\mathbf{C}) = \sum_{j=1}^n g_j(\mathbf{c}_j)$, and likewise for f .

Likewise, the framework applies to SSC- ℓ_0 using the same f , and modifying g to be:

$$g(\mathbf{C}) = \delta_{\mathcal{Y}_k}(\mathbf{C}) + \delta_{\mathcal{Y}_0}(\mathbf{C}) + \delta_{\mathcal{Y}_1}(\mathbf{C}) = \sum_{j=1}^n g_j(\mathbf{c}_j) \quad (12)$$

where $\mathcal{Y}_k = \{\mathbf{C} \mid \mathbf{C} = [\mathbf{c}_1, \dots, \mathbf{c}_n], \|\mathbf{c}_j\|_0 \leq k \forall j = 1, \dots, n\}$. This g is still separable in the columns of $\mathbf{C}$.

The generic proximal gradient algorithm to solve both problems is presented in Algorithm 1. We present a few standard convergence results about the proximal gradient descent algorithm in Section 3.4.

3.3. Proximity Operators for Each Case

We consider the computation of line 6 in Algorithm 1 in detail, for four cases of the g operator that arise from: (1) SSC- ℓ_1 without the $\mathbf{c}_j^T \mathbf{1} = 1$ constraint; (2) SSC- ℓ_1 with the $\mathbf{c}_j^T \mathbf{1} = 1$ constraint; (3) SSC- ℓ_0 without the $\mathbf{c}_j^T \mathbf{1} = 1$ constraint; (4) SSC- ℓ_0 with the $\mathbf{c}_j^T \mathbf{1} = 1$ constraint.

Remark 2. All projections involve the constraint $\mathcal{Y}_0 = \{\mathbf{C} \mid \text{diag}(\mathbf{C}) = 0\}$. For a given column $\mathbf{c}_j$, this can be enforced by

Algorithm 1 Prox. Gradient Descent for SSC- ℓ_1 and SSC- ℓ_0

Parameter: ϵ ▷ Stopping tolerance
Parameter: $\mathbf{C}^0$ ▷ Initialization
Require: $L = \lambda_e \|\mathbf{X}\|^2$ ▷ Lipschitz constant of gradient
1: $t \leftarrow 0$ ▷ Iteration counter
2: $\gamma \leftarrow L^{-1}$ (convex) or $.99L^{-1}$ (non-convex) ▷ Stepsize
3: **repeat**
4: $\tilde{\mathbf{C}} \leftarrow \mathbf{C}^t - \gamma \lambda_e \mathbf{X}^T (\mathbf{X} \mathbf{C}^t - \mathbf{X})$ ▷ Gradient step on f
5: **for** $j = 1, \dots, n$ **do**
6: $\mathbf{c}_j^{t+1} \leftarrow \text{prox}_{\gamma g_j}(\tilde{\mathbf{c}}_j)$ ▷ g as in (11) or (12)
7: $t \leftarrow t + 1$
8: **until** $\|\mathbf{C}^t - \mathbf{C}^{t+1}\|_F \leq \epsilon$

setting the appropriate entry $[\mathbf{c}_j]_j = 0$, and working with the $n - 1$ dimensional versions of the other constraints on the remaining indices. Hence, the dimensions of the columns are really $n - 1$. In this section, for simplicity of exposition, we assume each column $\mathbf{c}_j$ has already had the appropriate entry removed, and we denote its size with n rather than $n - 1$.

Remark 3. In all four cases for g , we can separate $g(\mathbf{C}) = \sum_{j=1}^n g_j(\mathbf{c}_j)$ over the columns. The proximity operator can be computed for each g_j separately and then combined (cf. [36, Prop. 24.11]), hence we only discuss the proximity operator for a single column $\mathbf{c}_j$, and denote this by $\mathbf{c}$ rather than $\mathbf{c}_j$ to unclutter notation. Specifically, with $\mathbf{C} = [\mathbf{c}_1, \dots, \mathbf{c}_n]$, then $\text{prox}_{\gamma g}(\mathbf{C}) = [\text{prox}_{\gamma g_1}(\mathbf{c}_1), \dots, \text{prox}_{\gamma g_n}(\mathbf{c}_n)]$.

3.3.1. ℓ_1 proximity operator

First, consider the SSC problem assuming all subspaces are true subspaces, and therefore pass through $\mathbf{0}$. In this case, there is no $\mathbf{C}^T \mathbf{1} = \mathbf{1}$ constraint, and the proximity operator is:

$$\text{prox}_{\gamma \|\cdot\|_1}(\mathbf{d}) \stackrel{\text{def}}{=} \arg \min_{\mathbf{c}} \frac{1}{2} \|\mathbf{c} - \mathbf{d}\|_2^2 + \gamma \|\mathbf{c}\|_1 \quad (13)$$

and it is well-known that the solution is component-wise soft-thresholding (also known as “shrinkage”):

$$[\text{prox}_{\gamma \|\cdot\|_1}(\mathbf{d})]_i = \text{sign}(d_i) \cdot [|d_i| - \gamma]_+ \quad (14)$$

where $[\tau]_+ \stackrel{\text{def}}{=} \max\{0, \tau\}$.

3.3.2. ℓ_1 proximity operator with affine constraint

Now, consider the full SSC problem with affine spaces. The proximity operator computation is to solve:

$$\arg \min_{\mathbf{c}} \frac{1}{2} \|\mathbf{c} - \mathbf{d}\|_2^2 + \gamma \|\mathbf{c}\|_1 \text{ s.t. } \mathbf{C}^T \mathbf{1} = \mathbf{1}. \quad (15)$$

Eq. (15) is a strongly convex minimization problem with a unique solution, but it is not separable, and the solution is not-obvious, yet it clearly has specific structure. Efficient algorithms for it have been proposed going back at least to the 1980s [37], and it has been rediscovered many times (e.g., [38–40]). In some incarnations, it is known as the “continuous knapsack” problem. It is related to other ℓ_1 problems, such as projection onto the ℓ_1 ball ([41], and rediscovered and/or improved in [42–46]) and trust-region or exact line search variants, as well as quasi-Newton variants [47]. Most formulations are reducible to each other, accounting for some of the duplications in the literature. The approaches fall into a few categories: reduction to low-dimensional linear or quadratic programs, fast median searches, or one-dimensional root-finding via bisection. We present below a derivation using a one-dimensional root-finding approach that has complexity $\mathcal{O}(n \log n)$. We suspect that fast median-finding ideas might enable a $\mathcal{O}(n)$ algorithm but

do not pursue this since theoretical $\mathcal{O}(n)$ median-finding algorithms are in practice slower than efficient implementations of $\mathcal{O}(n \log n)$ sorting algorithms until n is extremely large.

Proposition 4. *The problem (15) can be solved exactly in $\mathcal{O}(n \log n)$ flops.*

By “exact” solution, we mean there is no optimization error, though there is possibly roundoff error due to floating point computation unless exact arithmetic is used. As mentioned above, related results have appeared in the literature so we do not claim novelty, but the algorithms are not well known, so we give the proof below since it also explains the algorithm.

Proof. The standard Lagrangian for (15) is $\mathcal{L}(\mathbf{c}; \beta) = \frac{1}{2} \|\mathbf{c} - \mathbf{d}\|_2^2 + \gamma \|\mathbf{c}\|_1 + \beta(\mathbf{c}^T \mathbf{1} - 1)$ where the dual variable β is a scalar. Since the problem is convex and only has equality constraints, Slater's conditions are satisfied and the following two KKT conditions are necessary and sufficient for a point $\mathbf{c}$ to be optimal:

$$0 \in \partial \mathcal{L}(\mathbf{c}; \beta), \text{ i.e., } 0 \in \mathbf{c} - \mathbf{d} + \gamma \partial \|\mathbf{c}\|_1 + \beta \mathbf{1} \quad (16)$$

$$\mathbf{c}^T \mathbf{1} = 1 \quad (17)$$

where $\partial \mathcal{L}$ is the subdifferential. Observe that Fermat's rule for convex functions, namely that $y \in \arg \min F(x)$ if and only if $0 \in \partial F(y)$, applied to the objective in (13) implies that $\mathbf{c} = \text{prox}_{\gamma \|\cdot\|_1}(\mathbf{d})$ if and only if $0 \in \partial(\frac{1}{2} \|\mathbf{c} - \mathbf{d}\|_2^2 + \gamma \|\mathbf{c}\|_1) = \mathbf{c} - \mathbf{d} + \gamma \partial \|\mathbf{c}\|_1$ where the equality is true since both functions have full domain. Thus the condition (16) is equivalent to $\mathbf{c} = \text{prox}_{\gamma \|\cdot\|_1}(\mathbf{d} - \beta \mathbf{1})$. Substituting this into (17) gives that $\mathbf{1}^T \text{prox}_{\gamma \|\cdot\|_1}(\mathbf{d} - \beta \mathbf{1}) = 1$ is a necessary and sufficient condition in terms of only the scalar β . We can rewrite this condition as:

$$0 = f(\beta) \stackrel{\text{def}}{=} \sum_{i=1}^n \text{sign}(d_i - \beta) \cdot \lfloor |d_i - \beta| - \gamma \rfloor_+ - 1. \quad (18)$$

This is a one-dimensional, piecewise linear root-finding problem in β , and the linear regions occur between the break-points where $|d_i - \beta| = \gamma$, i.e., $\beta = d_i \pm \gamma$. In the linear regions, solving for β is just solving a 1D linear equation, so the only difficulty is finding the correct linear region. Each term in the sum of f is monotonically decreasing in β , therefore the function f is monotonically decreasing in β . There are $2n$ break-points of the form $\beta = d_i \pm \gamma$, so our algorithm sorts these $2n$ break-points, with cost $\mathcal{O}(n \log n)$ (e.g., using merge sort), and then does a bisection search on the regions defined by the break-points, with $\mathcal{O}(\log n)$ steps, and linear complexity per step. See Algorithm 2. $\square$

Algorithm 2 Algorithm to solve Eq. (15)

```

1:  $\text{prox}_{\gamma \|\cdot\|_1}$  defined as prox from Eq. (14)
2: Convention:  $b_0 = -\infty, b_{2n+1} = +\infty$ 
3: function PROX( $\mathbf{d} \in \mathbb{R}^n, \gamma \in \mathbb{R}^+$ )
4:    $i_{\min} = 0, i_{\max} = 2n + 1$ 
5:    $\mathbf{b} = \text{sort}(\{\mathbf{d} - \gamma\} \cup \{\mathbf{d} + \gamma\})$   $\triangleright b_1 \geq b_2 \dots \geq b_{2n}$ 
6:   while  $i_{\max} - i_{\min} > 1$  do
7:      $j \leftarrow \lfloor (i_{\min} + i_{\max})/2 \rfloor$   $\triangleright$  Round to an integer
8:      $\mathbf{c} \leftarrow \text{prox}_{\gamma \|\cdot\|_1}(\mathbf{d} - b_j \mathbf{1})$ 
9:     if  $\mathbf{c}^T \mathbf{1} > 1$  then  $i_{\max} \leftarrow j$ 
10:    else  $i_{\min} \leftarrow j$ 
11:   Choose any  $\beta \in (b_{i_{\min}}, b_{i_{\max}})$ 
12:    $\mathbf{c} \leftarrow \text{prox}_{\gamma \|\cdot\|_1}(\mathbf{d} - \beta \mathbf{1})$ 
13:    $\mathcal{S}^* \leftarrow \text{supp}(\mathbf{c})$   $\triangleright$  Find the support
14:    $\beta^* \leftarrow \frac{-1}{|\mathcal{S}^*|} (1 - \sum_{i \in \mathcal{S}^*} d_i - \gamma \text{sign}(c_i))$ 
15:    $\mathbf{c} \leftarrow \text{prox}_{\gamma \|\cdot\|_1}(\mathbf{d} - \beta^* \mathbf{1})$ 
16:   return  $\mathbf{c}$ 
```

3.3.3. ℓ_0 projection

Again, we first discuss the problem assuming all subspaces are true subspaces and not affine spaces, so there is no $\mathbf{c}^T \mathbf{1} = 1$ constraint. The relevant proximity operator reduces to the following Euclidean projection:

$$\arg \min_{\mathbf{c}} \frac{1}{2} \|\mathbf{c} - \mathbf{d}\|_2^2 \text{ s.t. } \|\mathbf{c}\|_0 \leq k. \quad (19)$$

While this is a non-convex problem, due to its simple structure, it is easy to solve. For example, one can sort the absolute value of all n terms ($|d_i|$) and then choose the top k largest (which may not be unique if there are duplicate values of $|d_i|$), at cost $\mathcal{O}(n \log(n))$. Alternatively, it may be faster to take the largest entry in absolute value, and repeat k times $\mathcal{O}(nk)$. Specialized implementations based on heapsort can also return the answer in $\mathcal{O}(n \log(k))$ [48].

3.3.4. ℓ_0 projection with affine constraint

Adding in the affine constraint $\mathbf{c}^T \mathbf{1} = 1$, the relevant proximity operator is:

$$\arg \min_{\mathbf{c}} \frac{1}{2} \|\mathbf{c} - \mathbf{d}\|_2^2 \text{ s.t. } \|\mathbf{c}\|_0 \leq k, \mathbf{c}^T \mathbf{1} = 1. \quad (20)$$

It is not obvious that there is an efficient algorithm to solve this non-convex problem, but in fact due to its special structure, there is a specific greedy algorithm, known as the “greedy selector and hyperplane projector” (GSHP), which has been shown to exactly solve (20) and take time complexity $\mathcal{O}(n \cdot k)$ [49]; pseudo-code is shown in Algorithm 3. For a set $\mathcal{S}$ and vector $\mathbf{d} \in \mathbb{R}^n$, the

Algorithm 3 GSHP to solve Eq. (20) [49]

```

1:  $\mathcal{P}(\mathbf{d}) \stackrel{\text{def}}{=} \mathbf{d} - \frac{1}{n}(\mathbf{d}^T \mathbf{1} - 1) \mathbf{1}$   $\triangleright$  Proj. onto  $\{\mathbf{c} \mid \mathbf{c}^T \mathbf{1} = 1\}$ 
2: function GSHP( $\mathbf{d} \in \mathbb{R}^n, k \in \mathbb{N}^+$ )
3:    $\ell = 1, \mathcal{S} = j, j \in \arg \max_i |d_i|$   $\triangleright$  Initialize
4:   repeat  $\ell \leftarrow \ell + 1, \mathcal{S} \leftarrow \mathcal{S} \cup \{j\}$ , where
5:      $j \in \arg \max_{i \in \mathcal{S}^c} \left| d_i - \frac{\sum_{j \in \mathcal{S}} d_j - 1}{\ell - 1} \right|$   $\triangleright$  Grow
6:   until  $\ell = k$ , set  $\mathcal{S}^* \leftarrow \mathcal{S}$ 
7:    $\mathbf{c}_{|\mathcal{S}^*} = \mathcal{P}(\mathbf{d}_{|\mathcal{S}^*})$ ,  $\mathbf{c}_{|\mathcal{S}^c} = 0$   $\triangleright$  Final projection
8:   return  $\mathbf{c}$ 
```

notation $\mathbf{d}_{|\mathcal{S}}$ refers to the vector created by restricting $\mathbf{d}$ to the entries in $\mathcal{S}$, and $\mathcal{S}^c = \{1, 2, \dots, n\} \setminus \mathcal{S}$.

3.4. Convergence Results

In this section, we provide convergence results for the proposed SSC- ℓ_1 and SSC- ℓ_0 solvers.

3.4.1. SSC- ℓ_1

Theorem 5. *Let $(\mathbf{C}^t)_{t \in \mathbb{N}}$ be the sequence of points generated by Algorithm 1, let $\mathbf{C}^*$ be any optimal solution to SSC- ℓ_1 (2), and let $F(\cdot)$ denote the objective function in (2). Then for any $t \in \mathbb{N}$, $\mathbf{C}^t$ is feasible for (2) and*

$$F(\mathbf{C}^t) - F(\mathbf{C}^*) \leq \frac{L}{2} \frac{1}{t} \|\mathbf{C}^0 - \mathbf{C}^*\|_F^2.$$

Furthermore, $(\mathbf{C}^t)_{t \in \mathbb{N}}$ converges to an optimal point.

This is a well-known result. See, for example, the textbook [33, Thm. 10.21] for the rate, and the textbook [36, Cor. 28.9] for the sequence convergence. We present this result for simplicity, but note that “Nesterov accelerated” variants of proximal gradient descent (also known as “FISTA”) have a very similar per-step computational cost and improve the convergence rate to $\mathcal{O}(1/t^2)$ instead of $\mathcal{O}(1/t)$. There are also variants that allow for variable step-sizes, rather than just $1/L$. If $\gamma = 1/L$ is used, L is not needed

to high accuracy, so it can be computed with a few iterations of the power method, or exactly in $\mathcal{O}(p^2n)$ time. In practice, for the SSC- ℓ_1 problem, we use the Nesterov accelerated variants provided in the TFOCS package [50] which also incorporates a line search for the stepsize.

Remark 6. Note that Algorithm 1 solves for all columns of $\mathbf{C}$ at once, requiring $\mathcal{O}(n^2)$ memory. If memory is a concern, the problem can be solved a single column at a time due to its separable nature, requiring only $\mathcal{O}(pn)$ memory (to store $\mathbf{X}$) for (2) or $\mathcal{O}(\text{nnz}(\mathbf{X}) + kn)$ for (3), and not changing the asymptotic computational cost. This should not be done unless necessary, since computing with all blocks at once allows for efficient level-3 BLAS operations which are optimized to reduce communication cost and greatly improve practical performance. In practice, a few columns at a time can be solved.

Remark 7. The convergence results for both convex and non-convex cases do not change whether one includes the $\mathbf{c}_j^T \mathbf{1} = 1$ constraint or not. Dropping the constraint only simplifies the computation of the proximity operator, as discussed in Section 3.3.

3.4.2. SSC- ℓ_0

This is a non-convex problem, so one would not expect a *a priori* global convergence guarantees. In particular, we cannot guarantee that for an arbitrary initialization, the sequence converges to a global minimizer, but the following theorem does show that the algorithm is at least *consistent* with the optimization problem. The theorem is actually unusually strong for non-convex problems, and relies on the results by Attouch et al. [34] on the Kurdyka-Łojasiewicz inequality. More traditional theory would only have been able to guarantee that, at best, any cluster point of the sequence is a stationary point of the optimization problem.

Theorem 8. Let $(\mathbf{C}^t)_{t \in \mathbb{N}}$ be the sequence of points generated by Algorithm 1. If the sequence $(\mathbf{C}^t)_{t \in \mathbb{N}}$ is bounded, then it converges to a stationary point $\bar{\mathbf{C}}$ of SSC- ℓ_0 (3), i.e., $\bar{\mathbf{C}}$ is feasible and

$$-\nabla f(\bar{\mathbf{C}}) \in N(\bar{\mathbf{C}})$$

where N is the normal cone of the set $\mathcal{Y} = \mathcal{Y}_k \cap \mathcal{Y}_0 \cap \mathcal{Y}_1$, i.e.,

$$\nabla f(\bar{\mathbf{C}})^T (\mathbf{C} - \bar{\mathbf{C}}) \geq 0 \quad \forall \mathbf{C} \in \mathcal{Y}$$

The proof follows from using $\epsilon = .01/L$ in [34, Thm. 5.3] and observing that f and g are semi-algebraic and all the sets $\mathcal{Y}$ are closed.

Remark 9. As in the convex case, we can solve for each column $\mathbf{c}_j$ one-by-one. If $\mathbf{X}$ is sparse, the memory savings are potentially very large, since for a single column, we only need a temporary memory of $\mathcal{O}(n)$ and $\mathcal{O}(\text{nnz}(\mathbf{X}) + k)$ for the variables.

4. Numerical Experiments

We compare the performance of our proposed methods from Section 3 with ADMM and OMP. Most of our experiments focus on the affine case, since there are fewer algorithms available to solve it, and some authors argue it is more powerful since it is a more general model. We implemented the proximal operators in MATLAB and C++, and then incorporated these into the generic proximal minimization framework of the software package TFOCS [50]. We write TFOCS in the legend of figures to mean our implementation of Algorithm 1 for either the SSC- ℓ_1 or SSC- ℓ_0 case.

The full clustering algorithm is shown in Alg. 4, which consists of running one of the four optimization solver described in the previous section, followed by spectral clustering. When n is large, we use Matlab's Krylov-subspace based solver `eigs` to compute the eigenvalue decomposition. The final K-means clustering is done

Algorithm 4 End-to-end algorithm including spectral clustering

Parameter: K ▷ Estimated number of clusters
Parameter: λ_e ▷ For SSC- ℓ_1 only
 1: $\mathbf{C} \leftarrow \mathbb{R}^{n \times n}$ ← Algorithm 1 ▷ SSC- ℓ_1 or SSC- ℓ_0 , affine or not
 2: $\mathbf{W} \leftarrow |\mathbf{C}| + |\mathbf{C}|^T$ ▷ Often very sparse
 3: $[\mathbf{D}]_{ii} = \sum_{j=1}^n [\mathbf{W}]_{ij}$, $\mathbf{D} \in \mathbb{R}^{n \times n}$ diagonal
 4: $\mathbf{V} \leftarrow \text{eig}(\mathbf{D}^{-\frac{1}{2}} \mathbf{W} \mathbf{D}^{-\frac{1}{2}}, K)$, $\mathbf{V} = [\mathbf{v}_1^T; \dots; \mathbf{v}_K^T] \in \mathbb{R}^{n \times K}$ ▷ Only need eigenvectors corresponding to K -largest eigenvalues
 5: $\mathbf{v}_i \leftarrow \mathbf{v}_i / \|\mathbf{v}_i\|_2$ for $i = 1, \dots, n$
 6: Cluster via `kmeans` ($\{\mathbf{v}_i\}_{i=1}^n, K$)

via Matlab's `kmeans` which uses Lloyd's algorithm and takes the best of 20 random initializations.

As explained in Remark 1, the implementation of ADMM in [5] has $\mathcal{O}(n^3)$ complexity. However, we provided a more efficient implementation using the matrix-inversion lemma that has reduced the per-iteration cost to $\mathcal{O}(n^2)$. The regularization parameter λ_e for SSC- ℓ_1 is controlled by some parameter $\alpha > 1$ as $\lambda_e = \alpha/\mu$, where μ is a quantity that depends on the given data set (cf. Eq. (7)). In all experiments with ℓ_1 norm regularization, TFOCS and ADMM share the same regularization parameter λ_e . However, ADMM requires the additional parameter ρ to be tuned. The default value for ρ in the implementation provided by the authors is $\rho = \alpha$. In agreement with the findings of many other papers, we observe that the choice of ρ can greatly impact the performance of ADMM. Thus, one should ideally tune the parameter ρ for each experiment, which increases the overall computational cost of ADMM for SSC- ℓ_1 . We also show that our proposed solver outperforms a new variant of ADMM, called "Adaptive-ADMM" (AADMM), which adaptively tunes the parameter ρ for fast convergence [31].

Throughout this section, we use real and synthetic data sets. The first real data set is the Extended Yale B data set [51]. This data set contains frontal face images of 38 individuals under 64 different illumination conditions. These images are downsampled to 48×42 pixels, thus the data points lie in $\mathbb{R}^p$ with $p = 2,016$.

The second real data set is the CoverType data set² which contains $n = 581,012$ observations of $p = 54$ features, where each observation is the forest cover type (lodgepole pine, cottonwood/willow, etc.) of a 30m by 30m section of Earth, and examples of features are elevation, aspect, etc. There are $K = 7$ possible forest cover types.

The synthetic data is based on the following statistical model that considers n data points in $\mathbb{R}^p$ drawn from a union of K affine subspaces $\{S_i\}_{i=1}^K$:

$$\mathbf{x}_i = \mathbf{U}^{(i)} \mathbf{z}_i + \boldsymbol{\mu}^{(i)} + \mathbf{v}_i, \quad \forall \mathbf{x}_i \in S_i, \quad (21)$$

where the columns of $\mathbf{U}^{(i)} \in \mathbb{R}^{p \times r_i}$ form an orthonormal basis of S_i , $\mathbf{z}_i \in \mathbb{R}^{r_i}$ is the low-dimensional representation of $\mathbf{x}_i$ with respect to $\mathbf{U}^{(i)}$, $\boldsymbol{\mu}^{(i)} \in \mathbb{R}^p$ is the intercept of S_i , and $\mathbf{v}_i \in \mathbb{R}^p$ is the noise vector. Thus, we can control the number of subspaces, their dimensions, intersections, and the amount of noise in order to gain insights on the performance of the aforementioned solvers. We test various SSC- ℓ_1 solvers on up to $n = 15,000$ data points.

4.1. SSC- ℓ_1 on the Extended Yale B Data Set

In the first experiment, we compare the performance of the proposed TFOCS solver with ADMM and adaptive ADMM (AADMM) for solving SSC- ℓ_1 on the Extended Yale B data set when the parameter α is set to be 1.1. For ADMM, we consider the recommended value of ρ , $\rho = \alpha$, as well as the alternatives

² <http://archive.ics.uci.edu/ml/datasets/Covertype>.

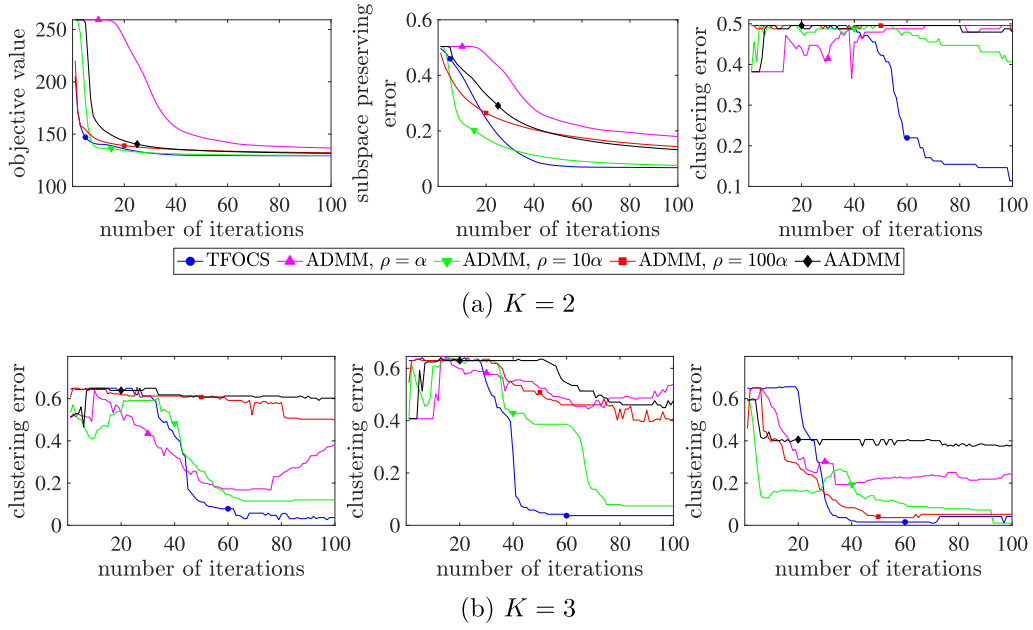


Fig. 1. SSC- ℓ_1 on the Extended Yale B data set for (a) $K = 2$ and (b) $K = 3$ clusters. For each case, three metrics are used from left to right: value of the objective function, subspace preserving error, and clustering error. The legends in (a) and (b) are the same.

$\rho = 10\alpha, 100\alpha$. Three metrics are used to demonstrate the performance of these solvers over 100 iterations (we report all three metrics because in our experience they are not necessarily correlated with each other): (1) value of the objective function in Eq. (2); (2) subspace preserving error [26], which is the average fraction of ℓ_1 norm of each representation vector in the data set that comes from other subspaces; and (3) clustering error, which is the fraction of misclustered points after applying spectral clustering to $\mathbf{W}$ [52].

Since we want to compare the three solvers in each iteration and the solution of ADMM is not necessarily feasible (e.g., $\mathbf{c}_j^T \mathbf{1}$ may not be 1), we find the closest feasible solution by first removing the j -th element of $\mathbf{c}_j$ to get $\tilde{\mathbf{c}}_j \in \mathbb{R}^{n-1}$. Then, we solve the following:

$$\mathbf{c}_j^* = \arg \min_{\mathbf{c} \in \mathbb{R}^{n-1}} \frac{1}{2} \|\mathbf{c} - \tilde{\mathbf{c}}_j\|_2^2 \text{ s.t. } \mathbf{c}^T \mathbf{1} = 1. \quad (22)$$

It is straightforward to show that the solution of this problem is $\mathbf{c}_j^* = \tilde{\mathbf{c}}_j - \nu \mathbf{1}$, where the scalar is $\nu = (\tilde{\mathbf{c}}_j^T \mathbf{1} - 1)/(n-1)$. The feasible representation vectors are only used for evaluating the three metrics in each iteration and they are not used for next iterations of ADMM.

In Figure 1a, the three metrics are plotted when $K = 2$ clusters are selected uniformly at random from 38 individuals. It is observed that the performance of ADMM depends heavily on the choice of the penalty parameter ρ . Interestingly, the choice of $\rho = \alpha$ is found to result in the worst performance. However, our proposed solver outperforms or has similar performance compared to ADMM without having to tune additional parameters. Moreover, the recently proposed AADMM which adaptively tunes ρ seems to be effective, but does not compete with our proposed solver. We also report clustering errors in Figure 1b for three independent trials when $K = 3$ clusters are randomly selected. Similar results are obtained when the feasibility projection in Eq. (22) is not performed.

We note that the clustering error we found for $K = 2$ is higher than found in the original sparse subspace clustering (SSC) paper [5]. The reason is that like many other papers, we used a subset of K individuals from the entire face data set. Thus the clustering error depends heavily on which subset is chosen (in general, clus-

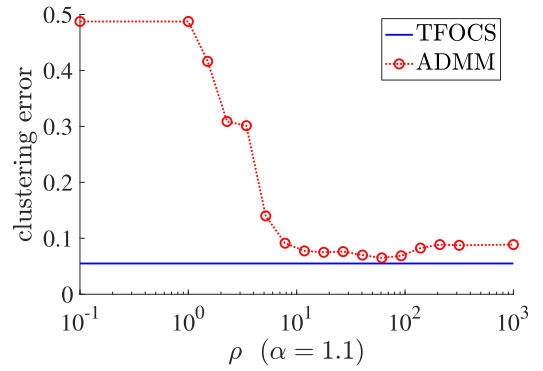


Fig. 2. Clustering error as a function of ρ .

tering error depends on the orientation of subspaces). To illustrate our point, we used the original SSC code (and the values that were originally recommended) and we observed that for $K = 2$, the clustering error can be as high as 0.5 depending on the selected subset.

4.2. Varying Values of ρ in ADMM

We note that larger values of ρ for ADMM does not necessarily improve performance. To demonstrate this point, we compare the performance of TFOCS and ADMM solvers for SSC- ℓ_1 when the maximum number of iterations is set to be 250. We set $\alpha = 1.1$ and consider various values of ρ from 0.1 to 1,000 (approximately from 0.09α to 909α) for a subset of $K = 2$ clusters with 400 data points chosen uniformly at random from each cluster of the Cover-Type data set. The clustering error results are shown in Figure 2. As we see, the performance of ADMM is close to our solver for a small interval of ρ , which again emphasizes the importance of tuning ρ for any given data set.

4.3. SSC- ℓ_1 on Synthetic Data Sets

We consider the statistical model described in Eq. (21). This model allows us to control the number of subspaces K , their di-

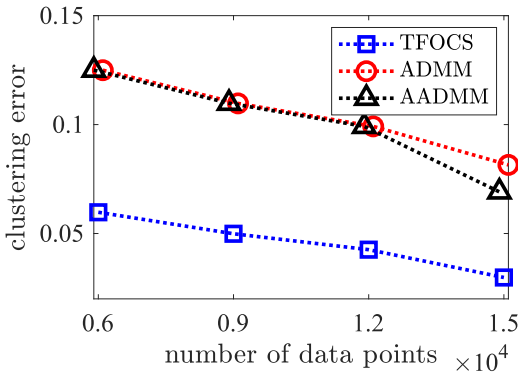


Fig. 3. Clustering error of $\text{SSC-}\ell_1$ on synthetic data.

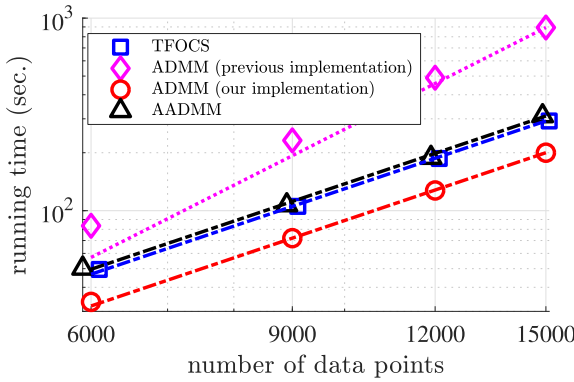


Fig. 4. Running time (logarithmic scale) of $\text{SSC-}\ell_1$ on synthetic data for varying n .

mensions r_l , orientations, and the amount of noise. We set parameters $p = 256$, $K = 10$, $r_l = 3$, and $\mu^{(l)} = \mathbf{0}$ for all $l \in \{1, \dots, 10\}$. The columns of the orthonormal matrices $\mathbf{U}^{(l)} \in \mathbb{R}^{p \times r_l}$ are drawn uniformly at random from a set of p orthonormal random vectors in $\mathbb{R}^p$. Each coefficient vector $\mathbf{z}_i \in \mathbb{R}^{r_l}$ is drawn i.i.d. from the standard normal distribution. The noise vectors $\mathbf{v}_i \in \mathbb{R}^p$, $i = 1, \dots, n$, are drawn i.i.d. according to $\mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$, where we set $\sigma = 0.1$. We sample 600 to 1,500 data points per subspace, which leads to the total number of data points from $n = 6,000$ to $n = 15,000$.

The clustering error results averaged over 10 independent trials are presented in Figure 3 for fixed $\alpha = 30$, $\rho = 10\alpha$ for ADMM, and the maximum number of iterations is set to be 50. We observe that our solver consistently outperforms both ADMM and AADMM. For example, when $n = 15,000$, the average errors are 0.03, 0.08, and 0.07 for our solver, ADMM, and AADMM, respectively.

To demonstrate the efficiency of the $\text{SSC-}\ell_1$ solvers, the average running times in seconds are plotted in Figure 4. These results verify our claim that both the proposed TFOCS solver and our implementation of ADMM scales quadratically with the number of data points n . However, the implementation of ADMM in [5] has complexity $\mathcal{O}(n^3)$. Although, the new implementation of ADMM is slightly faster than our proposed TFOCS solver by a constant factor, its performance depends crucially on the parameter ρ . Therefore, one can argue that the effective cost of ADMM is higher compared to the proposed solver in this work as our solver does not require any additional parameter tuning.

4.4. $\text{SSC-}\ell_1$ and $\text{SSC-}\ell_0$ on the CoverType Data Set

We test the subspace ℓ_0 model on the CoverType data set with $n = 581,012$, $p = 54$ and $K = 7$. Due to the size of n , the variable

$\mathbf{C} \in \mathbb{R}^{n \times n}$ can never be formed except as a sparse matrix. It was reported in [27] that for this data set, even using just two iterations of the solvers, that OMP took 783 minutes, and two $\text{SSC-}\ell_1$ methods (one based on the original SSC ADMM algorithm, without using Remark 1) either did not finish the two iterations within 7 days, or used more than 16 GB of memory.

The results of running our models on this data (after normalizing the features to z-scores), and taking 2 steps as in [27], is presented in Table 2. We do not include results for $\text{SSC-}\ell_1$ with affine constraints, as with default parameters this model does not lead to a sparse $\mathbf{C}$, and hence there are memory issues. The $\text{SSC-}\ell_0$ models are guaranteed to give a sparse output, and we test the non-affine variant on the full CoverType data, in addition to give results for randomly subsampling $n = 10^5$ data points (about 1 in 5). For $\text{SSC-}\ell_1$, α was set to 0.1 to encourage sparsity.

The time for the full $n = 5.81 \cdot 10^5$ data is $30.4 \times$ slower than for the $n = 10^5$ data. Based on the $\mathcal{O}(n^2)$ complexity, one would expect it to be $33.7 \times$ slower, which is in good agreement (to within factors such as cost of memory movement, CPU throttling and efficiencies of scale). It is notably faster than all the methods discussed in [27]. The experiment was run on a 6-core 2.6 GHz laptop with 16 GB of RAM.

4.5. $\text{SSC-}\ell_0$ on Synthetic Data Sets

In this experiment, we again use a synthetic data set generated based on the statistical model described in Eq. (21). The parameters are $p = 64$, $K = 3$, $r_l = 10$, $n = 600$, and $\mu^{(l)} = \mathbf{0}$ for all $l \in \{1, 2, 3\}$. We choose $\mu^{(l)} = \mathbf{0}$, i.e., subspace not affine space clustering, since we do not know of other algorithms that can handle the affine space case. The maximum number of iterations is set to be 100. Similar to one of the experiments in [6], we consider the case that every pair of subspaces intersects in at least 5 dimensions. To do so, the orthonormal bases are given by $\mathbf{U}^{(l)} = [\mathbf{U} \ \tilde{\mathbf{U}}^{(l)}] \in \mathbb{R}^{p \times 10}$, where matrices $\mathbf{U}$ and $\tilde{\mathbf{U}}^{(l)}$, $l = 1, 2, 3$, are chosen uniformly at random among all orthonormal matrices of size $p \times 5$. The noise term $\mathbf{v}_i \in \mathbb{R}^p$ is distributed according to $\mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$, where the noise level σ is varied from 0 to 1.0. Therefore, this synthetic data set allows us to study the impact of noise as well as the choice of k on clustering performance using our TFOCS and OMP methods for solving $\text{SSC-}\ell_0$.

The clustering error results, showing average and standard deviation over 20 independent trials, for two choices of sparsity $k = 10$ and $k = 20$, are plotted in Figure 5, where k is the sparsity parameter in Eq. (3). As expected, larger values of the noise level σ result in lower accuracy clustering results. However, we see that our TFOCS solver consistently outperforms OMP for both $k = 10$ and $k = 20$, and the effect is more pronounced for $k = 20$. Since each subspace in this example is 10-dimensional, it is worth pointing that the proposed TFOCS solver is less sensitive to the choice of sparsity k than OMP.

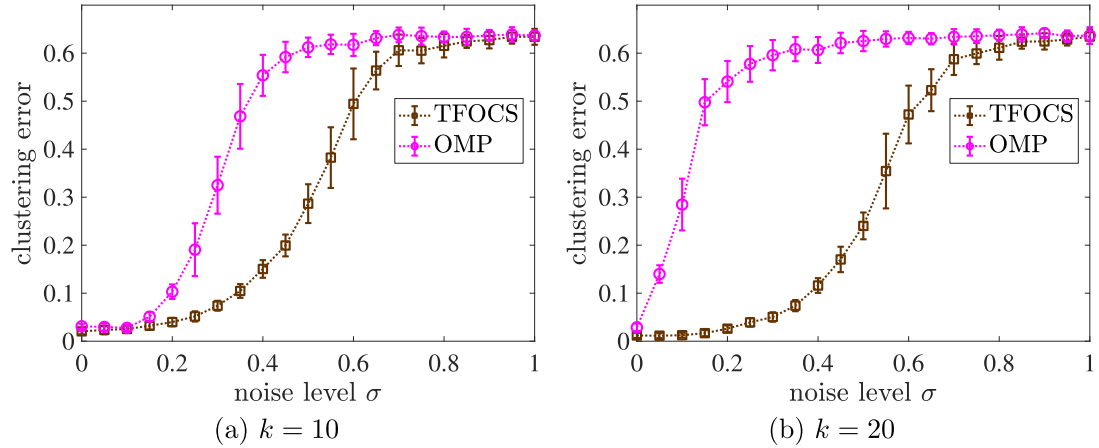
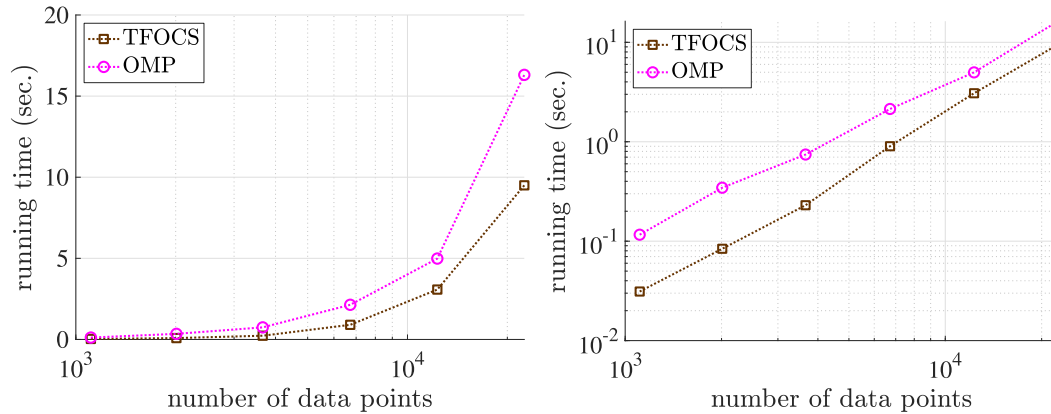
To compare the efficiency of our proximal gradient solver for the $\text{SSC-}\ell_0$ problem with OMP, the running times required to achieve a certain level of accuracy for various number of data points from $n = 600$ to $n = 22,500$ are plotted in Figure 6 (other parameters such as the dimension of subspaces and the ambient dimension are unchanged). To be more specific, we run OMP for 10 iterations and then run our TFOCS solver to match the clustering error produced by OMP, and report the corresponding running time. In this experiment, it is observed that TFOCS is faster to reach OMP's accuracy.

To summarize, our proximal $\text{SSC-}\ell_0$ algorithm is significantly more accurate than OMP when the noise is high and/or k is over-specified. Furthermore, our solver is the only known algorithm to solve the affine space variant of $\text{SSC-}\ell_0$.

Table 2

Results on CoverType data, $p = 54$. Times are in minutes. Time_{SC} is the time for the spectral clustering step. Avg sparsity is the avg number of nonzero entries per column of $\mathbf{C}$. There are 7 types of data, so accuracy for random guessing is 14%.

	Algorithm	Time	Time_{SC}	Avg sparsity	Accuracy
$n = 10^5$	SSC- ℓ_1 , not affine	3.7	0.04	35	42.24%
	SSC- ℓ_0 , not affine	4.6	0.15	7	41.41%
	SSC- ℓ_0 , affine	4.6	0.15	7	43.66%
$n = 5.81 \cdot 10^5$	SSC- ℓ_0 , not affine	139.9	2.7	7	43.41%

**Fig. 5.** Clustering error of SSC- ℓ_0 on synthetic data for varying σ (noise level).**Fig. 6.** Running time of SSC- ℓ_0 on synthetic data for varying number of data points n . Left: time on a linear scale. Right: same data, but time on a logarithmic scale.

5. Conclusion

We proposed two efficient proximal gradient methods for finding sparse representation vectors of data points that lie in or close to a union of affine subspaces. We also presented a detailed performance and complexity analysis of our proximal solvers. In addition, an efficient implementation of the popular ADMM technique for solving ℓ_1 norm regularized SSC optimization problems is provided. Overall, the two proposed proximal solvers and our implementation of ADMM substantially reduces the computational cost of solving large-scale SSC optimization problems. A key advantage of our proximal solver for SSC- ℓ_1 is the lack of additional parameter tuning, which makes it much more efficient than ADMM (if one does cross-validation to find the correct parameter). Experimentally, ADMM does appear to be sensitive to its additional parameter ρ . Finally, our proposed proximal solver for SSC- ℓ_0 has the ability to directly deal with the more general case of affine subspaces, and experimentally it appears to be less sensitive to the choice of sparsity parameter compared to the existing algorithm that uses OMP.

As a final note, it is worth pointing out our proposed solvers can be adopted in a recent line of work, e.g., [21], that uses exemplars or representative points to further achieve scalability to large data sets.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Supplementary material

Supplementary material associated with this article can be found, in the online version, at doi:[10.1016/j.sigpro.2020.107548](https://doi.org/10.1016/j.sigpro.2020.107548).

CRediT authorship contribution statement

Farhad Pourkamali-Anaraki: Conceptualization, Methodology, Investigation, Writing - original draft. **James Folberth:** Software,

Writing - review & editing. **Stephen Becker:** Methodology, Software, Writing - review & editing, Supervision.

References

- [1] F. Pourkamali-Anaraki, S. Becker, Preconditioned data sparsification for big data with applications to PCA and K-means, *IEEE Transactions on Information Theory* 63 (5) (2017) 2954–2974.
- [2] O. Bachem, M. Lucic, A. Krause, Scalable K-means clustering via lightweight coresets, in: *International Conference on Knowledge Discovery and Data Mining (KDD)*, 2018, Pp. 1119–1127.
- [3] M. Soltanolkotabi, E. Candès, A geometric analysis of subspace clustering with outliers, *The Annals of Statistics* 40 (4) (2012) 2195–2238.
- [4] R. Vidal, Y. Ma, S. Sastry, *Sparse and low-rank methods*, in: *Generalized Principal Component Analysis*, Springer, 2016, Pp. 291–346.
- [5] E. Elhamifar, R. Vidal, Sparse subspace clustering: Algorithm, theory, and applications, *IEEE Transactions on Pattern Analysis and Machine Intelligence* 35 (11) (2013) 2765–2781.
- [6] R. Heckel, M. Tschannen, H. Bölcskei, Dimensionality-reduced subspace clustering, *Information and Inference: A Journal of the IMA* 6 (3) (2017) 246–283.
- [7] E. Elhamifar, R. Vidal, Sparse subspace clustering, in: *IEEE Conference on Computer Vision and Pattern Recognition*, 2009, Pp. 2790–2797.
- [8] V. Patel, H.V. Nguyen, R. Vidal, Latent space sparse subspace clustering, in: *IEEE International Conference on Computer Vision*, 2013, Pp. 225–232.
- [9] C. Li, C. You, R. Vidal, On geometric analysis of affine sparse subspace clustering, *IEEE Journal of Selected Topics in Signal Processing* 12 (6) (2018) 1520–1533.
- [10] U.V. Luxburg, A tutorial on spectral clustering, *Statistics and computing* 17 (4) (2007) 395–416.
- [11] G. Schiebinger, M. Wainwright, B. Yu, The geometry of kernelized spectral clustering, *The Annals of Statistics* 43 (2) (2015) 819–846.
- [12] E. Arias-Castro, T.L. Gouic, Unconstrained and curvature-constrained shortest-path distances and their approximation, *Discrete & Computational Geometry* 62 (1) (2019) 1–28.
- [13] N. Tremblay, A. Loukas, Approximating spectral clustering via sampling: a review, in: *Sampling Techniques for Supervised or Unsupervised Tasks*, Springer, 2020, Pp. 129–183.
- [14] M. Soltanolkotabi, E. Elhamifar, E. Candès, Robust subspace clustering, *The Annals of Statistics* 42 (2) (2014) 669–699.
- [15] C. You, R. Vidal, Geometric conditions for subspace-sparse recovery, in: *International Conference on Machine Learning*, 2015, Pp. 1585–1593.
- [16] M. Tsakiris, R. Vidal, Theoretical analysis of sparse subspace clustering with missing entries, in: *International Conference on Machine Learning*, 2018, Pp. 4975–4984.
- [17] S. Boyd, N. Parikh, E. Chu, B. Peleato, J. Eckstein, Distributed optimization and statistical learning via the alternating direction method of multipliers, *Foundations and Trends in Machine Learning* 3 (1) (2011) 1–122.
- [18] Y. Xu, M. Liu, Q. Lin, T. Yang, ADMM without a fixed penalty parameter: Faster convergence with new adaptive penalization, in: *Advances in Neural Information Processing Systems*, 2017, Pp. 1267–1277.
- [19] A. Adler, M. Elad, Y. Hel-Or, Linear-time subspace clustering via bipartite graph modeling, *IEEE Transactions on Neural Networks and Learning Systems* 26 (10) (2015) 2234–2246.
- [20] P. Traganitis, G. Giannakis, Sketched subspace clustering, *IEEE Transactions on Signal Processing* 66 (7) (2017) 1663–1675.
- [21] C. You, C. Li, D. Robinson, R. Vidal, Scalable exemplar-based subspace clustering on class-imbalanced data, in: *European Conference on Computer Vision (ECCV)*, 2018, Pp. 67–83.
- [22] M. Abdolali, N. Gillis, M. Rahmati, Scalable and robust sparse subspace clustering using randomized clustering and multilayer graphs, *Signal Processing* 163 (2019) 166–180.
- [23] F. Pourkamali-Anaraki, Large-scale sparse subspace clustering using landmarks, in: *International Workshop on Machine Learning for Signal Processing (MLSP)*, 2019, Pp. 1–6.
- [24] E. Dyer, A. Sankaranarayanan, R. Baraniuk, Greedy feature selection for subspace clustering, *Journal of Machine Learning Research* 14 (1) (2013) 2487–2517.
- [25] Y. Chen, G. Li, Y. Gu, Active orthogonal matching pursuit for sparse subspace clustering, *IEEE Signal Processing Letters* 25 (2) (2018) 164–168.
- [26] C. You, D. Robinson, R. Vidal, 2016a. pages->3918–3927, unknown->Scalable sparse subspace clustering by orthogonal matching pursuit, in: *IEEE Conference on Computer Vision and Pattern Recognition*, pp.
- [27] C. You, C. Li, D. Robinson, R. Vidal, Oracle based active set algorithm for scalable elastic net subspace clustering, in: *IEEE Conference on Computer Vision and Pattern Recognition*, 2016b, Pp. 3928–3937.
- [28] C. You, C. Li, D. Robinson, R. Vidal, Is an affine constraint needed for affine subspace clustering? in: *IEEE International Conference on Computer Vision*, 2019, Pp. 9915–9924.
- [29] H. Jiang, D. Robinson, R. Vidal, C. You, A nonconvex formulation for low rank subspace clustering: algorithms and convergence analysis, *Computational Optimization and Applications* 70 (2) (2018) 395–418.
- [30] Y. Yang, J. Feng, N. Jovic, J. Yang, T. Huang, ℓ_0 -sparse subspace clustering, in: *European conference on computer vision*, 2016, Pp. 731–747.
- [31] Z. Xu, M. Figueiredo, T. Goldstein, Adaptive ADMM with spectral penalty parameter selection, in: *International Conference on Artificial Intelligence and Statistics*, 2017, Pp. 718–727.
- [32] P.L. Combettes, V.R. Wajs, Signal recovery by proximal forward-backward splitting, *SIAM Multiscale Model. Simul.* 4 (4) (2005) 1168–1200.
- [33] A. Beck, First-order methods in optimization, in: *MOS-SIAM Series on Optimization*, 2017.
- [34] H. Attouch, J. Bolte, B. Svaiter, Convergence of descent methods for semi-algebraic and tame problems: proximal algorithms, forward-backward splitting, and regularized Gauss-Seidel methods, *Mathematical Programming* (2011) 1–39.
- [35] Y. Nesterov, Introductory lectures on convex optimization: A basic course, in: *Vol. 87 of Applied Optimization*, Kluwer, Boston, 2004.
- [36] H.H. Bauschke, P.L. Combettes, *Convex analysis and monotone operator theory in hilbert spaces*, in: 2nd Edition, Springer-Verlag, New York, 2017.
- [37] J.-P. Dussault, J.A. Ferland, B. Lemaire, Convex quadratic programming with one constraint and bounded variables, *Mathematical Programming* 36 (1986) 90–104.
- [38] S. Stefanov, Polynomial algorithms for projecting a point onto a region defined by a linear constraint and box constraints in R^n , *Journal of Applied Mathematics* 2004 (5) (2004) 409–431.
- [39] K. Kiwiel, On linear-time algorithms for the continuous quadratic knapsack problem, *Journal of Optimization Theory and Applications* 134 (2007) 549–554.
- [40] L. Bayón, J.M. Grau, M.M. Ruiz, P.M. Suárez, An analytic solution for some separable convex quadratic programming problems with equality and inequality constraints, *Journal of Mathematical inequalities* 4 (3) (2010) 453–465.
- [41] P. Brucker, An $O(n)$ algorithm for quadratic knapsack problems, *Operations Research Letters* 3 (3) (1984) 163–166, doi:10.1016/0167-6377(84)90010-5.
- [42] J. Liu, J. Ye, Efficient Euclidean projections in linear time, in: *Proceedings of the 26th Annual International Conference on Machine Learning*, ACM, 2009, Pp. 657–664.
- [43] J. Duchi, S. Shalev-Schwartz, Y. Singer, T. Chandra, Efficient projections onto the l_1 -ball for learning in high dimensions, *International Conference on Machine Learning* (2008), Pp. 272–279.
- [44] E.V.d. Berg, M. Schmidt, M.P. Friedlander, K. Murphy, Group sparsity via linear time projection, 2008, Tech. Rep. TR-2008-09, Dept. Comp. Sci., U. British Columbia (June).
- [45] P.M. Pardalos, N. Kuvorov, An algorithm for a singly constrained class of quadratic programs subject to upper and lower bounds, *Mathematical Programming* 46 (1990) 321–328.
- [46] N. Maculan, J.R.G.G. de Paula, A linear-time median-finding algorithm for projecting a vector on the simplex of R^n , *Operations research letters* 8 (4) (1989) 219–222.
- [47] S. Becker, J. Fadili, A quasi-Newton proximal splitting method, in: *Neural Information Processing Systems*, 2012, Pp. 2618–2626.
- [48] D.E. Knuth, *The art of computer programming*, Vol. 3, Pearson Education (1997).
- [49] S. Becker, V. Cevher, C. Koch, A. Kyriklidis, Sparse projections onto the simplex, in: *International Conference on Machine Learning*, 2013, Pp. 235–243.
- [50] S. Becker, E.J. Candès, M. Grant, Templates for convex cone problems with applications to sparse signal recovery, *Mathematical programming computation* 3(3).
- [51] A. Georgiades, P. Belhumeur, D. Kriegman, From few to many: Illumination cone models for face recognition under variable lighting and pose, *IEEE Transactions on Pattern Analysis and Machine Intelligence* 23 (6) (2001) 643–660.
- [52] R. Heckel, H. Bölcskei, Robust subspace clustering via thresholding, *IEEE Transactions on Information Theory* 61 (11) (2015) 6320–6342.