

A lifted ℓ_1 framework for sparse recovery

YAGHOUB RAHIMI AND SUNG HA KANG

School of Mathematics, Georgia Institute of Technology, Atlanta, GA 30332, USA

AND

YIFEI LOU*

Department of Mathematics & School of Data Science and Society, The University of North Carolina at Chapel Hill, NC 27599, USA

*Corresponding author: yflou@unc.edu

[Received on 14 March 2023; revised on 10 August 2023; accepted on 18 November 2023]

We introduce a lifted ℓ_1 (LL1) regularization framework for the recovery of sparse signals. The proposed LL1 regularization is a generalization of several popular regularization methods in the field and is motivated by recent advancements in re-weighted ℓ_1 approaches for sparse recovery. Through a comprehensive analysis of the relationships between existing methods, we identify two distinct types of lifting functions that guarantee equivalence to the ℓ_0 minimization problem, which is a key objective in sparse signal recovery. To solve the LL1 regularization problem, we propose an algorithm based on the alternating direction method of multipliers and provide proof of convergence for the unconstrained formulation. Our experiments demonstrate the improved performance of the LL1 regularization compared with state-of-the-art methods, confirming the effectiveness of our proposed framework. In conclusion, the LL1 regularization presents a promising and flexible approach to sparse signal recovery and invites further research in this area.

Keywords: Compressed sensing; sparse recovery; reweighted L1; nonconvex minimization; alternating direction method of multipliers.

1. Introduction

Compressed sensing [14] plays an important role in many applications including signal processing, medical imaging, matrix completion, feature selection and machine learning [18,24,33]. One important assumption in compressed sensing is that a signal of interest is sparse or compressible. This allows for an efficient representation of high-dimensional data only by a few meaningful subsets. Compressed sensing often involves sparse recovery from an under-determined linear system that can be formulated mathematically by minimizing the ℓ_0 ‘pseudo-norm’, i.e.

$$\arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{x}\|_0 \quad \text{s.t.} \quad \mathbf{A}\mathbf{x} = \mathbf{b}, \quad (1.1)$$

where $\mathbf{A} \in \mathbb{R}^{m \times n}$ is called a sensing matrix and $\mathbf{b} \in \mathbb{R}^m$ is a measurement vector. Since the minimization problem (1.1) is NP-hard [34], various regularization functionals are proposed to approximate the ℓ_0 penalty. The ℓ_1 norm is widely used as a convex relaxation of ℓ_0 , which is called LASSO [46] in statistics or basis pursuit [13] in signal processing. Due to its convex nature, the ℓ_1 norm is computationally traceable to optimize with exact recovery guarantees based on restricted isometry property (RIP) and/or null space property (NSP) [9,15,47]. Alternatively, there are non-convex models, i.e. concave with respect

to the positive cone, that outperform the convex ℓ_1 approach in terms of improved accuracy of identifying sparse solutions. For example, ℓ_p ($0 < p < 1$) [12,26,56,57], smoothly clipped absolute deviation (SCAD) [17], minimax concave penalty (MCP) [61], capped ℓ_1 (CL1) [30,43,64] and transformed ℓ_1 (TL1) [31,62,63] are separable and concave penalties. Some non-separable non-convex penalties include sorted ℓ_1 [23], $\ell_1 - \ell_2$ [4,28,29,58,59] and ℓ_1/ℓ_2 [39,51,55]. Properties of sparsity-promoting functions are developed in [42]. In addition, a large volume of literature discusses greedy algorithms [25,33] that minimize the ℓ_0 regularization, such as orthogonal matching pursuit [49], CoSAMP [35], iterative hard thresholding (IHT) [7], conjugate gradient IHT [6] and algebraic pursuits [11].

In this paper, we generalize a type of re-weighted approaches [10,20,53] for sparse recovery based on the fact that a convex function can be smoothly approximated by subtracting a strongly convex function from the objective function [36]. In particular, we propose a *lifted regularization*, which we referred to as *lifted ℓ_1 (LLI)*,

$$F_{g,\alpha}^U(\mathbf{x}) = \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}), \quad (1.2)$$

where we denote $|\mathbf{x}| = [|x_1|, \dots, |x_n|] \in \mathbb{R}^n$ and $\langle \cdot, \cdot \rangle$ is the standard Euclidean inner product. In (1.2), the variable $\mathbf{u}$ plays the role of weights with U as a set of restrictions on $\mathbf{u}$, e.g. $U = [0, \infty)^n$ or $U = [0, 1]^n$. The function $g: \mathbb{R}^n \rightarrow \mathbb{R}$ is decreasing near zero (please refer to Definition 1 for multi-variable decreasing function) to ensure a non-trivial solution of (1.2), and $\alpha > 0$ is a parameter. We can consider a more general function $g(\mathbf{u}; \alpha)$, instead of $\alpha g(\mathbf{u})$ when making a connection of the proposed model to several existing sparse-promoting regularizations in Section 2.2. We focus on two specific types of g functions (see Definition 2) that ensure the existence and uniqueness of the minimum.

To find a sparse signal from a set of linear measurements, we propose the following constrained minimization problem:

$$\min_{\mathbf{x}, \mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) \quad \text{s.t.} \quad \mathbf{A}\mathbf{x} = \mathbf{b}. \quad (1.3)$$

We lift the dimension of a single vector $\mathbf{x} \in \mathbb{R}^n$ in the original ℓ_0 problem (1.1) into a joint optimization over $\mathbf{x}$ and $\mathbf{u}$ in the proposed model (1.3). We can establish the equivalence between the two if the function g and the set U satisfy certain conditions. For the measurements with noise, we consider an unconstrained formulation, i.e.

$$\min_{\mathbf{x}, \mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) + \frac{\gamma}{2} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2^2, \quad (1.4)$$

where $\gamma > 0$ is a balancing parameter.

From an algorithmic point of view, the lifted form enables us to minimize two variables each with fast implementation, and can lead to a better local minimum in a higher dimension than the original formulation [2,27,36,60]. The idea of updating $\mathbf{u}$ while finding the sparsest solution $\mathbf{x}$ is also related to the classic graduated non-convex algorithm [5,41] that constructs a family of functions (1.3) approaching to the true cost function (1.1). To minimize (1.3) and (1.4), we apply the alternating direction method of multipliers (ADMM) [8,19], and conduct convergence analysis of ADMM for solving the unconstrained problem (1.4). We demonstrate in experiments that the proposed approach outperforms the state of the art. Our contributions are threefold,

- We propose a unified model that generalizes many existing regularizations and re-weighted algorithms for sparse recovery problems;

- We establish the equivalence between the proposed model (1.3) and the ℓ_0 model (1.1);
- We present an efficient algorithm that is supported by a convergence analysis.

The rest of the paper is organized as follows. In Section 2, we present details of the proposed model, including its properties with an exact sparse recovery guarantee. Section 3 describes the ADMM implementation and its convergence. Section 4 presents experimental results to demonstrate how this lifted model improves the sparse recovery over the state of the art. Finally, concluding remarks are given in Section 5.

2. The Proposed Lifted L1 Framework

We first introduce definitions and notations that are used throughout the paper. The connection to well-known sparsity-promoting regularizations is presented in Section 2.2. We present useful properties of LL1 in Section 2.3, and establish its equivalence to the original ℓ_0 problem in Section 2.4.

2.1 Definitions and Notations

We mark any vector in bold, specifically $\mathbf{1}$ denotes the all-ones vector and $\mathbf{0}$ denotes the zero vector. For a vector $\mathbf{x} \in \mathbb{R}^n$, we define $|x|_{(i)}$ as the i -th largest component of $|\mathbf{x}|$. We say that a vector $\mathbf{x} \in \mathbb{R}^n$ majorizes $\mathbf{x}' \in \mathbb{R}^n$ with the notation $|\mathbf{x}| \succ |\mathbf{x}'|$, if we have $\sum_{i \leq j} |x|_{(i)} \leq \sum_{i \leq j} |x'|_{(i)} \forall 1 \leq j < n$ and $\sum_{i=1}^n |x_i| = \sum_{i=1}^n |x'_i|$. For two vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, the notation $\mathbf{x} \leq \mathbf{y}$ means each element of $\mathbf{x}$ is smaller than or equal to the corresponding element in $\mathbf{y}$; similarly for $\mathbf{x} \geq \mathbf{y}$, $\mathbf{x} > \mathbf{y}$ and $\mathbf{x} < \mathbf{y}$. The positive cone refers to the set $\mathbb{R}_+^n = \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{x} \geq \mathbf{0}\}$. A rectangular subset of $\mathbb{R}^n$ is a Cartesian product of intervals aligned with the canonical axes. We define two elementwise operators, $\max(\mathbf{x}, \mathbf{y})$ and $\mathbf{x} \odot \mathbf{y}$, both returning a vector form by taking maximum and multiplication, respectively, for every component. The proposed regularization (1.2) can be expressed by the δ -function,

$$F_{g,\alpha}^U(\mathbf{x}) = \min_{\mathbf{u}} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) + \delta_U(\mathbf{u}), \quad \text{where } \delta_U(\mathbf{x}) = \begin{cases} 0 & \text{if } \mathbf{x} \in U \\ +\infty & \text{if } \mathbf{x} \notin U. \end{cases} \quad (2.1)$$

We use (1.2) and (2.1) interchangeably throughout the paper. The constrained LL1 problem (1.3) is equivalent to

$$\min_{\mathbf{x}, \mathbf{u}} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) + \delta_U(\mathbf{u}) + \delta_\Omega(\mathbf{x}), \quad (2.2)$$

where $\Omega = \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{A}\mathbf{x} = \mathbf{b}\}$. We denote $[n]$ as the set $\{1, 2, \dots, n\}$, S_n as the symmetric group of n elements and a permutation $\pi \in S_n$ of a vector $\mathbf{x}$ is defined as $\mathbf{x} \circ \pi = (x_{\pi(1)}, \dots, x_{\pi(n)})$. We summarize the relevant properties of a function as follows.

DEFINITION 1. Let $g : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty, -\infty\}$ be a function, we say that

- The function g is *separable* if there exists a set of functions $\{g_i : \mathbb{R} \rightarrow \mathbb{R}\}$ for $i \in [n]$ s.t.

$$g(\mathbf{x}) = \sum_{i=1}^n g_i(x_i), \quad \forall \mathbf{x} = [x_1, \dots, x_n].$$

- The function g is *strongly convex* with parameter $\mu > 0$ if

$$g(\mathbf{y}) \geq g(\mathbf{x}) + \langle \nabla g(\mathbf{x}), \mathbf{y} - \mathbf{x} \rangle + \frac{\mu}{2} \|\mathbf{y} - \mathbf{x}\|_2^2, \quad \forall \mathbf{x}, \mathbf{y}. \quad (2.3)$$

- The function g is *symmetric* if $g(\mathbf{x}) = g(\mathbf{x} \circ \pi)$, $\forall \mathbf{x} \in \mathbb{R}^n$ and $\forall \pi \in S_n$.
- The function g is *coercive* if $g(\mathbf{x}) \rightarrow +\infty$ as $\|\mathbf{x}\| \rightarrow +\infty$.
- The function g is *decreasing* on U if $g(\mathbf{x}) \leq g(\mathbf{y})$, $\forall \mathbf{x} \geq \mathbf{y}$ and $\mathbf{x}, \mathbf{y} \in U$.

2.2 Connections to Sparsity Promoting Models

Many existing models can be understood as special cases of the proposed LL1 model (1.3). We start with two recent works. One is a joint minimization model [65] between the weights and the vector,

$$\min_{\mathbf{x}, \mathbf{u}} \langle \mathbf{u}, |\mathbf{x}| - \epsilon \rangle \quad \text{s.t.} \quad A\mathbf{x} = \mathbf{b}, \mathbf{u} \in \{0, 1\}^n, \quad (2.4)$$

where $\epsilon > 0$ is a fixed number and $|\mathbf{x}| - \epsilon$ is a vector subtracting ϵ from all components of $|\mathbf{x}|$. With the assumption that the weights are binary, the authors [65] established the equivalence between (1.1) and (2.4) for a small enough ϵ . Another related work is the trimmed LASSO [1,3],

$$\min_{\|\mathbf{y}\|_0 \leq k} \|\mathbf{x} - \mathbf{y}\|_1 = \sum_{i=k+1}^n |x|_{(i)} = \min_{\mathbf{u}} \langle \mathbf{u}, |\mathbf{x}| \rangle + \delta_U(\mathbf{u}),$$

which is equivalent to the LL1 form on the right with $U = \{\mathbf{u} \in \{0, 1\}^n \mid \|\mathbf{u}\|_1 = n - k\}$. In the middle, the sum is over the $n - k$ smallest components of the vector $\mathbf{x}$ for a given sparsity k .

In what follows, we consider a more general form of $g(\mathbf{u}; \alpha)$, as opposed to $\alpha g(\mathbf{u})$, i.e.

$$F_g^U(\mathbf{x}) = \min_{\mathbf{u}} \langle \mathbf{u}, |\mathbf{x}| \rangle + g(\mathbf{u}; \alpha) + \delta_U(\mathbf{u}). \quad (2.5)$$

As a consequence of the *Fenchel–Moreau’s* Theorem, regularizations that are concave on the positive cone are of the form (2.5). In other words, we have the following theorem:

THEOREM 1. Any proper and lower semi-continuous function $J(\mathbf{x})$ that is concave on the positive cone can be lifted by a convex function $g : \mathbb{R}^n \rightarrow \mathbb{R}$ and a set U such that $J(|\mathbf{x}|) = F_g^U(\mathbf{x})$ as in (2.5). Here $g(\mathbf{u}) := \sup_{\mathbf{x} \geq \mathbf{0}} \langle \mathbf{x}, -\mathbf{u} \rangle + J(\mathbf{x})$ and $U = \{\mathbf{u} \in \mathbb{R}^n \mid g(\mathbf{u}) \neq +\infty\}$.

Please refer to Appendix A and Appendix B for the proof of Theorem 1 and detailed computations, respectively. As many sparsity-promoting functions satisfy the assumptions in Theorem 1, we can rewrite $J(\mathbf{x})$ in the form of F_g^U . Using Theorem 1, we relate (2.5) to the following functionals that are widely used to promote sparsity.

- The ℓ_p model [12,26] is defined as $J^{\ell_p}(\mathbf{x}) = \sum_{i=1}^n |x_i|^p$. As $p \rightarrow 0$, ℓ_p^p approaches to ℓ_0 , while it reduces to ℓ_1 for $p = 1$. Much research focuses on $p = 1/2$, due to a closed-form solution [56,57] in the optimization process. By choosing $g(\mathbf{u}; p) = \frac{1-p}{p} \sum_{i=1}^n u_i^{\frac{p}{p-1}}$ and $U = \mathbb{R}_+^n$, we

can express the ℓ_p regularization as

$$\min_{\mathbf{u}} \left\{ \sum_{i=1}^n \left(u_i |x_i| + \frac{1-p}{p} u_i^{\frac{p}{p-1}} \right) \mid u_i \geq 0 \right\}.$$

- (ii) The log-sum penalty [10,38] is given by $J_a^{\log}(\mathbf{x}) = \sum_{i=1}^n \log(|x_i| + a)$, for a small positive parameter a to make the logarithmic function well-defined. The re-weighted ℓ_1 algorithm [10] (IRL1) minimizes $J_a^{\log}(\mathbf{x})$, which is equivalent to (2.5) in that

$$\min_{\mathbf{u}} \left\{ \sum_{i=1}^n (u_i |x_i| + a u_i - \log(u_i)) \mid u_i \geq 0 \right\}.$$

- (iii) Smoothly clipped absolute deviation (SCAD) [17] is defined by

$$J_{a,b}^{\text{SCAD}}(\mathbf{x}) = \sum_{i=1}^n f_{a,b}^{\text{SCAD}}(x_i), \text{ where } f_{a,b}^{\text{SCAD}}(t) = \begin{cases} a|t| & \text{if } |t| \leq a, \\ \frac{2ab|t| - t^2 - a^2}{2(b-1)} & \text{if } a < |t| \leq ab, \\ \frac{(b+1)a^2}{2} & \text{if } |t| > ab. \end{cases} \quad (2.6)$$

This penalty is designed to reduce the bias of the ℓ_1 model. For $g(\mathbf{u}, a, b) = -ab\|\mathbf{u}\|_1 + (b-1)\frac{\|\mathbf{u}\|_2^2}{2}$ and $U = [0, a]^n$, we have $J_{a,b}^{\text{SCAD}}$ is equivalent to

$$\min_{\mathbf{u}} \left\{ \sum_{i=1}^n \left(u_i |x_i| - ab u_i + (b-1) \frac{u_i^2}{2} \right) \mid u_i \in [0, a] \right\}.$$

- (iv) Mini-max concave penalty (MCP) [61] is defined by

$$J_{a,b}^{\text{MCP}}(\mathbf{x}) = \sum_{i=1}^n f_{a,b}^{\text{MCP}}(x_i), \text{ where } f_{a,b}^{\text{MCP}}(t) = \begin{cases} a|t| - \frac{t^2}{2b} & \text{if } |t| \leq ab, \\ \frac{1}{2}ba^2 & \text{if } |t| > ab. \end{cases} \quad (2.7)$$

With the same purpose of reducing bias as SCAD, MCP consists of two piece-wise defined functions, which is simpler than SCAD. For $g(\mathbf{u}, a, b) = -b(a\|\mathbf{u}\|_1 - \frac{\|\mathbf{u}\|_2^2}{2})$ and $U = \mathbb{R}_+^n$, we can rewrite $J_{a,b}^{\text{MCP}}$ as

$$\min_{\mathbf{u}} \left\{ \sum_{i=1}^n \left(u_i |x_i| + b \left(\frac{u_i^2}{2} - a u_i \right) \right) \mid u_i \geq 0 \right\}.$$

- (v) Capped ℓ_1 (CL1) [64] is defined as $J_a^{\text{CL1}}(\mathbf{x}) = \sum_{i=1}^n \min\{|x_i|, a\}$, with a positive parameter $a > 0$. As $a \rightarrow 0$ the function F_a^{CL1}/a approaches to ℓ_0 . The CL1 penalty is unbiased and has

fewer internal parameters than SCAD/MCP. For $g(\mathbf{u}, a) = -a\|\mathbf{u}\|_1$ and $U = [0, 1]^n$, the CL1 regularization can be expressed as

$$\min_{\mathbf{u}} \left\{ \sum_{i=1}^n (u_i |x_i| - au_i) \mid 0 \leq u_i \leq 1 \right\}.$$

It is similar to the model in [65], except for a binary restriction on $\mathbf{u}$ as in (2.4).

- (vi) Transformed ℓ_1 (TL1) [31] is defined as $J_a^{\text{TL1}}(\mathbf{x}) = \sum_{i=1}^n \frac{(a+1)|x_i|}{a+|x_i|}$. It reduces to ℓ_0 for $a = 0$, and converges to ℓ_1 as $a \rightarrow \infty$. The TL1 regularization is Lipschitz continuous, unbiased and equivalent to

$$\min_{\mathbf{u}} \left\{ \sum_{i=1}^n (u_i |x_i| + au_i - 2\sqrt{au_i}) \mid u_i \geq 0 \right\}.$$

- (vii) Error function penalty (ERF) [20] is defined by

$$J_{\sigma}^{\text{ERF}}(\mathbf{x}) = \sum_{i=1}^n f_{\sigma}^{\text{ERF}}(|x_i|), \text{ where } f_{\sigma}^{\text{ERF}}(t) = \int_0^t e^{-\tau^2/\sigma^2} d\tau. \quad (2.8)$$

This model is less biased than ℓ_1 , and gives a good approximation to ℓ_0 for a small value of σ . Let $h(t) = \int_t^1 \sqrt{-\log(\tau)} d\tau$, then for $g(\mathbf{u}, \sigma) = \sigma \sum_{i=1}^n h(u_i)$ and $U = [0, 1]^n$, the ERF function is equivalent to

$$\min_{\mathbf{u}} \left\{ \sum_{i=1}^n (u_i |x_i| + \sigma h(u_i)) \mid u_i \in [0, 1] \right\}.$$

- (viii) The $\ell_1 - \ell_2$ penalty, defined as $J^{\text{L1-L2}}(\mathbf{x}) = \|\mathbf{x}\|_1 - \|\mathbf{x}\|_2$, is an example of non-separable regularization, which has been proven to perform very well when the matrix A is highly coherent [29,59]. The $\ell_1 - \ell_2$ regularization is equivalent to

$$\min_{\mathbf{u}} \left\{ \sum_{i=1}^n u_i |x_i| \mid \|\mathbf{1} - \mathbf{u}\|_2 \leq 1 \right\}.$$

In this case, the corresponding g function is zero, and the set U is a non-rectangular set, thus non-separable.

We summarize these existing regularizations with their corresponding g function and U set in Table 1. All these g functions are separable, and hence we can plot g as a univariate function. As illustrated in Fig. 1, each g function is decreasing on a small interval near the origin, thus motivating the conditions on the function g presented in Definition 2 as well as Theorems 5 and 6 for exact recovery guarantees.

TABLE 1 Summary of various regularizations and their corresponding g function and U set

Model	Regularization	Function g	Set U
ℓ_p	$\sum_{i=1}^n x_i ^p$	$\frac{1-p}{p} \sum_{i=1}^n u_i^{\frac{p}{p-1}}$	$\mathbb{R}_+^n$
log-sum	$\sum_{i=1}^n \log(x_i + a)$	$a\ \mathbf{u}\ _1 - \sum_{i=1}^n \log(u_i)$	$\mathbb{R}_+^n$
SCAD	(2.6)	$-ab\ \mathbf{u}\ _1 + (b-1)\frac{\ \mathbf{u}\ _2^2}{2}$	$[0, a]^n$
MCP	(2.7)	$-b(a\ \mathbf{u}\ _1 - \frac{\ \mathbf{u}\ _2^2}{2})$	$\mathbb{R}_+^n$
CL1	$\sum_{i=1}^n \min\{ x_i , a\}$	$-a\ \mathbf{u}\ _1$	$[0, 1]^n$
TL1	$\sum_{i=1}^n \frac{(a+1) x_i }{a+ x_i }$	$a\ \mathbf{u}\ _1 - 2 \sum_i (\sqrt{au_i})$	$\mathbb{R}_+^n$
ERF	(2.8)	$\sigma \sum_i h(u_i)$	$[0, 1]^n$
$\ell_1 - \ell_2$	$\ \mathbf{x}\ _1 - \ \mathbf{x}\ _2$	0	$\{\ \mathbf{1} - \mathbf{u}\ _2 \leq 1\}$

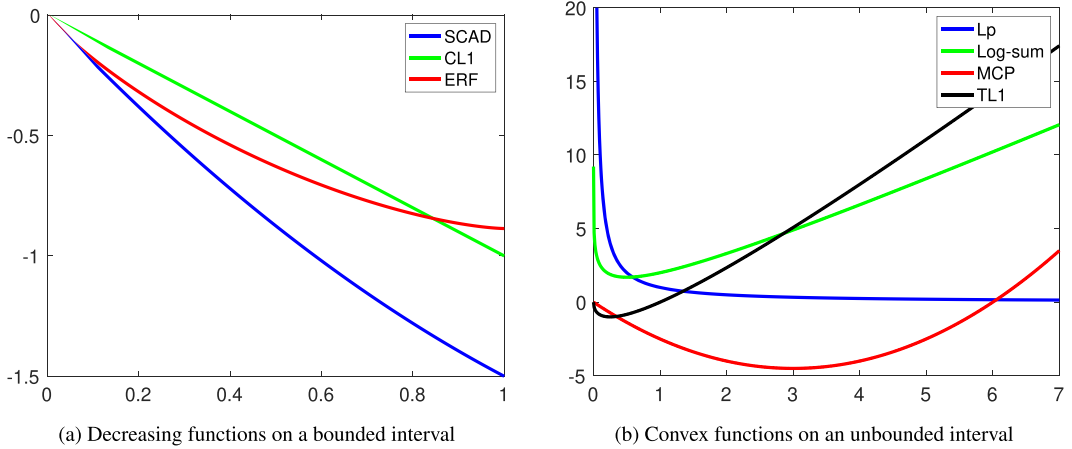


FIG. 1. Plotting the associated g function for models listed in Table 1 into two categories: (a) SCAD with $a = 1, b = 2$, CL1 with $a = 1$ and ERF with $\sigma = 1$; and (b) ℓ_p with $p = 1/2$, log-sum with $a = 1$, MCP with $a = 1, b = 3$ and TL1 with $a = 4$. This motivates Type B and Type C in Definition 2.

2.3 Properties of the Proposed Regularization

There is a wide range of analysis related to concave and symmetric regularization functions based on RIP and NSP conditions [9, 15, 47, 48] regarding the sensing matrix A . Our general model $F_{g,\alpha}^U$ satisfies all the NSP-related conditions discussed in [48] so that the exact sparse recovery can be guaranteed. Theorem 2 summarizes some important properties of the proposed regularization.

THEOREM 2. For any $\mathbf{x} \in \mathbb{R}^n$, $\alpha > 0$, and a feasible set U on the weights $\mathbf{u}$, $F_{g,\alpha}^U(\mathbf{x})$ defined in (1.2) has the following properties:

- (i) The function $F_{g,\alpha}^U(\mathbf{x}) = -(\alpha g + \delta_U)^*(-|\mathbf{x}|)$, where f^* denotes the *convex conjugate* of a function f , thus $F_{g,\alpha}^U(\mathbf{x})$ is concave in the positive cone.
- (ii) If g is symmetric on U , then $F_{g,\alpha}^U$ is symmetric on $\mathbb{R}^n$.

- (iii) If g is separable on U , then $F_{g,\alpha}^U$ is separable on $\mathbb{R}^n$.
- (iv) If g is separable and symmetric on U , $F_{g,\alpha}^U$ satisfies the increasing property on $\mathbb{R}_+^n$, i.e. $F_{g,\alpha}^U(|\mathbf{x}|) \geq F_{g,\alpha}^U(|\mathbf{x}'|)$ for any $|\mathbf{x}| \geq |\mathbf{x}'|$, and reverses the order of majorization, i.e. $F_{g,\alpha}^U(|\mathbf{x}|) \leq F_{g,\alpha}^U(|\mathbf{x}'|)$ if $|\mathbf{x}| \succ |\mathbf{x}'|$, where $\succ$ is defined in Section 2.1.
- (v) If g is separable and U is rectangular, then $F_{g,\alpha}^U$ satisfies the sub-additive property on $\mathbb{R}^n$, i.e.

$$F_{g,\alpha}^U(\mathbf{x}_1 + \mathbf{x}_2) \leq F_{g,\alpha}^U(\mathbf{x}_1) + F_{g,\alpha}^U(\mathbf{x}_2), \quad \forall \mathbf{x}_1, \mathbf{x}_2 \in \mathbb{R}^n.$$

The equality holds if $\mathbf{x}_1$ and $\mathbf{x}_2$ have disjoint support, and each coordinate of g has the same minimum.

- (vi) Let U be compact and g continuous. Then $F_{g,\alpha}^U$ is continuous and the set of sub-differentials of $-F_{g,\alpha}^U$ at the point $\mathbf{x} \geq 0$ is given by

$$\partial(-F_{g,\alpha}^U)(\mathbf{x}) = -\text{Conv} \left(\arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) \right),$$

where Conv is the convex hull of the points. In addition, the function $F_{g,\alpha}^U$ is differentiable at $\mathbf{x}$ if there exists a unique solution $\mathbf{u}$ for minimizing $F_{g,\alpha}^U(\mathbf{x})$. Consequently, if g is strongly convex, $F_{g,\alpha}^U$ is continuously differentiable on the positive cone.

- (vii) If $g(0) = 0$ and g takes its minimum value at some point in $U \subseteq \mathbb{R}_+^n$, then we have for $\alpha_1 > \alpha_2$ that

$$F_{g,\alpha_1}^U(\mathbf{x}) \leq F_{g,\alpha_2}^U(\mathbf{x}), \quad \forall \mathbf{x} \in \mathbb{R}^n.$$

Proof.

- (i) Recall the convex conjugate of a function f is defined as $f^*(\mathbf{y}) = \sup\{\langle \mathbf{x}, \mathbf{y} \rangle - f(\mathbf{x})\}$. Comparing it with the definition $F_{g,\alpha}^U(\mathbf{x})$ in (1.2), we have $F_{g,\alpha}^U(\mathbf{x}) = -(\alpha g + \delta_U)^*(-|\mathbf{x}|)$. As convex conjugate is always convex, $F_{g,\alpha}^U(\mathbf{x})$ is concave on the positive cone.
- (ii) If g is symmetric, it is straightforward that $F_{g,\alpha}^U$ is also symmetric.
- (iii) Since g is a separable function, $\min_{\mathbf{u} \in U} f(\mathbf{x}, \mathbf{u})$ breaks down into n scalar problems, and hence $F_{g,\alpha}^U$ is separable.
- (iv) The increasing property follows from the fact that we have for every $\mathbf{u} \in U$

$$\langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) \geq \langle \mathbf{u}, |\mathbf{x}'| \rangle + \alpha g(\mathbf{u}) \quad \forall |\mathbf{x}| \geq |\mathbf{x}'|.$$

Taking the minimum of both sides with respect to $\mathbf{u}$ proves the increasing property. The reverse order of majorization can be proved in the same way as in [48, Proposition 2.10].

- (v) The sub-additive property can be proved in the same way as in [48, Lemma 2.7].

(vi) It is straightforward that

$$-F_{g,\alpha}^U(\mathbf{x}) = -\min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) = \max_{\mathbf{u} \in U} -\langle \mathbf{u}, |\mathbf{x}| \rangle - \alpha g(\mathbf{u}), \quad (2.9)$$

for each $\mathbf{u} \in U$. We set $f_{\mathbf{u}}(\mathbf{x}) = -\langle \mathbf{u}, |\mathbf{x}| \rangle - \alpha g(\mathbf{u})$. As g is continuous, the map $\mathbf{u} \mapsto f_{\mathbf{u}}(\mathbf{x})$ is continuous for every $\mathbf{x}$, and for each $\mathbf{u} \in U$, the function $f_{\mathbf{u}}$ is continuous. For $\mathbf{x} \geq 0$, it follows from the *Ioffe–Tikhomirov*'s Theorem [21, Proposition 6.3] that

$$\begin{aligned} \partial(-F_{g,\alpha}^U)(\mathbf{x}) &= \text{Conv} \left\{ \bigcup_{\mathbf{u} \in T(\mathbf{x})} \partial f_{\mathbf{u}}(\mathbf{x}) \right\} = \text{Conv} \left\{ \bigcup_{\mathbf{u} \in T(\mathbf{x})} \{-\mathbf{u}\} \right\} \\ &= \text{Conv} \{-\mathbf{u} \mid -F_{g,\alpha}^U(\mathbf{x}) = f_{\mathbf{u}}(\mathbf{x})\} = \text{Conv} \{-\mathbf{u} \mid \mathbf{u} \in \arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u})\} \\ &= -\text{Conv} \{\arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u})\}, \end{aligned}$$

where $T(\mathbf{x}) = \{\mathbf{u} \mid -F_{g,\alpha}^U(\mathbf{x}) = f_{\mathbf{u}}(\mathbf{x})\}$.

(vii) Given $\mathbf{x} \in \mathbb{R}^n$ with α_2 , we denote $\mathbf{u}_2^* \in \arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha_2 g(\mathbf{u})$, which implies $F_{g,\alpha_2}^U(\mathbf{x}) = \langle \mathbf{u}_2^*, |\mathbf{x}| \rangle + \alpha_2 g(\mathbf{u}_2^*)$ and $\langle \mathbf{u}_2^*, |\mathbf{x}| \rangle + \alpha_2 g(\mathbf{u}_2^*) \leq \langle \mathbf{0}, |\mathbf{x}| \rangle + \alpha_2 g(\mathbf{0}) = 0$. It further follows from $\langle \mathbf{u}_2^*, |\mathbf{x}| \rangle \geq 0$ that $g(\mathbf{u}_2^*) \leq 0$. For $\alpha_1 > \alpha_2$, we have $\alpha_1 g(\mathbf{u}_2^*) \leq \alpha_2 g(\mathbf{u}_2^*)$ and hence $F_{g,\alpha_1}^U(\mathbf{x}) \leq \langle \mathbf{u}_2^*, |\mathbf{x}| \rangle + \alpha_1 g(\mathbf{u}_2^*) \leq \langle \mathbf{u}_2^*, |\mathbf{x}| \rangle + \alpha_2 g(\mathbf{u}_2^*) = F_{g,\alpha_2}^U(\mathbf{x})$. $\square$

Using proprieties (i)–(v), we can prove that every s -sparse vector $\mathbf{x}$ is the unique solution to (1.3) if and only if $F_{g,\alpha}$ satisfies the generalized null space property (gNSP) [48] of order s . A function F satisfies the gNSP of order s corresponding to a matrix A if

$$\ker(A) \setminus \{0\} \subset \{\mathbf{v} \in \mathbf{R}^n \mid F(\mathbf{v}_S) < F(\mathbf{v}_{\bar{S}}), \text{ for all } S \text{ with } |S| < s\}.$$

Note that $S \subseteq \{1, \dots, n\}$, $\bar{S}$ is the complement of S , and $\mathbf{v}_S$ refers to the vector with the same coordinates as $\mathbf{v}$ except zero for indices in $\bar{S}$. Please refer to [48, Theorem 4.3] for more details on gNSP. The property (vii) has algorithmic benefits, as many optimization algorithms are designed for continuously differentiable functions. We show in Theorem 3 that $F_{g,\alpha}^U$ is related to ℓ_0 and ℓ_1 if g is separable (without the assumption of strong convexity on g). The relationship of $F_{g,\alpha}^U$ to iteratively re-weighted algorithms, e.g. [10,20] is characterized in Theorem 4.

THEOREM 3. Suppose $U = [0, 1]^n$ and g is separable, i.e. $g(\mathbf{u}) = \sum_{i=1}^n g_i(u_i)$ with each g_i being a strictly decreasing function on $[0, 1]$ with a bounded derivative. If $g_i(0) = 0$ and $g_i(1) = -1$ for $1 \leq i \leq n$, we have that for $\mathbf{x} \in \mathbb{R}^n$ there are $\alpha_0 \leq \alpha_1$ such that

- (i) $\frac{1}{\alpha} F_{g,\alpha}^U(\mathbf{x}) + n = \|\mathbf{x}\|_0$, for all $0 < \alpha \leq \alpha_0$;
- (ii) $F_{g,\alpha}^U(\mathbf{x}) - \alpha g(\mathbf{1}) = \|\mathbf{x}\|_1$, for all $\alpha \geq \alpha_1$.

and consequently, we have the functional convergent results

- (i) $\frac{1}{\alpha} F_{g,\alpha}^U + n \rightarrow \ell_0$, as $\alpha \rightarrow 0$;
- (ii) $F_{g,\alpha}^U - \alpha g(\mathbf{1}) \rightarrow \ell_1$, as $\alpha \rightarrow +\infty$.

Note that the function $F_{g,\alpha}^U$ is defined in (1.2).

Proof.

- (i) For any fixed $\mathbf{x} \in \mathbb{R}^n$, we consider the derivative of $F_{g,\alpha}^U$ with respect to each of its component, i.e. $f_i := |x_i| + \alpha g'(u_i)$. If $x_i = 0$, then f_i is negative due to decreasing g , and hence the minimum is attained at $u_i = 1$. If $x_i \neq 0$, f_i is positive for a small enough α , due to the assumption that g' is bounded. Then positive derivative implies that the function is increasing, and hence the minimum is attained at $u_i = 0$. In summary, if α is sufficiently small, then we obtain that $u_i = 1$ if $x_i = 0$ and $u_i = 0$ if $x_i \neq 0$. For this choice of $\mathbf{u}$, we get that $\langle \mathbf{u}, |\mathbf{x}| \rangle = 0$ and

$$\sum_i g_i(u_i) = \sum_{x_i=0} g_i(u_i) = \sum_{x_i=0} (-1) = \|\mathbf{x}\|_0 - n,$$

which implies that $F_{g,\alpha}(\mathbf{x}) = \alpha(\|\mathbf{x}\|_0 - n)$ for a small enough α that depends on the choice of $\mathbf{x}$. By letting $\alpha \rightarrow 0$, we have $F_{g,\alpha}(\mathbf{x})/\alpha = \|\mathbf{x}\|_0 - n$ for all $\mathbf{x}$.

- (ii) Since $g(1) < g(0)$ there exists a value of $u_i \in (0, 1]$ with $g'(u_i) < 0$ and so for large enough α , the derivative $|x_i| + \alpha g'(u_i)$ is always negative. It further follows from the decreasing function g that the minimizer is always attained at $\mathbf{u} = \mathbf{1}$ to reach to the desired result. Similarly to (i), by letting $\alpha \rightarrow +\infty$, the analysis holds for all $\mathbf{x}$. $\square$

Theorem 3 presents the ideal choice for the weight, i.e. $u_i = 1$ if $x_i = 0$ and $u_i = 0$ if $x_i \neq 0$. A similar idea of zero weights on the known support was explored in [32,40,50], which is referred to as weighted ℓ_1 . In addition, Theorem 3 implies that the function $\frac{1}{\alpha} F_{g,\alpha}^U + n$ approximates the ℓ_0 norm from below. We can define a function of $H(\mathbf{x}, \alpha) := \frac{1}{\alpha} F_{g,\alpha}^U(\mathbf{x}) + n : \mathbb{R}^n \times [0, \alpha_0] \rightarrow \mathbb{R}$ as a transformation between $\|\mathbf{x}\|_0$ and $\frac{1}{\alpha_0} F_{g,\alpha_0}^U(\mathbf{x}) + n$ for a fixed α_0 . As characterized in Corollary 1, this relationship motivates us to consider a type of homotopic algorithm (discussed in Section 3) to better approximate the desired ℓ_0 norm, although $H(\mathbf{x}, \alpha)$ is not homotopy itself (it is not continuous with respect to $\mathbf{x}$).

COROLLARY 1. If $\{\alpha_i\}$ is a decreasing sequence converging to zero and g satisfies the conditions in Theorem 3, then a sequence of functions $\{\frac{1}{\alpha_i} F_{g,\alpha_i}^U\}$ are increasing, i.e.

$$\frac{1}{\alpha_0} F_{g,\alpha_0}^U(\mathbf{x}) \leq \frac{1}{\alpha_1} F_{g,\alpha_1}^U(\mathbf{x}) \leq \dots \rightarrow \|\mathbf{x}\|_0 - n, \quad \forall \mathbf{x}.$$

If we do not restrict all the g_i functions attain the same value at 1 as in Theorem 3, but rather $g_i(1)$ can take different values, then the proposed regularization $F_{g,\alpha}^U$ is equivalent to a weighted ℓ_1 model with a certain shift of $g(\mathbf{1})$; see Theorem 4.

THEOREM 4. Suppose $U = [0, 1]^n$ and g is separable, i.e. $g(\mathbf{u}) = \sum_{i=1}^n g_i(u_i)$ with each g_i being a strictly decreasing function on $[0, 1]$ with a bounded derivative. Let $\mathbf{w} = \langle w_1, \dots, w_n \rangle \geq \mathbf{0}$, then if $g_i(0) = 0$ and $g_i(1) = -w_i$ for $1 \leq i \leq n$, we have for all $\mathbf{x} \in \mathbb{R}^n$,

$$F_{g,\alpha}^U(\mathbf{x}) - \alpha g(\mathbf{1}) \rightarrow \langle \mathbf{w}, \mathbf{x} \rangle, \quad \text{as } \alpha \rightarrow +\infty. \quad (2.10)$$

In another words, for a sufficiently large α , $F_{g,\alpha}^U$ is approaching to a weighted ℓ_1 model.

Proof. Since each derivative g'_i is bounded, there exists a positive number $M_{\mathbf{x}}$ (depending on $\mathbf{x}$) such that $|x_i| + \alpha g'(u_i)$ is negative for $\alpha > M_{\mathbf{x}}$. As a result, the minimizer of $\arg \min_{u_i} \langle u_i, |x_i| \rangle + \alpha g_i(u_i)$ is

attained at $u_i = 1$. For a sufficiently large α , the minimizer $\mathbf{u}^* = \mathbf{1}$. By letting $\alpha \rightarrow +\infty$, (2.10) holds for all $\mathbf{x}$. $\square$

2.4 Exact Recovery Analysis

There are many models approximating the ℓ_0 minimization problem (1.1), and yet only a few of them have exact recovery guarantees. Motivated by the equivalence [65] between the ℓ_0 model (1.1) and (2.4) with a sufficiently small parameter ϵ , we give conditions on the g function to establish the equivalence between our proposed model (1.3) and (1.1). Note that the weight vector $\mathbf{u}$ in our formulation is not binary, but takes continuous values. By taking Table 1 and Fig. 1 into account, we consider two types of g functions defined as follows:

DEFINITION 2. Let $g : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{+\infty, -\infty\}$ be a separable function with $g(\mathbf{u}) = \sum_{i=1}^n g_i(u_i)$, for $\mathbf{u} = [u_1, \dots, u_n] \in \mathbb{R}^n$. We define

Type B: All g_i functions have bounded derivatives on $[0, 1]$, and are strictly decreasing on $[0, 1]$ with the same value at 0 and 1, i.e. there exist two constants $a, b \in \mathbb{R}$ such that $g_i(0) = a, g_i(1) = b, \forall i \in [n]$.

Type C: All g_i functions are convex on $[0, \infty)$ with the same value at 0 and the same minimum at a point other than zero, i.e. there exist two constants $a > b \in \mathbb{R}$ such that $g_i(0) = a$ and $\min_{t \geq 0} g_i(t) = b, \forall i \in [n]$.

An important characteristic both types of functions share is that they are decreasing near zero. Type B functions are defined on a bounded interval, and we enforce a box constraint on $\mathbf{u}$ for strictly decreasing g . Type C refers to convex functions defined on an unbounded interval due to $\lim_{t \rightarrow \infty} g(t) > g(0)$. Note that Theorems 3 and 4 hold when g is a Type B or Type C function. We establish the equivalent between (1.3) and (1.1) for Type B and Type C functions in Theorems 5 and 6, respectively.

THEOREM 5. Suppose g is a Type B function in Definition 2 on $U = [0, 1]^n$. There exists $\alpha > 0$ such that the model (1.3) is equivalent to (1.1), i.e. if $(\mathbf{x}^*, \mathbf{u}^*)$ is a minimizer of (1.3), then $\mathbf{x}^*$ is a minimizer of (1.1); conversely, if $\mathbf{x}^*$ is a minimizer of (1.1) then by taking $u_i^* = 1$ for $x_i^* = 0$ and $u_i^* = 0$ otherwise, $(\mathbf{x}^*, \mathbf{u}^*)$ is a minimizer of (1.3).

Proof. Since g is a Type B function, we represent $g(\mathbf{u}) = \sum_{i=1}^n g_i(u_i)$ with each g_i strictly decreasing and having bounded derivatives on $[0, 1]$. Without loss of generality, we assume that $g_i(0) = 0$ and $g_i(1) = -1, \forall i \in [n]$. Denote $s := \min_{\mathbf{x}} \{\|\mathbf{x}\|_0 \mid \mathbf{A}\mathbf{x} = \mathbf{b}\}$ and $\epsilon_0 := \min_{\mathbf{x}} \{|\mathbf{x}|_{(s)} \mid \mathbf{A}\mathbf{x} = \mathbf{b}\}$. Here $\epsilon_0 > 0$; otherwise there exists a solution to (1.1) with sparsity less than s . Since g_i has bounded derivatives, there exists a scalar $\alpha > 0$ such that $-\frac{\epsilon_0}{\alpha} < g'_i(u), \forall u \in [0, 1]$ and $i \in [n]$.

Let $(\mathbf{x}^*, \mathbf{u}^*)$ be a solution of (1.3). If $|x_i^*| \geq \epsilon_0$, we obtain $\frac{\partial f}{\partial u_i} = |x_i^*| + \alpha g'_i(u_i) > 0$. Therefore, $h(t) := t|x_i^*| + \alpha g_i(t)$ is an increasing function on $[0, 1]$, thus attaining its minimum at $t = 0$. As a result, we have $u_i^* = 0$; otherwise $(\mathbf{x}^*, \mathbf{u}^*)$ is not a minimizer of (1.3). In addition, we have $u|x| + \alpha g_i(u) \geq \alpha g_i(u) \geq \alpha g_i(1), \forall u \in [0, 1]$ and $\forall x$, as g_i is strictly decreasing. By combining two cases of $|x_i^*| \geq \epsilon_0$ and $|x_i^*| < \epsilon_0$, we estimate a lower bound of

$$\begin{aligned} f(\mathbf{x}^*, \mathbf{u}^*) &:= \sum_{i=1}^n \left[u_i^* |x_i^*| + \alpha g_i(u_i^*) \right] = \sum_{\{i \mid |x_i^*| \geq \epsilon_0\}} \alpha g_i(0) + \sum_{\{i \mid |x_i^*| < \epsilon_0\}} \left[u_i^* |x_i^*| + \alpha g_i(u_i^*) \right] \\ &\geq \sum_{\{i \mid |x_i^*| \geq \epsilon_0\}} \alpha g_i(0) + \sum_{\{i \mid |x_i^*| < \epsilon_0\}} \alpha g_i(1) = -\alpha |\{i \mid |x_i^*| < \epsilon_0\}| \geq -\alpha(n - s), \end{aligned}$$

where we use the assumptions of $g_i(0) = 0$ and $g_i(1) = -1$ together with $|\{i \mid |x_i| < \epsilon_0\}| \leq n - s$ by the definitions of s and ϵ_0 . On the other hand, the lower bound $-\alpha(n - s)$ for $f(\mathbf{x}^*, \mathbf{u}^*)$ can be achieved by any solution $\mathbf{x}$ of $\mathbf{Ax} = \mathbf{b}$ with sparsity s , by choosing $u_i = 0$ for $x_i \neq 0$ and $u_i = 1$ otherwise. Therefore, we have $f(\mathbf{x}^*, \mathbf{u}^*) = -\alpha(n - s)$.

Next we show that $\mathbf{x}^*$ must have sparsity s . If $|x_i^*| \in (0, \epsilon_0)$ for some i , then we have $u_i|x_i| + \alpha g_i(u_i) > \alpha g_i(1)$. Note that the inequality is strict, forcing $f(\mathbf{x}^*, \mathbf{u}^*)$ to be strictly greater than the lower bound $-\alpha(n - s)$, which is a contradiction. Therefore, if $|x_i^*| < \epsilon_0$, we must have $x_i = 0$. Based on the definition of ϵ_0 , we have $\|\mathbf{x}^*\|_0 = s$, which implies $\mathbf{x}^*$ is a minimizer of (1.1).

Conversely, if $\mathbf{x}^*$ is a solution of (1.1), then $\mathbf{x}^*$ satisfies $\mathbf{Ax}^* = \mathbf{b}$. With the choice of $u_i^* = 0$ for $|x_i^*| \neq 0$ and $u_i^* = 1$ otherwise, we get $(\mathbf{x}^*, \mathbf{u}^*)$ is a minimizer of (1.3) such that the objective function attains the minimum value $-\alpha(n - s)$. $\square$

THEOREM 6. Suppose $U = [0, \infty)^n$ and g is a Type C function in Definition 2. Then there exists a constant $\alpha > 0$ such that the model (1.3) is equivalent to (1.1).

Proof. We define $f(\mathbf{x}, \mathbf{u})$, s and ϵ_0 in the same way as in the proof of Theorem 5. Since g_i is convex, g'_i is increasing and hence $g'_i(u) \geq g'_i(0), \forall u \in [0, \infty)$ and $i \in [n]$. Then there exists a scalar $\alpha > 0$ such that $g'_i(0) > -\frac{\epsilon_0}{\alpha}, \forall i \in [n]$. Without loss of generality, we assume $a = 0, b = \arg \min_{t \geq 0} g_i(t) = -1$. In this setting, we get

$$u|x| + \alpha g_i(u) \geq \alpha g_i(u) \geq \alpha g_i(d_i) = -\alpha, \quad \forall u \in [0, \infty), \quad \forall x, i \in [n],$$

which implies that $f(\mathbf{x}^*, \mathbf{u}^*) \geq -\alpha(n - s)$. The rest of the proof follows exactly from the one of Theorem 5, thus omitted. $\square$

REMARK 1. In Theorems 5 and 6, we consider a linear constraint set $\Omega = \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{Ax} = \mathbf{b}\}$. All the analysis can be extended to a feasible set of inequality constraints, e.g. $\Omega_\epsilon = \{\mathbf{x} \in \mathbb{R}^n \mid \|\mathbf{Ax} - \mathbf{b}\| \leq \epsilon\}$ for $\epsilon \geq 0$. In this case, we can show that our model (2.2) with a given Ω_ϵ is equivalent to the following ℓ_0 formulation:

$$\arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{x}\|_0 + \delta_{\Omega_\epsilon}(\mathbf{x}). \quad (2.11)$$

3. Numerical Algorithms and Convergence Analysis

We describe in Section 3.1 the ADMM [8, 19] for solving the general model, with convergence analysis presented in Section 3.2. In Section 3.3, we discuss closed-form solutions of the $\mathbf{u}$ -subproblem for two specific choices of g .

3.1 The Proposed Algorithm

We define a function $\psi(\cdot)$ to unify the constrained and the unconstrained formulations, i.e. $\psi(\mathbf{x}) = \delta_\Omega(\mathbf{x})$ for (1.3) and $\psi(\mathbf{x}) = \frac{\gamma}{2} \|\mathbf{Ax} - \mathbf{b}\|_2^2$ for (1.4). We introduce a new variable $\mathbf{y}$ in order to apply ADMM to minimize

$$\arg \min_{\mathbf{u}, \mathbf{x}, \mathbf{y}} \{\langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) + \delta_U(\mathbf{u}) + \psi(\mathbf{y}) \mid \mathbf{y} = \mathbf{x}\}. \quad (3.1)$$

The corresponding augmented Lagrangian becomes

$$L_\rho(\mathbf{u}, \mathbf{x}, \mathbf{y}; \mathbf{v}) = \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}) + \delta_U(\mathbf{u}) + \psi(\mathbf{y}) + \langle \mathbf{v}, \mathbf{x} - \mathbf{y} \rangle + \frac{\rho}{2} \|\mathbf{x} - \mathbf{y}\|_2^2, \quad (3.2)$$

where $\mathbf{v}$ is the Lagrangian dual variable and ρ is a positive parameter. The ADMM scheme involves the following iterations:

$$\begin{cases} \mathbf{u}^{k+1} = \arg \min_{\mathbf{u}} L_\rho(\mathbf{u}, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k) \\ \mathbf{x}^{k+1} = \arg \min_{\mathbf{x}} L_\rho(\mathbf{u}^{k+1}, \mathbf{x}, \mathbf{y}^k; \mathbf{v}^k) \\ \mathbf{y}^{k+1} = \arg \min_{\mathbf{y}} L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^{k+1}, \mathbf{y}; \mathbf{v}^k) \\ \mathbf{v}^{k+1} = \mathbf{v}^k + \rho(\mathbf{x}^{k+1} - \mathbf{y}^{k+1}). \end{cases} \quad (3.3)$$

The original problem (1.3) jointly minimizes $\mathbf{u}$ and $\mathbf{x}$, which are updated separately in (3.3). In particular, the $\mathbf{u}$ -subproblem can be expressed as

$$\mathbf{u}^{k+1} = \arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}^k| \rangle + \alpha g(\mathbf{u}). \quad (3.4)$$

In general, one may not find a closed-form solution to (3.4). For a separable function g and a rectangular set U , the $\mathbf{u}$ -update simplifies into n one-dimensional minimization problems; refer to Section 3.3 for the $\mathbf{u}$ -update with two specific g functions that are used in experiments. For the $\mathbf{x}$ -update, it has a closed-form solution given by

$$\mathbf{x}^{k+1} = \text{shrink} \left(\mathbf{y}^k - \frac{1}{\rho} \mathbf{v}^k, \frac{1}{\rho} \mathbf{u}^{k+1} \right),$$

where $\text{shrink}(\mathbf{v}, \mathbf{u}) = \text{sign}(\mathbf{v}) \odot \max(|\mathbf{v}| - \mathbf{u}, \mathbf{0})$.

For the constrained formulation, i.e. $\psi(\mathbf{y}) = \delta_\Omega(\mathbf{y})$, the $\mathbf{y}$ -subproblem becomes

$$\arg \min_{\mathbf{y}} \left\{ \frac{1}{2} \|\mathbf{y} - (\mathbf{x}^{k+1} + \frac{1}{\rho} \mathbf{v}^k)\|_2^2 \mid \mathbf{A}\mathbf{y} = \mathbf{b} \right\}.$$

It is equivalent to a projection into the affine solution of $\mathbf{A}\mathbf{x} = \mathbf{b}$, which has a closed-form solution,

$$\mathbf{y}^{k+1} = \left(I_n - A^T(AA^T)^{-1}A \right) \left(\mathbf{x}^{k+1} + \frac{1}{\rho} \mathbf{v}^k \right) + A^T(AA^T)^{-1}\mathbf{b}.$$

For the unconstrained formulation, $\psi(\mathbf{y}) = \frac{\gamma}{2} \|\mathbf{A}\mathbf{y} - \mathbf{b}\|_2^2$, the $\mathbf{y}$ -subproblem also has a closed-form solution by solving a linear system

$$\begin{aligned} \mathbf{y}^{k+1} &= \arg \min_{\mathbf{y}} \frac{\gamma}{2} \|\mathbf{A}\mathbf{y} - \mathbf{b}\|_2^2 + \frac{\rho}{2} \|\mathbf{y} - (\mathbf{x}^{k+1} + \frac{1}{\rho} \mathbf{v}^k)\|_2^2 \\ &= \left(\rho I_n + \gamma A^T A \right)^{-1} \left(\rho \mathbf{x}^{k+1} + \mathbf{v}^k + \gamma A^T \mathbf{b} \right). \end{aligned} \quad (3.5)$$

It is worth noting that for the unconstrained formulation, a more accurate choice for ψ is $\psi(\mathbf{y}) = \frac{\alpha\gamma}{2} \|\mathbf{A}\mathbf{y} - \mathbf{b}\|_2^2$. The reason for this is further explained in Remark 2.

Algorithm 1 Adaptive ADMM algorithm for solving the general model (3.1)

Require: $\mathbf{A}$ and $\mathbf{b}$. Set parameters: ρ, η and Max .

Initialize: $\mathbf{x}^0, \mathbf{y}^0, \alpha_0$.

for $k = 1$ to Max **do**

$$\mathbf{u}^{k+1} = \arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}^k| \rangle + \alpha^k g(\mathbf{u})$$

$$\mathbf{x}^{k+1} = \text{shrink} \left(\mathbf{y}^k - \frac{1}{\rho} \mathbf{v}^k, \frac{1}{\rho} \mathbf{u}^{k+1} \right)$$

$$\mathbf{y}^{k+1} = \arg \min_{\mathbf{y}} \psi(\mathbf{y}) + \frac{\rho}{2} \|\mathbf{y} - (\mathbf{x}^{k+1} + \frac{1}{\rho} \mathbf{v}^k)\|_2^2$$

$$\mathbf{v}^{k+1} = \mathbf{v}^k + \rho(\mathbf{x}^{k+1} - \mathbf{y}^{k+1})$$

$$\alpha^{k+1} = (1 - \eta)\alpha^k$$

end for

The ADMM iterations (3.3) minimize the general model (2.2) for a fixed value of α . Following Theorem 3 and Corollary 1, we consider a type of homotopy optimization (also known as continuation approach) [16,52] to update α in order to better approximate the ℓ_0 norm. In particular, we gradually decrease α to 0 while optimizing (3.1) for each α . Algorithm 1 summarizes the overall iteration for the proposed approach.

REMARK 2. We remark that letting α approach to 0 is not exactly a homotopy algorithm, as the transformation between $\|\mathbf{x}\|_0$ and $\frac{1}{\alpha_0} F_{g, \alpha_0}^U(\mathbf{x}) + n$ is not continuous. We observe empirically the rate that α decays to zero plays a critical role in the performance of sparse recovery. On the other hand, we should minimize $\frac{1}{\alpha} F_{g, \alpha_0}^U(\mathbf{x}) + \frac{\gamma}{2} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2^2$ to approximate the ℓ_0 norm. This formulation requires the inversion of $(\rho I_n + \gamma \alpha A^T A)$ for different α in the $\mathbf{y}$ -update, which is computationally expensive, as opposed to pre-computing the inverse of $(\rho I_n + \gamma A^T A)$ with a fixed value γ .

3.2 Convergence Analysis for ADMM

We prove the convergence for the ADMM method (3.3) for the unconstrained case, i.e. $\psi(\mathbf{y}) = \frac{\gamma}{2} \|\mathbf{A}\mathbf{y} - \mathbf{b}\|_2^2$. The steps that we take to prove the convergence of (3.3) are similar to the convergence of the standard ADMM method. Yet, since we update $\mathbf{x}$ and $\mathbf{u}$ separately, the method (3.3) is not actually the ADMM iterations for the Lagrangian (3.2), and the function in the optimization (3.1) is not jointly convex with respect to $\mathbf{u}$ and $\mathbf{x}$. Note that we apply an adaptive α update in Algorithm 1, but the convergence analysis is restricted to a fixed α value. In addition, we assume that g is a Type B or Type C function that is continuously differentiable on U .

In the case of Type C functions with $U = [0, \infty)^n$, it follows from optimality conditions for each sub-problem in (3.3) that there exists $\mathbf{s}^{k+1} \geq 0$ and $\mathbf{p}^{k+1} \in \partial |\mathbf{x}^{k+1}|$ such that

$$\begin{cases} \mathbf{0} = |\mathbf{x}^k| + \alpha \nabla g(\mathbf{u}^{k+1}) - \mathbf{s}^{k+1} \\ \mathbf{0} = \mathbf{u}^{k+1} \odot \mathbf{p}^{k+1} + \mathbf{v}^k + \rho(\mathbf{x}^{k+1} - \mathbf{y}^k) \\ \mathbf{0} = \nabla \psi(\mathbf{y}^{k+1}) - \mathbf{v}^k - \rho(\mathbf{x}^{k+1} - \mathbf{y}^{k+1}), \end{cases} \quad (3.6)$$

with $\mathbf{s}^{k+1} \odot \mathbf{u}^{k+1} = \mathbf{0}$.

Note that for the Type B functions with $U = [0, 1]^n$, the optimality condition for u sub-problem is that there exists $\mathbf{s}^{k+1}, \mathbf{t}^{k+1} \geq 0$ such that $\mathbf{s}^{k+1} \odot \mathbf{u}^{k+1} = 0$, $\mathbf{t}^{k+1} \odot (\mathbf{u}^{k+1} - \mathbf{1}) = \mathbf{0}$, and $\mathbf{0} = |\mathbf{x}^k| + \alpha \nabla g(\mathbf{u}^{k+1}) - \mathbf{s}^{k+1} + \mathbf{t}^{k+1}$.

LEMMA 1. Suppose the sequence $\{\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k, \mathbf{v}^k\}$ is generated by (3.3) and denote $C_A := \|A^T A\|_2$, then the following inequality holds:

$$\|\mathbf{x}^{k+1} - \mathbf{y}^{k+1}\|_2 \leq \frac{\gamma C_A}{\rho} \|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2.$$

Proof. It is straightforward that ψ has Lipschitz continuous gradient with parameter γC_A . From the $\mathbf{v}$ -update in (3.3) and the last optimality condition in (3.6), we have $\mathbf{v}^{k+1} = \nabla \psi(\mathbf{y}^{k+1})$, and hence $\|\nabla \psi(\mathbf{y}^{k+1}) - \nabla \psi(\mathbf{y}^k)\|_2 \leq \gamma C_A \|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2$, which implies that

$$\|\mathbf{x}^{k+1} - \mathbf{y}^{k+1}\|_2 = \frac{1}{\rho} \|\mathbf{v}^{k+1} - \mathbf{v}^k\|_2 = \frac{1}{\rho} \|\nabla \psi(\mathbf{y}^{k+1}) - \nabla \psi(\mathbf{y}^k)\|_2 \leq \frac{\gamma C_A}{\rho} \|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2.$$

□

THEOREM 7. Sufficient decrease condition Suppose the sequence $\{\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k, \mathbf{v}^k\}$ is generated by (3.3). Let $\rho > \sqrt{2}\gamma C_A$, then there exists a constant $C > 0$ such that the augmented Lagrangian $L_\rho(\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k)$ defined in (3.2) satisfies

$$\begin{aligned} & L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^{k+1}, \mathbf{y}^{k+1}, \mathbf{v}^{k+1}) \\ & \leq L_\rho(\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k) - \frac{\gamma}{2} \|A(\mathbf{y}^{k+1} - \mathbf{y}^k)\|_2^2 - C \|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2^2 - \frac{\rho}{2} \|\mathbf{x}^{k+1} - \mathbf{x}^k\|_2^2, \end{aligned} \quad (3.7)$$

which implies that L_ρ decreases sufficiently.

Proof. The $\mathbf{v}$ -update in (3.3) and Lemma 1 lead to

$$\begin{aligned} & L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^{k+1}, \mathbf{y}^{k+1}, \mathbf{v}^{k+1}) - L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^{k+1}, \mathbf{y}^{k+1}, \mathbf{v}^k) \\ & = \langle \mathbf{v}^{k+1} - \mathbf{v}^k, \mathbf{x}^{k+1} - \mathbf{y}^{k+1} \rangle = \rho \|\mathbf{x}^{k+1} - \mathbf{y}^{k+1}\|_2^2 \leq \frac{\gamma^2 C_A^2}{\rho} \|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2^2. \end{aligned}$$

Using the $\mathbf{y}$ -update in (3.3) and the last optimality condition in (3.6), we have

$$\begin{aligned} & L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^{k+1}, \mathbf{y}^{k+1}, \mathbf{v}^k) - L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^{k+1}, \mathbf{y}^k, \mathbf{v}^k) \\ & = \frac{\gamma}{2} (\|A\mathbf{y}^{k+1}\|_2^2 - \|A\mathbf{y}^k\|_2^2) + \frac{\rho}{2} (\|\mathbf{y}^{k+1}\|_2^2 - \|\mathbf{y}^k\|_2^2) + \langle \gamma A^T \mathbf{b} + \rho \mathbf{x}^{k+1} + \mathbf{v}^k, \mathbf{y}^k - \mathbf{y}^{k+1} \rangle \\ & = \frac{\gamma}{2} (\|A\mathbf{y}^{k+1}\|_2^2 - \|A\mathbf{y}^k\|_2^2) + \frac{\rho}{2} (\|\mathbf{y}^{k+1}\|_2^2 - \|\mathbf{y}^k\|_2^2) + \langle \gamma \rho \mathbf{y}^{k+1} + \gamma A^T A \mathbf{y}^{k+1}, \mathbf{y}^k - \mathbf{y}^{k+1} \rangle \\ & = -\frac{\gamma}{2} \|A\mathbf{y}^{k+1} - A\mathbf{y}^k\|_2^2 - \frac{\rho}{2} \|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2^2. \end{aligned}$$

As for the $\mathbf{x}$ -update, we get

$$\begin{aligned}
& L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^{k+1}, \mathbf{y}^k; \mathbf{v}^k) - L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k) \\
&= \left\langle \mathbf{u}^{k+1}, |\mathbf{x}^{k+1}| - |\mathbf{x}^k| \right\rangle + \left\langle \mathbf{v}^k, \mathbf{x}^{k+1} - \mathbf{x}^k \right\rangle + \frac{\rho}{2} \|\mathbf{x}^{k+1} - \mathbf{y}^k\|_2^2 - \frac{\rho}{2} \|\mathbf{x}^k - \mathbf{y}^k\|_2^2 \\
&= \left\langle \mathbf{u}^{k+1}, |\mathbf{x}^{k+1}| - |\mathbf{x}^k| \right\rangle - \left\langle \mathbf{u}^{k+1} \odot \mathbf{p}^{k+1} + \rho \mathbf{x}^{k+1}, \mathbf{x}^{k+1} - \mathbf{x}^k \right\rangle + \frac{\rho}{2} \|\mathbf{x}^{k+1}\|_2^2 - \frac{\rho}{2} \|\mathbf{x}^k\|_2^2 \\
&= \left[\left\langle \mathbf{u}^{k+1}, |\mathbf{x}^{k+1}| - |\mathbf{x}^k| \right\rangle - \left\langle \mathbf{u}^{k+1} \odot \mathbf{p}^{k+1}, \mathbf{x}^{k+1} - \mathbf{x}^k \right\rangle \right] - \frac{\rho}{2} \|\mathbf{x}^{k+1} - \mathbf{x}^k\|_2^2 \\
&\leq -\frac{\rho}{2} \|\mathbf{x}^{k+1} - \mathbf{x}^k\|_2^2,
\end{aligned}$$

where the last inequality comes from the definition of subgradient. For the $\mathbf{u}$ -update, we use the fact that $\mathbf{u}^{k+1}$ is a minimizer, hence

$$L_\rho(\mathbf{u}^{k+1}, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k) - L_\rho(\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k) \leq 0.$$

Adding all these inequalities yields the desired inequality (3.7) with $C = (\frac{\rho}{2} - \frac{\gamma^2 C_A^2}{\rho})$. If $\rho > \sqrt{2}\gamma C_A$, then $C > 0$, leading to the sufficient decreasing of L_ρ . $\square$

THEOREM 8. [Residue convergent] Suppose the sequence $\{\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k, \mathbf{v}^k\}$ is generated by (3.3). If g is a Type B or Type C function and $\rho > \sqrt{2}\gamma C_A$, the following hold as $k \rightarrow \infty$

$$\mathbf{x}^{k+1} - \mathbf{x}^k \rightarrow \mathbf{0}, \quad \mathbf{y}^{k+1} - \mathbf{y}^k \rightarrow \mathbf{0}, \quad \text{and} \quad \mathbf{r}^k := \mathbf{x}^k - \mathbf{y}^k \rightarrow \mathbf{0}.$$

Proof. Since g is a Type B or Type C, then it is bounded below, and hence we denote $m_g := \min_{\mathbf{u} \in U} g(\mathbf{u})$. By telescoping summation of (3.7) from $k = 1$ to N , we obtain

$$\begin{aligned}
& \sum_{k=0}^N \left(\frac{\gamma}{2} \|A(\mathbf{y}^{k+1} - \mathbf{y}^k)\|_2^2 + C \|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2^2 + \frac{\rho}{2} \|\mathbf{x}^{k+1} - \mathbf{x}^k\|_2^2 \right) \\
& \leq \sum_{k=0}^N L_\rho(\mathbf{x}^k, \mathbf{y}^k, \mathbf{u}^k, \mathbf{v}^k) - L_\rho(\mathbf{x}^{k+1}, \mathbf{y}^{k+1}, \mathbf{u}^{k+1}, \mathbf{v}^{k+1}) \\
& = L_\rho(\mathbf{x}^0, \mathbf{y}^0, \mathbf{u}^0, \mathbf{v}^0) - L_\rho(\mathbf{x}^{N+1}, \mathbf{y}^{N+1}, \mathbf{u}^{N+1}, \mathbf{v}^{N+1}) \leq L_\rho(\mathbf{x}^0, \mathbf{y}^0, \mathbf{u}^0, \mathbf{v}^0) - \alpha m_g < \infty,
\end{aligned}$$

which implies that $\sum_{k=0}^{\infty} \|\mathbf{x}^{k+1} - \mathbf{x}^k\|_2^2 \leq \infty$ and $\sum_{k=0}^{\infty} \|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2^2 \leq \infty$. Therefore, we must have $\|\mathbf{x}^{k+1} - \mathbf{x}^k\|_2^2 \rightarrow 0$ and $\|\mathbf{y}^{k+1} - \mathbf{y}^k\|_2^2 \rightarrow 0$, as $k \rightarrow \infty$. It further follows from Lemma 1 that $\|\mathbf{x}^k - \mathbf{y}^k\|_2 \rightarrow 0$, which completes the proof. $\square$

THEOREM 9. [Stationary points] Suppose the sequence $\{\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k, \mathbf{v}^k\}$ is generated by (3.3). If $g \in \mathcal{C}^1(\mathbb{R})$ is a Type B or Type C function and $\rho > \sqrt{2}\gamma C_A$, and $(\mathbf{u}^k, \mathbf{x}^k)$ is bounded, then every limit point of $\{\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k, \mathbf{v}^k\}$, denoted by $\{\mathbf{u}^*, \mathbf{x}^*, \mathbf{y}^*, \mathbf{v}^*\}$, is a stationary point of $L_\rho(\mathbf{u}, \mathbf{x}, \mathbf{y}; \mathbf{v})$ and also $\{\mathbf{x}^*, \mathbf{u}^*\}$ is a stationary point of (1.4).

Proof. Using $\mathbf{v}^k = \nabla \psi(\mathbf{y}^k)$ from Lemma 1, i.e. $\mathbf{v}^k = \gamma A^T(A\mathbf{y}^k - \mathbf{b})$, we have

$$\begin{aligned} \psi(\mathbf{y}^k) + \langle \mathbf{v}^k, \mathbf{x}^k - \mathbf{y}^k \rangle &= \frac{\gamma}{2} \|A\mathbf{y}^k - \mathbf{b}\|_2^2 + \langle \gamma A^T(A\mathbf{y}^k - \mathbf{b}), \mathbf{x}^k - \mathbf{y}^k \rangle \\ &= \frac{\gamma}{2} \|A\mathbf{x}^k - \mathbf{b}\|_2^2 - \frac{\gamma}{2} \|A(\mathbf{x}^k - \mathbf{y}^k)\|_2^2 \geq \frac{\gamma}{2} \|A\mathbf{x}^k - \mathbf{b}\|_2^2 - \frac{\gamma C_A}{2} \|\mathbf{x}^k - \mathbf{y}^k\|_2^2. \end{aligned}$$

Consequently, we obtain that

$$L_\rho(\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k) \geq \langle \mathbf{u}^k, |\mathbf{x}^k| \rangle + \alpha g(\mathbf{u}^k) + \frac{\gamma}{2} \|A\mathbf{x}^k - \mathbf{b}\|_2^2 + \frac{\rho - \gamma C_A}{2} \|\mathbf{x}^k - \mathbf{y}^k\|_2^2. \quad (3.8)$$

If $\rho > \gamma C_A$, then $L_\rho(\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k) \geq \alpha g(\mathbf{u}^k) \geq \alpha m_g$, where $m_g := \min_{\mathbf{u} \in U} g(\mathbf{u})$, showing L_ρ is lower bounded. On the other hand, Theorem 7 gives an upper bound of $L_\rho(\mathbf{u}^k, \mathbf{x}^k, \mathbf{y}^k; \mathbf{v}^k)$, i.e. $L_\rho(\mathbf{u}^0, \mathbf{x}^0, \mathbf{y}^0; \mathbf{v}^0)$.

The boundedness of L_ρ , $\mathbf{u}^k$ and $\mathbf{x}^k$ together with (3.8) implies that $\mathbf{y}^k$ is bounded, and hence $\mathbf{v}^k$ is bounded due to $\mathbf{v}^k = \nabla \psi(\mathbf{y}^k)$. Then the Bolzano–Weierstrass Theorem guarantees that there exists a subsequence, denoted by $\{\mathbf{x}^{k_s}, \mathbf{y}^{k_s}, \mathbf{u}^{k_s}, \mathbf{v}^{k_s}\}$, that converges to a limit point, i.e. $(\mathbf{x}^{k_s}, \mathbf{y}^{k_s}, \mathbf{u}^{k_s}, \mathbf{v}^{k_s}) \rightarrow (\mathbf{x}^*, \mathbf{y}^*, \mathbf{u}^*, \mathbf{v}^*)$.

By Theorem 8, we get $\mathbf{x}^{k_s} - \mathbf{y}^{k_s} \rightarrow \mathbf{0}$, leading to $\mathbf{x}^* = \mathbf{y}^*$, and $(\mathbf{x}^{k_s-1}, \mathbf{y}^{k_s-1}) \rightarrow (\mathbf{x}^*, \mathbf{y}^*)$, and hence we have $\mathbf{v}^{k_s-1} \rightarrow \mathbf{v}^*$. Let $\mathbf{p}^{k_s}$ be the corresponding variables in the optimality condition (3.6). As $\mathbf{p}^{k_s} \in \partial|\mathbf{x}^{k_s}|$, we know $\mathbf{p}^{k_s}$ is bounded by $[-1, 1]$. Therefore, there exists a limit point of the sequence $\mathbf{p}^{k_s}$. Without loss of generality, we assume it is the sequence itself, i.e. $\mathbf{p}^{k_s} \rightarrow \mathbf{p}^*$, and hence we have $\mathbf{p}^* \in \partial|\mathbf{x}^*|$.

Type C: If g is a Type C function then $\mathbf{u} \in [0, \infty)^n$ and hence the optimality condition for the $\mathbf{u}$ -update is $\mathbf{0} = |\mathbf{x}^k| + \alpha \nabla g(\mathbf{u}^{k+1}) - \mathbf{s}^{k+1}$, with $\mathbf{s}^{k+1} \odot \mathbf{u}^{k+1} = \mathbf{0}$. We define

$$\mathbf{s}^* := \lim_{s \rightarrow \infty} |\mathbf{x}^{k_s-1}| + \alpha \nabla g(\mathbf{u}^{k_s}),$$

and so $\mathbf{s}^* \geq \mathbf{0}$ and $\mathbf{s}^{k_s} \rightarrow \mathbf{s}^*$ (since g is continuously differentiable). The optimality condition $\mathbf{s}^{k_s} \odot \mathbf{u}^{k_s} = \mathbf{0}$ implies that $\mathbf{s}^* \odot \mathbf{u}^* = \mathbf{0}$. Since all the equations in (3.6) are continuous, we can replace k by $k_s - 1$ and take the limit as $k_s \rightarrow \infty$ to get

$$\begin{cases} \mathbf{0} = |\mathbf{x}^*| + \alpha \nabla g(\mathbf{u}^*) - \mathbf{s}^* \\ \mathbf{0} = \mathbf{u}^* \odot \mathbf{p}^* + \mathbf{v}^* + \rho(\mathbf{x}^* - \mathbf{y}^*) \\ \mathbf{0} = \nabla \psi(\mathbf{y}^*) - \mathbf{v}^* - \rho(\mathbf{x}^* - \mathbf{y}^*), \end{cases}$$

where $\mathbf{s}^* \geq \mathbf{0}$ with $\mathbf{s}^* \odot \mathbf{u}^* = \mathbf{0}$, and $\mathbf{p}^* \in \partial|\mathbf{x}^*|$. Hence, $(\mathbf{x}^*, \mathbf{y}^*, \mathbf{u}^*, \mathbf{v}^*)$ is a stationary point of $L_\rho(\mathbf{u}, \mathbf{x}, \mathbf{y}; \mathbf{v})$. Furthermore, we have $\mathbf{v}^{k_s} = \nabla \psi(\mathbf{y}^{k_s})$ from the proof of Lemma 1, leading to $\mathbf{v}^* = \nabla \psi(\mathbf{y}^*)$. Together with $\mathbf{x}^* = \mathbf{y}^*$, we get

$$\begin{cases} \mathbf{0} = |\mathbf{x}^*| + \alpha \nabla g(\mathbf{u}^*) - \mathbf{s}^*, \\ \mathbf{0} = \mathbf{u}^* \odot \mathbf{p}^* + \nabla \psi(\mathbf{y}^*), \end{cases}$$

which means that $(\mathbf{x}^*, \mathbf{u}^*)$ is a stationary point of (1.4) for $U = [0, \infty)^n$.

Type B: If g is a type B function then $\mathbf{u} \in [0, 1]^n$ and hence the optimality condition for the $\mathbf{u}$ -update is $\mathbf{0} = |\mathbf{x}^k| + \alpha \nabla g(\mathbf{u}^{k+1}) - \mathbf{s}^{k+1} + \mathbf{t}^{k+1}$, with $\mathbf{s}^{k+1} \odot \mathbf{u}^{k+1} = \mathbf{0}$ and $\mathbf{t}^{k+1} \odot (\mathbf{u}^{k+1} - 1) = \mathbf{0}$ with $\mathbf{s}^{k+1} \geq \mathbf{0}$ and $\mathbf{t}^{k+1} \geq \mathbf{0}$.

Note that we have $\mathbf{s}^{k_s} - \mathbf{t}^{k_s} = |\mathbf{x}^{k_s-1}| + \alpha \nabla g(\mathbf{u}^{k_s})$. Since $\mathbf{u}^{k_s} \rightarrow \mathbf{u}^*$, $\mathbf{x}^{k_s-1} \rightarrow \mathbf{x}^*$ and g is continuously differentiable, the sequence $\mathbf{s}^{k_s} - \mathbf{t}^{k_s}$ is bounded and converges to the limit $|\mathbf{x}^*| + \alpha \nabla g(\mathbf{u}^*)$. Combining the boundedness of $\mathbf{s}^{k_s} - \mathbf{t}^{k_s}$ together with the optimality conditions, the sequences $\mathbf{s}^{k_s}$ and $\mathbf{t}^{k_s}$ must be bounded. Therefore, each sequence has a convergent sub-sequence and without loss of generality, we may assume it is the sequence itself, i.e. $\mathbf{s}^{k_s} \rightarrow \mathbf{s}^*$, and $\mathbf{t}^{k_s} \rightarrow \mathbf{t}^*$. We must have

$$\mathbf{s}^* - \mathbf{t}^* = \lim_{s \rightarrow \infty} |\mathbf{x}^{k_s-1}| + \alpha \nabla g(\mathbf{u}^{k_s}),$$

and $\mathbf{s}^* \geq \mathbf{0}$ and $\mathbf{t}^* \geq \mathbf{0}$ with the conditions $\mathbf{s}^* \odot \mathbf{u}^* = \mathbf{0}$ and $\mathbf{t}^* \odot (\mathbf{u}^* - 1) = \mathbf{0}$. The rest of the analysis is similar to the Type C functions and we get

$$\begin{cases} \mathbf{0} = |\mathbf{x}^*| + \alpha \nabla g(\mathbf{u}^*) - \mathbf{s}^* + \mathbf{t}^* \\ \mathbf{0} = \mathbf{u}^* \odot \mathbf{p}^* + \mathbf{v}^* + \rho(\mathbf{x}^* - \mathbf{y}^*) \\ \mathbf{0} = \nabla \psi(\mathbf{y}^*) - \mathbf{v}^* - \rho(\mathbf{x}^* - \mathbf{y}^*), \end{cases}$$

which means $(\mathbf{x}^*, \mathbf{y}^*, \mathbf{u}^*, \mathbf{v}^*)$ is a stationary point of $L_\rho(\mathbf{u}, \mathbf{x}, \mathbf{y}; \mathbf{v})$ and

$$\begin{cases} \mathbf{0} = |\mathbf{x}^*| + \alpha \nabla g(\mathbf{u}^*) - \mathbf{s}^* + \mathbf{t}^*, \\ \mathbf{0} = \mathbf{u}^* \odot \mathbf{p}^* + \nabla \psi(\mathbf{y}^*), \end{cases}$$

which means that $(\mathbf{x}^*, \mathbf{u}^*)$ is a stationary point of (1.4) for $U = [0, 1]^n$. □

3.3 Algorithm Updates for Different Lifting Functions

Here we consider two examples of g functions, with which the $\mathbf{u}$ -subproblem has a closed-form solution. We define one function as $g_1(\mathbf{u}) = -\frac{1}{2} \|\mathbf{u}\|_2^2$, a Type B function with $U_1 = [0, 1]^n$ and a Type C function $g_2(\mathbf{u}) = \frac{1}{2} \|\mathbf{u}\|_2^2 - \|\mathbf{u}\|_1$ with $U_2 = [0, \infty)^n$. For these combinations the update for (3.4) simplifies to

$$\text{For } g_1, U_1, \quad u_i^{k+1} = \begin{cases} 1 & \text{if } |x_i^k| \leq \frac{\alpha}{2}, \\ 0 & \text{if } |x_i^k| \geq \frac{\alpha}{2}, \end{cases} \quad \text{and for } g_2, U_2, \quad u_i^{k+1} = \max \left\{ 1 - \frac{|x_i^k|}{\alpha}, 0 \right\}.$$

Note that for this choice of g_2 , the proposed model simplifies to

$$\min_{\mathbf{x} \in \mathbb{R}^n, \mathbf{u} \in \mathbb{R}_+^n} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha \left(\frac{1}{2} \|\mathbf{u}\|_2^2 - \|\mathbf{u}\|_1 \right) \quad \text{s.t.} \quad \mathbf{A}\mathbf{x} = \mathbf{b},$$

which can be solved by a quadratic programming with linear constraints.

4. Numerical Experiments

We demonstrate the performance of Algorithm 1 with $\epsilon = 0.01$ and two specific g functions discussed in Section 3.3. We compare with the following sparsity promoting regularizations: ℓ_1 [13], $\ell_{1/2}$ [57],

transformed ℓ_1 (TL1) [62], $\ell_1 - \ell_2$ [59] and ERF [20]. For more related models, see [54]. For each model, we consider both constrained and unconstrained formulations. Specifically for the ℓ_p model, we adopt the iteratively reweighted least-squares algorithm [12] in the constrained case and use the half thresholding [57] as a proximal operator for minimizing the unconstrained $\ell_{1/2}$ formulation. Both $\ell_1 - \ell_2$ and TL1 are minimized by the difference of convex algorithm (DCA) for the best performance as reported in [59,62]. We use the online code provided by the authors of [20] to solve for the ERF model. We use the default values of model parameters suggested in respective papers; note that ℓ_1 and $\ell_1 - \ell_2$ do not involve any parameters. In Appendix C, we present a comparison between using ADMM and DCA for the proposed model. All the experiments are conducted on a Windows desktop with CPU (Intel i7-6700, 3.19GHz) and MATLAB (R2021a).

4.1 Constrained Models

We examine the performance of finding a sparse solution that satisfies the constraint $\mathbf{Ax} = \mathbf{b}$. We consider two types of sensing matrices, Gaussian and over-sampled discrete cosine transform (DCT). The Gaussian matrix is generated based on the multivariate normal distribution $\mathcal{N}(0, \Sigma)$, where $\Sigma_{ij} = (1 - r)\delta(i = j) + r$ for a parameter $r > 0$. Note that $\delta(i = j)$ is 1 if $i = j$ and zero otherwise. The over-sampled DCT matrix is defined by $A = [\mathbf{a}_1, \dots, \mathbf{a}_n] \in \mathbb{R}^{m \times n}$ with each column defined as

$$\mathbf{a}_j := \frac{1}{\sqrt{m}} \cos\left(\frac{2\pi \mathbf{w}j}{F}\right),$$

where $\mathbf{w}$ is a uniformly random vector and $F \in \mathbb{R}_+$ is a scalar. The larger the F is, the larger the coherence of the matrix A is, thus more challenging to find a sparse solution.

We fix the dimension as 64×1024 for Gaussian and DCT matrix, while generating Gaussian matrices with $r \in \{0, 0.2, 0.8\}$ and DCT matrices with $F \in \{1, 5, 10\}$. The ground truth vector $\mathbf{x}_g \in \mathbb{R}^n$ is simulated as s -sparse signal, where s is the total number of non-zero entries each drawn from normal distribution $\mathcal{N}(0, 1)$ and the support index set is also drawn randomly. We evaluate the performance by success rates where a ‘successful’ reconstruction refers to the case when the distance of the output vector $\mathbf{x}$ and the ground truth $\mathbf{x}_g$ is less than 10^{-2} , i.e.

$$\frac{\|\mathbf{x} - \mathbf{x}_g\|_2}{\|\mathbf{x}_g\|_2} \leq 10^{-2}.$$

Figure 2 presents success rates for both Gaussian and DCT matrices, and demonstrates that the proposed LL1 outperforms the state of the art in all the testing cases. For the Gaussian matrices, the parameter r has little affect on the performance, as we observe the same ranking of these models under various r values. As for the DCT matrices, the parameter F influences the coherence of the resulting matrix. For smaller F value, ℓ_p is the second best, while TL1 and $\ell_1 - \ell_2$ perform well for coherent matrices (for $F = 10$). With a well-chosen g function, the proposed LL1 framework always achieves the best results among the competing methods. The results of LL1 using g_1 with U_1 and g_2 with U_2 are similar. This phenomenon illustrates that our model works best as it is equivalent to the ℓ_0 model for small enough α .

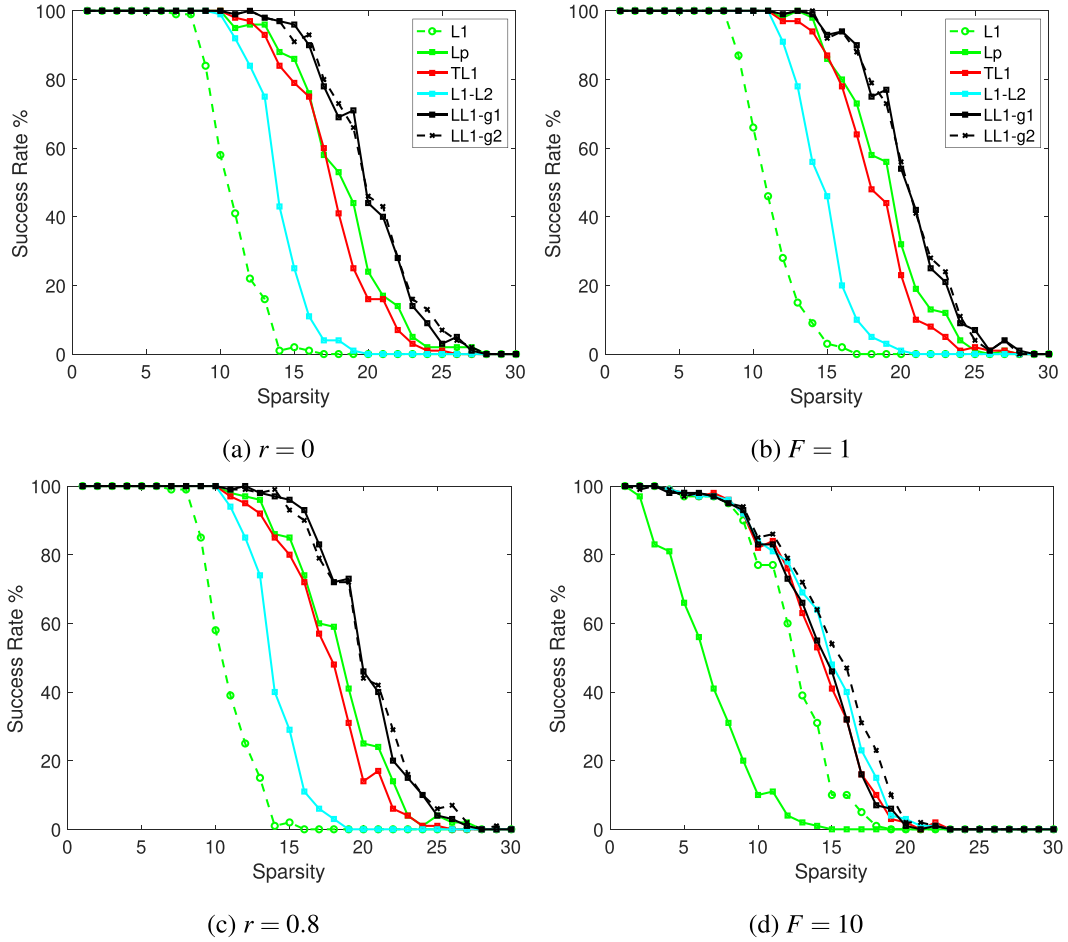


FIG. 2. Success rate comparison among all the competing methods based on Gaussian matrices (left) with $r = 0, 0.8$ and DCT matrices (right) with $F = 1, 10$.

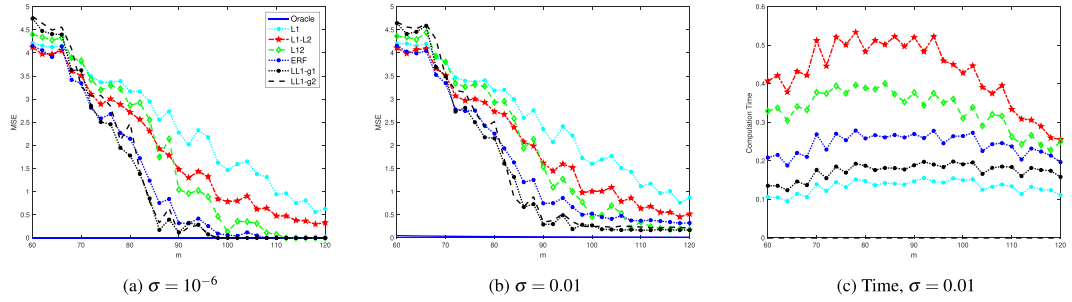
4.2 Unconstrained Models

We consider the unconstrained ℓ_0 model for comparison on noisy data:

$$\arg \min_{\mathbf{x} \in \mathbb{R}^n} \|\mathbf{x}\|_0 + \frac{\gamma}{2} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2^2, \quad (4.1)$$

where γ is a regularization parameter. We consider signals of length 512 with sparsity 130, and m measurements $\mathbf{b}$, determined by a Gaussian sensing matrix $\mathbf{A}$. The columns of $\mathbf{A}$ are normalized with mean zero and unit norm. A Gaussian noise with means zero and standard deviation σ is also added to the measurements. To evaluate the success rate of algorithms, we consider the mean square error (MSE) of the output signal $\mathbf{x}$ with the ground-truth solution $\mathbf{x}^*$ using the formula

$$\text{MSE}(\mathbf{x}) = \|\mathbf{x} - \mathbf{x}^*\|_2.$$

FIG. 3. Comparison of all algorithms for $m \times 512$ matrices.

For each algorithm, we compute the average of MSE for 100 realizations by ranging the number of measurements in the range $60 < m < 120$. Figure 3 presents the comparison results for two noise levels $\sigma \in \{10^{-6}, 0.01\}$. All the algorithms perform badly with a few measurements, and as the number of measurements m increases, their MSE error decreases. For the smaller amount of the noise ($\sigma = 10^{-6}$), our approach almost works perfectly in around 100 measurements, while other algorithms either require more measurements to achieve the nearly perfect MSE or are unable to do so. Figure 3c presents the computational times, which suggests that LL1 performs as fast as the ℓ_1 model and at the same time it has the lowest recovery error.

When the noise level is high, for instance $\sigma = 0.1$, then it is almost impossible to reconstruct the ground-truth signal using any number of measurements. In such cases, our algorithm finds a signal that is sparser and has smaller objective for any choice of the regularization parameter γ .

5. Concluding Remarks

In this paper, we proposed a lifted ℓ_1 model for sparse recovery, which describes a class of regularizations. Specifically we established the connections of this framework to various existing methods that aim to promote sparsity of the model solution. Furthermore, we proved that our method can exactly recover the sparsest solution under a constrained formulation. We promoted the use of ADMM to solve for the proposed model with convergence analysis. An alternative approach of using DCA was discussed in Section C, showing the efficiency of ADMM over DCA. Experimental results on both noise-free and noisy cases illustrate that the proposed framework outperforms the state-of-the-art methods in terms of accuracy and is comparable with the convex ℓ_1 approach in terms of computational time.

One future work involves the convergence analysis of ADMM for solving the constrained model. One difficulty lies in the fact that the corresponding function $\psi(\cdot)$ is a δ -function, which is neither differentiable nor *coercive*, and as a result, the proof presented in Section 3.2 for the unconstrained minimization is not applicable for the constrained case. We observe that the ADMM algorithm for the constrained case does converge and the augmented Lagrangian is decreasing. This empirical evidence suggests the potential to prove the convergence or the sufficient decrease of the augmented Lagrangian, which will be left as future work.

Funding

NSF CAREER 1846690; Simons Foundation grant 584960.

Data Availability Statement

No new data were generated or analysed in support of this review.

REFERENCES

1. AMIR, T., BASRI, R. & NADLER, B. (2021) The trimmed lasso: sparse recovery guarantees and practical optimization by the generalized soft-min penalty. *SIAM J. Math. Data Sci.*, **3**, 900–929.
2. ASKARI, A., NEGIAR, G., SAMBHARYA, R. & EL GHAOU, L. (2018) Lifted neural networks arXiv preprint arXiv:1805.01532.
3. BERTSIMAS, D., COPENHAVER, M. S. & MAZUMDER, R. (2017) The trimmed lasso: sparsity and robustness arXiv preprint arXiv:1708.04527.
4. BI, N. & TANG, W.-S. (2022) A necessary and sufficient condition for sparse vector recovery via $l_1 - l_2$ minimization. *Appl. Comput. Harmon. Anal.*, **56**, 337–350.
5. BLAKE, A. & ZISSERMAN, A. (1987) *Visual reconstruction*. MIT press.
6. BLANCHARD, J. D., TANNER, J. & WEI, K. (2015) CGIHT: conjugate gradient iterative hard thresholding for compressed sensing and matrix completion. *Inf. Inference*, **4**, 289–327.
7. BLUMENSATH, T. & DAVIES, M. E. (2009) Iterative hard thresholding for compressed sensing. *Appl. Comput. Harmon. Anal.*, **27**, 265–274.
8. BOYD, S., PARIKH, N., CHU, E., PELEATO, B. & ECKSTEIN, J. (2011) Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends® in Machine learning*, **3**, 1–122.
9. CANDÈS, E. J., ROMBERG, J. K. & TAO, T. (2006) Stable signal recovery from incomplete and inaccurate measurements. *Commun. Pure Appl. Math.*, **59**, 1207–1223.
10. CANDÈS, E. J., WAKIN, M. B. & BOYD, S. P. (2008) Enhancing sparsity by reweighted l_1 minimization. *J. Fourier Anal. Appl.*, **14**, 877–905.
11. CEVHER, V. An ALPS view of sparse recovery. In *2011 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 5808–5811. IEEE, 2011.
12. CHARTRAND, R. & YIN, W. Iteratively reweighted algorithms for compressive sensing. In *2008 IEEE international conference on acoustics, speech and signal processing*, pages 3869–3872. IEEE, 2008.
13. CHEN, S. S., DONOHO, D. L. & SAUNDERS, M. A. (2001) Atomic decomposition by basis pursuit. *SIAM Rev.*, **43**, 129–159.
14. DONOHO, D. L. (2006) Compressed sensing. *IEEE Trans. Inf. Theory*, **52**, 1289–1306.
15. DONOHO, D. L. & HUO, X. (2001) Uncertainty principles and ideal atomic decomposition. *IEEE Trans. Inf. Theory*, **47**, 2845–2862.
16. DUNLAVY, D. M. & O’LEARY, D. P. (2005) *Homotopy optimization methods for global optimization* Report SAND2005-7495, Sandia National Laboratories.
17. FAN, J. & LI, R. (2001) Variable selection via nonconcave penalized likelihood and its oracle properties. *J. Amer. Statist. Assoc.*, **96**, 1348–1360.
18. FOUCART, S. & RAUHUT, H. (2013) *An invitation to compressive sensing*. Springer.
19. GABAY, D. & MERCIER, B. (1976) A dual algorithm for the solution of nonlinear variational problems via finite element approximation. *Comput. Math. Appl.*, **2**, 17–40.
20. GUO, W., LOU, Y., QIN, J. & YAN, M. (2021) A novel regularization based on the error function for sparse recovery. *J. Sci. Comput.*, **87**, 31.
21. HANTOUTE, A. & LÓPEZ, M. A. (2008) Characterizations of the subdifferential of the supremum of convex functions. *J. Convex Anal.*, **15**, 831–858.
22. HORST, R. & THOAI, N. V. (1999) DC programming: overview. *J. Optim. Theory Appl.*, **103**, 1–43.
23. HUANG, X.-L., SHI, L. & YAN, M. (2015) Nonconvex sorted l_1 minimization for sparse approximation. *J. Oper. Res. Soc. China*, **3**, 207–229.
24. KUTYNIOK, G. (2013) Theory and applications of compressed sensing. *GAMM-Mitt.*, **36**, 79–101.

25. LAI, M.-J. & WANG, Y. (2021) *Sparse solutions of underdetermined linear systems and their applications*. SIAM.
26. LAI, M.-J., XU, Y. & YIN, W. (2013) Improved iteratively reweighted least squares for unconstrained smoothed l_q minimization. *SIAM J. Numer. Anal.*, **51**, 927–957.
27. LI, J., XIAO, M., FANG, C., DAI, Y., CHAO, X. & LIN, Z. (2020) Training neural networks by lifted proximal operator machines. *IEEE Trans. Pattern Anal. Mach. Intell.*, **44**, 3334–3348.
28. LOU, Y. & YAN, M. (2018) Fast L_1 – L_2 minimization via a proximal operator. *J. Sci. Comput.*, **74**, 767–785.
29. LOU, Y., YIN, P., HE, Q. & XIN, J. (2015) Computing sparse representation in a highly coherent dictionary based on difference of L_1 and L_2 . *J. Sci. Comput.*, **64**, 178–196.
30. LOU, Y., YIN, P. & XIN, J. (2016) Point source super-resolution via non-convex L_1 based methods. *J. Sci. Comput.*, **68**, 1082–1100.
31. LV, J. & FAN, Y. (2009) A unified approach to model selection and sparse recovery using regularized least squares. *Ann. Stat.*, **37**, 3498–3528.
32. MANSOUR, H. & SAAB, R. (2017) Recovery analysis for weighted l_1 -minimization using the null space property. *Appl. Comput. Harmon. Anal.*, **43**, 23–38.
33. MARQUES, E. C., MACIEL, N., NAVINER, L., CAI, H. & YANG, J. (2018) A review of sparse recovery algorithms. *IEEE Access*, **7**, 1300–1322.
34. NATARAJAN, B. K. (1995) Sparse approximate solutions to linear systems. *SIAM J. Comput.*, **24**, 227–234.
35. NEEDELL, D. & TROPP, J. A. (2009) Cosamp: iterative signal recovery from incomplete and inaccurate samples. *Appl. Comput. Harmon. Anal.*, **26**, 301–321.
36. NESTEROV, Y. (2005) Smooth minimization of non-smooth functions. *Math. Program.*, **103**, 127–152.
37. OCHS, P., DOSOVITSKIY, A., BROX, T. & POCK, T. (2015) On iteratively reweighted algorithms for nonsmooth nonconvex optimization in computer vision. *SIAM J. Imaging Sci.*, **8**, 331–372.
38. PRATER-BENNETTE, A., SHEN, L. & TRIPP, E. E. (2022) The proximity operator of the log-sum penalty. *J. Sci. Comput.*, **93**.
39. RAHIMI, Y., WANG, C., DONG, H. & LOU, Y. (2019) A scale-invariant approach for sparse signal recovery. *SIAM J. Sci. Comput.*, **41**, A3649–A3672.
40. RAUHUT, H. & WARD, R. (2016) Interpolation via weighted l_1 minimization. *Appl. Comput. Harmon. Anal.*, **40**, 321–351.
41. NAOKI SAITO. Superresolution of noisy band-limited data by data adaptive regularization and its application to seismic trace inversion. In *International Conference on Acoustics, Speech, and Signal Processing*, pages 1237–1240. IEEE, 1990.
42. SHEN, L., SUTER, B. W. & TRIPP, E. E. (2019) Structured sparsity promoting functions. *J. Optim. Theory Appl.*, **183**, 386–421.
43. SHEN, X., PAN, W. & ZHU, Y. (2012) Likelihood-based selection and sharp parameter estimation. *J. Am. Stat. Assoc.*, **107**, 223–232.
44. TAO, P. D. & AN, L. T. H. (1997) Convex analysis approach to DC programming: theory, algorithms and applications. *Acta Math. Vietnam.*, **22**, 289–355.
45. TAO, P. D. & AN, L. T. H. (1998) A DC optimization algorithm for solving the trust-region subproblem. *SIAM J. Optim.*, **8**, 476–505.
46. TIBSHIRANI, R. (1996) Regression shrinkage and selection via the lasso. *J. R. Stat. Soc. B*, **58**, 267–288.
47. TILLMANN, A. M. & PFETSCH, M. E. (2013) The computational complexity of the restricted isometry property, the nullspace property, and related concepts in compressed sensing. *IEEE Trans. Inf. Theory*, **60**, 1248–1259.
48. TRAN, H. & WEBSTER, C. (2019) A class of null space conditions for sparse recovery via nonconvex, non-separable minimizations. *Results Appl. Math.*, **3**, 100011.
49. TROPP, J. A. & GILBERT, A. C. (2007) Signal recovery from random measurements via orthogonal matching pursuit. *IEEE Trans. Inf. Theory*, **53**, 4655–4666.
50. VASWANI, N. & WEI, L. (2010) Modified-CS: modifying compressive sensing for problems with partially known support. *IEEE Trans. Signal Process.*, **58**, 4595–4607.

51. WANG, C., YAN, M., RAHIMI, Y. & LOU, Y. (2020) Accelerated schemes for the L_1/L_2 minimization. *IEEE Trans. Signal Process.*, **68**, 2660–2669.
52. WATSON, L. T. (2001) Theory of globally convergent probability-one homotopies for nonlinear programming. *SIAM J. Optim.*, **11**, 761–780.
53. WIPF, D. & NAGARAJAN, S. (2010) Iterative reweighted ℓ_1 and ℓ_2 methods for finding sparse solutions. *IEEE J. Sel. Top. Signal Process.*, **4**, 317–329.
54. FANDING, X., DUAN, J. & LIU, W. (2023) Comparative study of non-convex penalties and related algorithms in compressed sensing. *Digit. Signal Process.*, **135**, 103937.
55. YIMING, X., NARAYAN, A., TRAN, H. & WEBSTER, C. G. (2021) Analysis of the ratio of ℓ_1 and ℓ_2 norms in compressed sensing. *Appl. Comput. Harmon. Anal.*, **55**, 486–511.
56. ZONG-BEN, X. U., GUO, H.-L., WANG, Y. & ZHANG, H. (2012) Representative of $l_{1/2}$ regularization among l_q ($0 < q < 1$) regularizations: an experimental study based on phase diagram. *Acta Automat. Sinica*, **38**, 1225–1228.
57. XU, Z., CHANG, X., XU, F. & ZHANG, H. $L_{1/2}$ regularization: a thresholding representation theory and a fast solver. *IEEE Trans. Neural Netw. Learn. Syst.*, **23**:1013–1027, 2012.
58. YIN, P., ESSER, E. & XIN, J. (2014) Ratio and difference of ℓ_1 and ℓ_2 norms and sparse representation with coherent dictionaries. *Commun. Inf. Syst.*, **14**, 87–109.
59. YIN, P., LOU, Y., HE, Q. & XIN, J. (2015) Minimization of ℓ_{1-2} for compressed sensing. *SIAM J. Sci. Comput.*, **37**, A536–A563.
60. ZACH, C. & BOURMAUD, G. (2017) Iterated lifting for robust cost optimization. In *British Machine Vision Conference (BMVC)*.
61. ZHANG, C.-H. (2010) Nearly unbiased variable selection under minimax concave penalty. *Ann. Stat.*, **38**, 894–942.
62. ZHANG, S. & XIN, J. (2017) Transformed Schatten-1 iterative thresholding algorithms for low rank matrix completion. *Commun. Math. Sci.*, **15**, 839–862.
63. ZHANG, S. & XIN, J. (2018) Minimization of transformed ℓ_1 penalty: theory, difference of convex function algorithm, and robust application in compressed sensing. *Math. Program.*, **169**, 307–336.
64. TONG ZHANG. Multi-stage convex relaxation for learning with sparse regularization. In *Advances in neural information processing systems*, volume **21**, 2008.
65. ZHU, W., HUANG, Z., CHEN, J. & PENG, Z. (2021) Iteratively weighted thresholding homotopy method for the sparse solution of underdetermined linear equations. *Sci. China Math.*, **64**, 639–664.

A. Proof of Theorem 1

Proof. Set

$$f(\mathbf{x}) = \begin{cases} -J(\mathbf{x}) & \mathbf{x} \geq \mathbf{0} \\ +\infty & \text{otherwise.} \end{cases}$$

The function f is proper, lower semi-continuous and convex, hence by the *Fenchel–Moreau’s* theorem we have that $f = f^{**}$. Also, we have

$$f^*(\mathbf{y}) = \sup_{\mathbf{x} \in \mathbb{R}^n} \langle \mathbf{x}, \mathbf{y} \rangle - f(\mathbf{x}) = \sup_{\mathbf{x} \geq \mathbf{0}} \langle \mathbf{x}, \mathbf{y} \rangle + J(\mathbf{x}) = g(-\mathbf{y})$$

using $g(\mathbf{y}) := \sup_{\mathbf{x} \geq \mathbf{0}} \langle \mathbf{x}, -\mathbf{y} \rangle + J(\mathbf{x})$. In addition,

$$f(\mathbf{x}) = f^{**}(\mathbf{x}) = \sup_{\mathbf{y} \in \mathbb{R}^n} \langle \mathbf{x}, \mathbf{y} \rangle - f^*(\mathbf{y}) = \sup_{\mathbf{y} \in \mathbb{R}^n} \langle \mathbf{x}, \mathbf{y} \rangle - g(-\mathbf{y}).$$

Therefore, since $|\mathbf{x}| \geq \mathbf{0}$

$$\begin{aligned} J(|\mathbf{x}|) = -f(|\mathbf{x}|) &= -\left(\sup_{\mathbf{y} \in \mathbb{R}^n} \langle |\mathbf{x}|, \mathbf{y} \rangle - g(-\mathbf{y}) \right) = \inf_{\mathbf{y} \in \mathbb{R}^n} \langle |\mathbf{x}|, -\mathbf{y} \rangle + g(-\mathbf{y}) \\ &= \inf_{\mathbf{y} \in \mathbb{R}^n} \langle |\mathbf{x}|, \mathbf{y} \rangle + g(\mathbf{y}) = \inf_{\mathbf{y} \in U} \langle |\mathbf{x}|, \mathbf{y} \rangle + g(\mathbf{y}), \end{aligned}$$

where $U \subseteq \{\mathbf{y} \in \mathbb{R}^n \mid g(\mathbf{y}) \neq +\infty\}$. □

B. Finding the Lift Corresponding to Existing Models

Given a regularization function $J(\mathbf{x})$, we want to find a proper g function and a set U such that $J(\mathbf{x}) = F_g^U(\mathbf{x})$ up to a constant. We assume $J(\mathbf{x}) = J(|\mathbf{x}|)$ since F_g^U satisfies this condition. Consequently, we only need to consider the case $\mathbf{x} \geq \mathbf{0}$ so that the notation $\frac{\partial J}{\partial |x_i|}$ should be in place of $\frac{\partial J}{\partial x_i}$ for $x_i \geq 0$. Using Theorem 1, we can directly find F_g^U ; however, sometimes it might be easier to use the following observation which leads to simpler computations. Suppose F_g^U has a unique minimizer $\mathbf{u}$, and hence $\mathbf{u}$ satisfies $\mathbf{u} = \nabla_{\mathbf{x}} F_g^U(\mathbf{x}) = \nabla_{\mathbf{x}} J(\mathbf{x})$. Assuming that the minimum of (2.5) is finite, the optimality condition gives $|\mathbf{x}| + \nabla_{\mathbf{u}} g = 0$ for $\mathbf{u} \in \text{int}(U)$, where $\text{int}(U)$ denotes the interior of the set U (Note that $|\mathbf{x}| + \nabla_{\mathbf{u}} g$ can have non-zero coordinates on the boundary of U .) Thus, we only need to solve the following two equations for a function g with respect to $\mathbf{u}$ on the feasible set U :

$$\begin{cases} \mathbf{u} = \nabla_{\mathbf{x}} J(\mathbf{x}), \\ |\mathbf{x}| + \nabla_{\mathbf{u}} g = 0, \mathbf{u} \in \text{int}(U). \end{cases} \quad (\text{B.1})$$

- (i) **ℓ_p model:** Consider $J = J^{\ell_p}/p$, and note that $\frac{\partial J}{\partial |x_i|} = |x_i|^{p-1}$. For $g(\mathbf{u}) = \sum g_i(u_i)$ and $\mathbf{x} \in \mathbb{R}^n$, the (B.1) simplifies into

$$\begin{cases} u_i = |x_i|^{p-1}, \\ |x_i| + g'_i(u_i) = 0, \end{cases}$$

for all i . From the first equation we get that $|x_i| = u_i^{\frac{1}{p-1}}$ and then from the second equation we get $g'_i(u_i) = -u_i^{\frac{1}{p-1}}$. A solution for g is $g_i(u_i) = \frac{1-p}{p} u_i^{\frac{p}{p-1}}$ for $u_i \geq 0$. Finally taking $U = \mathbb{R}_+^n$ and $g(\mathbf{u}) = \sum_i \frac{1-p}{p} u_i^{\frac{p}{p-1}}$, one can check that $F_g^U = J$.

- (ii) **log-sum penalty:** Consider $J = J_a^{\log}$, and note that $\frac{\partial J}{\partial |x_i|} = \frac{1}{|x_i|+a}$. For $g(\mathbf{u}) = \sum g_i(u_i)$ and $\mathbf{x} \in \mathbb{R}^n$, the (B.1) simplifies into

$$\begin{cases} u_i = \frac{1}{|x_i|+a}, \\ |x_i| + g'_i(u_i) = 0, \end{cases}$$

for all i . From the first equation we get that $|x_i| = \frac{1}{u_i} - a$ and then from the second equation we get $g'_i(u_i) = a - \frac{1}{u_i}$. A solution for g is $g_i(u_i) = au_i - \log(u_i)$ for $u_i > 0$. Finally taking $U = \mathbb{R}_{>0}^n$ and $g(\mathbf{u}) = \sum_i (au_i - \log(u_i))$, one can check that $F_g^U + 1 = J$.

- (iii) **Smoothly clipped lasso model:** Consider $J = J_{\gamma,\lambda}^{\text{SCAD}}$, and note that

$$\frac{\partial f_{\lambda,\gamma}^{\text{SCAD}}(t)}{\partial |t|} = \begin{cases} \lambda & \text{if } |t| \leq \lambda, \\ \frac{\lambda\gamma - t}{\gamma - 1} & \text{if } \lambda < |t| \leq \gamma\lambda, \\ 0 & \text{if } |t| > \gamma\lambda. \end{cases} \quad (\text{B.2})$$

For $g(\mathbf{u}) = \sum g_i(u_i)$ and $\mathbf{x} \in \mathbb{R}^n$, the first equation in (B.1) simplifies into

$$u_i = \begin{cases} \lambda & \text{if } |x_i| \leq \lambda, \\ \frac{\lambda\gamma - |x_i|}{\gamma - 1} & \text{if } \lambda < |x_i| \leq \gamma\lambda, \\ 0 & \text{if } |x_i| > \gamma\lambda. \end{cases}$$

for all i . In the case of $\lambda < |x_i| \leq \gamma\lambda$, we get that $|x_i| = \gamma\lambda - (\gamma - 1)u_i$, which means we should have $g'_i(u_i) = -\gamma\lambda + (\gamma - 1)u_i$ and $u_i \leq \lambda$. By taking $U = [0, \lambda]^n$ and $g(\mathbf{u}) = \sum_i \left(-\gamma\lambda u_i + (\gamma - 1)\frac{u_i^2}{2} \right)$, one can check that $F_g^U + \frac{(\gamma+1)\lambda^2}{2} = J$.

- (iv) **Mini-max concave penalty:** Consider $J = J_{\gamma,\lambda}^{\text{MCP}}$, and note that

$$\frac{\partial f_{\lambda,\gamma}^{\text{MCP}}(t)}{\partial |t|} = \begin{cases} \lambda - \frac{t}{\gamma} & \text{if } |t| \leq \gamma\lambda, \\ 0 & \text{if } |t| > \gamma\lambda. \end{cases} \quad (\text{B.3})$$

For $g(\mathbf{u}) = \sum g_i(u_i)$ and $\mathbf{x} \in \mathbb{R}^n$, the first equation in (B.1) simplifies into

$$u_i = \begin{cases} \lambda - \frac{|x_i|}{\gamma} & \text{if } |x_i| \leq \gamma\lambda, \\ 0 & \text{if } |x_i| > \gamma\lambda, \end{cases}$$

for all i . When $|x_i| \leq \gamma\lambda$, we obtain $|x_i| = \gamma(\lambda - u_i)$, which implies that $g'_i(u_i) = -\gamma(\lambda - u_i)$ from the second equation in (B.1). Therefore, we set $U = [0, \infty)^n$ and $g(\mathbf{u}) = \sum_i \left(-\gamma(\lambda u_i - \frac{u_i^2}{2}) \right)$, leading to $F_g^U + \frac{1}{2}\gamma\lambda^2 = J$.

- (v) **Capped ℓ_1 model:** Consider $J = J_a^{\text{CL1}}$, and note that

$$\frac{\partial J}{\partial |x_i|} = \begin{cases} 1 & \text{if } |x_i| < a, \\ 0 & \text{if } |x_i| > a. \end{cases} \quad (\text{B.4})$$

For $g(\mathbf{u}) = \sum g_i(u_i)$ and $\mathbf{x} \in \mathbb{R}^n$, the first equation in (B.1) simplifies into

$$u_i = \begin{cases} 1 & \text{if } |x_i| < a, \\ 0 & \text{if } |x_i| > a, \end{cases}$$

for all i . Note that the second equation in (B.1) only happens if the minimizer is in the interior of the set U . Consider $U = [0, 1]^n$, therefore for this case since the minimizer is on the boundary, therefore we need a g function which is non-zero in the interior of U and for $|x_i| < a$ we have $|x_i| + g'_i(u_i) < 0$ and for $|x_i| > a$ we have $|x_i| + g'_i(u_i) > 0$. Therefore $g'_i(u_i) = -a$ and a solution for this is $g_i(u_i) = -au_i$. Finally taking $U = [0, 1]^n$ and $g(\mathbf{u}) = \sum_i (-au_i)$, one can check that $F_g^U + a = J$.

- (vi) **Transformed ℓ_1 model:** Consider $J = J_a^{\text{TL1}}/(a+1)$, and note that $\frac{\partial J}{\partial |x_i|} = \frac{a}{(a+|x_i|)^2}$. For $g(\mathbf{u}) = \sum g_i(u_i)$ and $\mathbf{x} \in \mathbb{R}^n$, the (B.1) simplifies into

$$\begin{cases} u_i = \frac{a}{(a+|x_i|)^2}, \\ |x_i| + g'_i(u_i) = 0, \end{cases}$$

for all i . From the first equation we get that $|x_i| = \sqrt{\frac{a}{u_i}} - a$ and then from the second equation we get $g'_i(u_i) = a - \sqrt{\frac{a}{u_i}}$. A solution for g is $g_i(u_i) = au_i - 2\sqrt{au_i}$ for $u_i \geq 0$. Finally taking $U = \mathbb{R}_+^n$ and $g(\mathbf{u}) = \sum_i au_i - 2\sqrt{au_i}$, one can check that $F_g^U + 1 = J$.

- (vii) **Error function penalty:** Consider $J = J_\sigma^{\text{ERF}}$, and note that $\frac{\partial J}{\partial |x_i|} = e^{-x_i^2/\sigma^2}$ and $e^{-x_i^2/\sigma^2} \in (0, 1]$. For $g(\mathbf{u}) = \sum g_i(u_i)$ and $\mathbf{x} \in \mathbb{R}^n$, the (B.1) simplifies into

$$\begin{cases} u_i = e^{-x_i^2/\sigma^2}, \\ |x_i| + g'_i(u_i) = 0, \end{cases}$$

for all i . From the first equation we get that $|x_i| = \sigma\sqrt{-\log(u_i)}$ and then from the second equation we get $g'_i(u_i) = -\sigma\sqrt{-\log(u_i)}$. A solution for g is $g_i(u_i) = \sigma \int_{u_i}^1 \sqrt{-\log(\tau)} d\tau$ for $u_i \in (0, 1]$. Finally taking $U = [0, 1]^n$ and $g(\mathbf{u}) = \sigma \sum_i \int_{u_i}^1 \sqrt{-\log(\tau)} d\tau$, one can check that $F_g^U = J$.

- (viii) $\ell_1 - \ell_2$ **penalty**: As for $J = J^{L1-L2}$ working with derivatives is a bit challenging so we directly apply Theorem 1 for finding the lifted form. First note that

$$\sup_{\mathbf{x} \geq 0} \langle \mathbf{x}, \mathbf{y} \rangle - \|\mathbf{x}\|_2 = \begin{cases} +\infty & \|\mathbf{y}_+\|_2 > 1, \\ 0 & \|\mathbf{y}_+\|_2 \leq 1. \end{cases}$$

Hence we get that

$$g(\mathbf{u}) = \sup_{\mathbf{x} \geq 0} \langle \mathbf{x}, -\mathbf{u} \rangle + \|\mathbf{x}\|_1 - \|\mathbf{x}\|_2 = \begin{cases} +\infty & \|(\mathbf{1} - \mathbf{u})_+\|_2 > 1, \\ 0 & \|(\mathbf{1} - \mathbf{u})_+\|_2 \leq 1. \end{cases}$$

Therefore, the set $U = \{\mathbf{u} \mid \|(\mathbf{1} - \mathbf{u})_+\|_2 \leq 1\}$ and the function $g = 0$. Notice that we can relax the set U to $\{\mathbf{u} \mid \|\mathbf{1} - \mathbf{u}\|_2 \leq 1\}$ and still get the same J function.

C. Comparing the ADMM and DCA based algorithms

Alternatively to ADMM, one can minimize the proposed model (1.3) and (1.4) by the graduated non-convexity algorithm [5] and the DCA [22,44,45]. Specifically for DCA, since our function $-F_{g,\alpha}^U$ is convex on the positive cone, then we use the algorithm introduced in [37], where we only need to find the sub-differential of the function on $\mathbb{R}_+^n$. We use the function $\psi(\cdot)$ to unify the constrained and the unconstrained formulations and we have the model

$$\arg \min_{\mathbf{x} \in \mathbb{R}_+^n} F_{g,\alpha}^U(\mathbf{x}) + \psi(\mathbf{x}). \quad (\text{C.1})$$

Since the function $F_{g,\alpha}^U$ is concave on $\mathbb{R}_+^n$, it can be written as a difference of two convex functions, i.e. $F_{g,\alpha}^U(\mathbf{x}) + \psi(\mathbf{x}) = h_1(\mathbf{x}) - h_2(\mathbf{x})$, where $h_1(\mathbf{x}) = \psi(\mathbf{x}) + \frac{\beta}{2} \|\mathbf{x}\|_2^2$ and $h_2(\mathbf{x}) = \frac{\beta}{2} \|\mathbf{x}\|_2^2 - F_{g,\alpha}^U(\mathbf{x})$ for $\beta \geq 0$. An interesting fact about using a DCA form is that if g is a Type B or Type C then for $\mathbf{x} \geq 0$ we have sub-differentials of the form

$$\partial(-F_{g,\alpha}^U)(\mathbf{x}) = -\arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}| \rangle + \alpha g(\mathbf{u}).$$

For $\mathbf{x} \in \mathbb{R}^n$, take $\mathbf{u}_{g,\alpha}(\mathbf{x}^k) \in \arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}^k| \rangle + \alpha g(\mathbf{u})$ then the DCA iterations become

$$\mathbf{x}^{k+1} = \arg \min_{\mathbf{x} \in \mathbb{R}_+^n} \psi(\mathbf{x}) + \frac{\beta}{2} \|\mathbf{x}\|_2^2 - \left\langle \beta |\mathbf{x}^k| - \mathbf{u}_{g,\alpha}(\mathbf{x}^k), |\mathbf{x}| \right\rangle. \quad (\text{C.2})$$

We implement the DCA iterations of (C.2) for $\beta = 0$ for its simplicity and efficiency as opposed to $\beta > 0$. In addition, we can consider an adaptive scheme to update α , which is adopted in Algorithm 2.

We compare ADMM (Algorithm 1) and DCA (Algorithm 2) for minimizing the same constrained formulation (1.3) with g_1 and g_2 discussed in Section 3.3. We are particularly interested in the algorithmic behaviours when dynamically updating α . As mentioned in Theorem 5, α is supposed to be small enough to approximate the ℓ_0 solution. A common way involves an exponential decay in the form of $\alpha^{k+1} = (1 - \eta)\alpha^k$, for $\eta \in (0, 1)$. If the parameter η is close to 1 then α converges to zero too quickly and hence

Algorithm 2 Homotopy-based DCA algorithm

Require: $\mathbf{x}^0, \alpha^0 > 0, \eta \geq 0$, function g , set U and MaxOuter

for $k = 1$ to MaxOuter **do**

$$\mathbf{u}^k = \arg \min_{\mathbf{u} \in U} \langle \mathbf{u}, |\mathbf{x}^k| \rangle + \alpha^k g(\mathbf{u}).$$

$$\mathbf{x}^{k+1} = \arg \min_{\mathbf{x} \in \mathbb{R}^n} \psi(\mathbf{x}) + \langle \mathbf{u}_{g,\alpha}(\mathbf{x}^k), |\mathbf{x}| \rangle.$$

$$\alpha^{k+1} = (1 - \eta)\alpha^k$$

end for

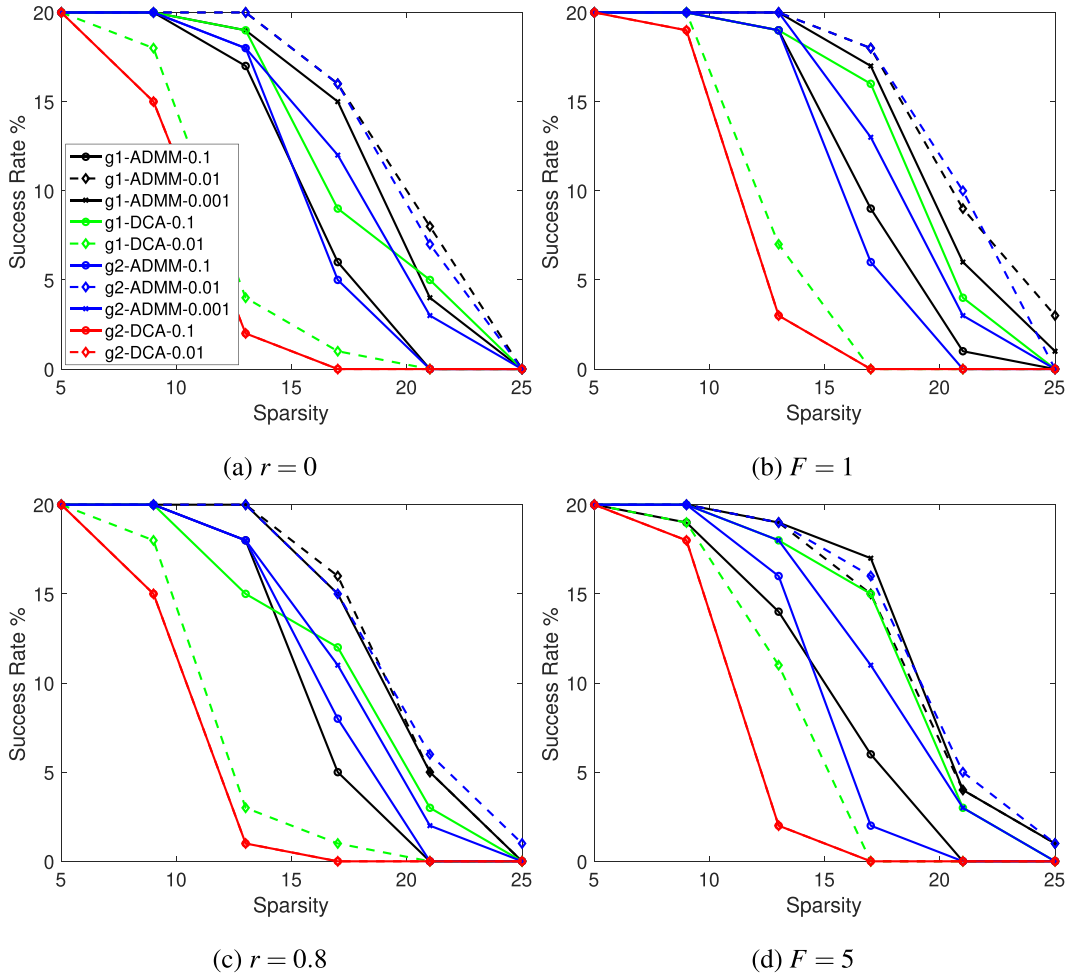


FIG. C1. Success rate comparison of ADMM and DCA with two g functions and different η for Gaussian matrices (left) with $r = 0$ and $r = 0.8$ and DCT matrices (right) with $F = 1, 5$.

the algorithm cannot converge to a good local minimum as it is equivalent to having $\alpha = 0$ in the very beginning. On the other hand, if η is close to 0 then α slowly decreases to zero; and as a result, Algorithm 1 may terminate before α is small enough for $F_{g,\alpha}^U$ to approximate the ℓ_0 norm.

The comparison between ADMM and DCA on Gaussian and DCT matrices is presented in Fig. C1. By fixing $\eta = 0.1$, we select the optimal $\rho = 2^j$ for ADMM with $j \in \{-1, 0, 1, 2, 3, 4, 5, 6, 7, 8\}$ that achieves the smallest relative error to the ground-truth. Then using the optimal parameters ρ and γ , Fig. C1 presents the ADMM results for $\eta \in \{0.001, 0.01, 0.1\}$ and the DCA ones for $\eta \in \{0.01, 0.1\}$. For all the cases, ADMM is superior to DCA in that it is less sensitive to η . In addition, DCA consists of two loops and hence it is generally slower than ADMM. Our experiment shows that a suitable choice for our experiments is $\eta = 0.01$. These comparisons show that the lifted form works well with ADMM, but not necessarily with DCA. This may be due to the fact that DCA usually converges in a few iterations and this is not enough for α to get close to 0. One may incorporate the decreasing of α into the $\mathbf{x}$ -update iteration, or update α according to the result of each iteration, which is left for future work.