



Fast communication

Accelerated reweighted nuclear norm minimization algorithm for low rank matrix recovery



Xiaofan Lin*, Gang Wei

School of Electronic and Information Engineering, South China University of Technology, Wushan Road, Tianhe District, Guangzhou 510641, PR China

ARTICLE INFO

Article history:

Received 16 August 2014

Received in revised form

3 November 2014

Accepted 5 February 2015

Available online 14 February 2015

Keywords:

Matrix rank minimization

Matrix completion

Compressed sensing

ABSTRACT

In this paper we propose an accelerated reweighted nuclear norm minimization algorithm to recover a low rank matrix. Our approach differs from other iterative reweighted algorithms, as we design an accelerated procedure which makes the objective function descend further at every iteration. The proposed algorithm is the accelerated version of a state-of-the-art algorithm. We provide a new analysis of the original algorithm to derive our own accelerated version, and prove that our algorithm is guaranteed to converge to a stationary point of the reweighted nuclear norm minimization problem. Numerical results show that our algorithm requires distinctly fewer iterations and less computational time than the original one to achieve the same (or very close) accuracy, in some problem instances even require only about 50% computational time of the original one, and is also notably faster than several other state-of-the-art algorithms.

© 2015 Elsevier B.V. All rights reserved.

1. Introduction

There is a rapidly growing interest in the low rank matrix recovery problem or the rank minimization problem, which recovers an unknown low rank matrix from very limited information [1–5]. This paper deals with the following affine rank minimization problem:

$$\min_{\mathbf{X}} \text{rank}(\mathbf{X}) \quad \text{s.t. } \mathcal{A}(\mathbf{X}) = \mathbf{b}, \quad (1)$$

where the linear map $\mathcal{A}: \mathbb{R}^{m \times n} \rightarrow \mathbb{R}^S$ and the vector $\mathbf{b} \in \mathbb{R}^S$ are known.

The above problem aims to find a matrix of minimum rank that satisfies a given system of linear equality constraints, which is a useful idea that can be applied to various applications such as signal or image processing [6,7], subspace segmentation [8], collaborative filtering [9] and system

identification [10]. Although (1) is simple in form, it is a challenging optimization problem due to the discrete nature of the rank function.

A commonly used heuristic introduced in [11] is replacing the rank function with the nuclear norm, which is the sum of the singular values of the matrix. This technique is based on the fact that the nuclear norm minimization is the tightest convex relaxation of the rank minimization problem. The new formula can be given by

$$\min_{\mathbf{X}} \|\mathbf{X}\|_* \quad \text{s.t. } \mathcal{A}(\mathbf{X}) = \mathbf{b}, \quad (2)$$

which can be rewritten as follows under some conditions (recall the Lagrangian relaxation):

$$\min_{\mathbf{X}} \lambda \|\mathbf{X}\|_* + \frac{1}{2} \|\mathcal{A}(\mathbf{X}) - \mathbf{b}\|_2^2, \quad (3)$$

where $\|\mathbf{X}\|_* := \sum_{i=1}^r \sigma_i(\mathbf{X})$ denotes the nuclear norm.

However, the nuclear norm minimization requires more measurements for exact recovery of the low rank solution. Recently, a tighter relaxation, called the Schatten

* Corresponding author.

E-mail address: walkerlin@foxmail.com (X. Lin).

p -norm¹ minimization [12–14], was introduced instead of the previous nuclear norm minimization in order to give a better approximation to the original problem. The Schatten p -norm with $0 < p \leq 1$ is defined as

$$\|\mathbf{X}\|_p := \left(\sum_{i=1}^r \sigma_i(\mathbf{X})^p \right)^{1/p}. \quad (4)$$

It is easy to see that when $p=1$, the Schatten p -norm is exactly the nuclear norm, and when p tends towards 0, the Schatten p -norm becomes closer to the rank function. Thus, Schatten p -norm is more general than the nuclear norm. Theoretical analysis in [12,13,15] shows that the Schatten p -norm needs fewer measurements with small p for exact recovery.

With the above definition, the Schatten p -norm minimization can be formulated as follows:

$$\min_{\mathbf{X}} P_0(\mathbf{X}) := \lambda p \|\mathbf{X}\|_p^p + \frac{1}{2} \|\mathbf{A}(\mathbf{X}) - \mathbf{b}\|_2^2, \quad (5)$$

When $0 < p < 1$, the above problem is nonconvex. To solve this nonconvex problem efficiently, some iterative reweighted algorithms have been proposed and analyzed in recent published literature [2,3,14]. The key idea of iterative reweighted technique is to solve a convex problem with a given weight at each iteration, and update the weight at every turn. This idea is commonly used to solve Schatten p -norm or L_p norm minimization.

In this paper we propose a novel iterative reweighted algorithm that solves the Schatten p -norm minimization problem (i.e., (5)) very efficiently. Different from other iterative reweighted algorithms, we add an accelerated procedure which makes the objective function descend further at every iteration. The proposed algorithm is the accelerated version of a state-of-the-art algorithm, called the reweighted nuclear norm minimization (RNNM) algorithm [16]. However, our accelerated version is notably faster than the original one, and achieves the same (or very close) accuracy.

The rest of this paper is organized as follows. In Section 2, we give some preliminary results that are used in our analysis. In Section 3, we give our own analysis of the original RNNM algorithm, and Section 4 uses this analysis to derive a sketch of the accelerated version. Some algorithm details are discussed in Sections 5 and 6. The convergence analysis, which proves that our algorithm is guaranteed to converge to a stationary point of $P_0(\mathbf{X})$ (see (5)), is given in Section 7. Numerical results are reported in Section 8. In Section 9, we give some conclusions.

2. Preliminaries

We first introduce the following unconstrained smooth nonconvex problem to approximate (5):

$$\min_{\mathbf{X}} P(\mathbf{X}, \varepsilon) := \lambda p \|\mathbf{X}\|_{p,\varepsilon}^p + \frac{1}{2} \|\mathbf{A}(\mathbf{X}) - \mathbf{b}\|_2^2, \quad (6)$$

where $\mathbf{X} \in \mathbb{R}^{m \times n}$ and

$$\|\mathbf{X}\|_{p,\varepsilon} := \left(\sum_{i=1}^m (\sigma_i(\mathbf{X}) + \varepsilon)^p \right)^{1/p}.^2 \quad (7)$$

with $\varepsilon > 0$. Here we assume $m \leq n$ without loss of generality. Compare (7) with (4), one can see that with ε tends to 0, (7) will get closer to (4).

The following definition will be used frequently in this paper.

Definition 1. Let $\mathbf{X} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T$ be the reduced singular value decomposition (SVD) of $\mathbf{X}$ with rank r . For each $v \geq 0$, the weighted singular value shrinkage operator $D_v(\cdot, \mathbf{w})$ is defined by

$$D_v(\mathbf{X}, \mathbf{w}) = \mathbf{U} \text{Diag}\{(\sigma_i(\mathbf{X}) - v w_i)_+, i = 1, 2, \dots, r\} \mathbf{V}^T,$$

where $(\cdot)_+ := \max\{a, 0\}$ for any $a \in \mathbb{R}$; w_i , $i = 1, 2, \dots, r$ are elements of vector $\mathbf{w}$; $\sigma_i(\mathbf{X})$, $i = 1, 2, \dots, r$, are singular values of $\mathbf{X}$.

With Definition 1, we give the following lemma, which has been proved in [16].

Lemma 1. For each $v \geq 0$ and a constant matrix $\mathbf{Y} \in \mathbb{R}^{m \times n}$, the weighted singular value shrinkage operator $D_v(\mathbf{Y}, \mathbf{w})$ is the optimal solution of the following problem:

$$\min_{\mathbf{X}} \sum_{i=1}^m w_i \sigma_i(\mathbf{X}) + \frac{1}{2} \|\mathbf{X} - \mathbf{Y}\|_F^2.$$

For convenience, the abbreviation ORNNM refers to the original RNNM algorithm that we aim to accelerate in the sequel.

3. A new analysis for original RNNM algorithm

In this paper, we propose an accelerated procedure to improve the performance of ORNNM. This algorithm has been analyzed in [16]. However, to derive our accelerated version, here we give our own analysis.

Inspired by [17,4], where auxiliary functions for analyzing the L_1 norm/Schatten p -norm minimization were constructed, we define a different auxiliary function as follows:

$$Q(\mathbf{X}, \mathbf{Y}, \mathbf{w}, \varepsilon) = F(\mathbf{X}, \mathbf{w}, \varepsilon) + G(\mathbf{X}, \mathbf{Y}) \quad (8)$$

where

$$F(\mathbf{X}, \mathbf{w}, \varepsilon) := \lambda \sum_{i=1}^m \left(p w_i (\sigma_i(\mathbf{X}) + \varepsilon) + (1-p) w_i^{p/(p-1)} \right) + \frac{1}{2} \|\mathbf{A}(\mathbf{X}) - \mathbf{b}\|_2^2 \quad (9)$$

and

$$G(\mathbf{X}, \mathbf{Y}) := -\frac{1}{2} \|\mathbf{A}(\mathbf{X} - \mathbf{Y})\|_2^2 + \frac{1}{2} \|\mathbf{X} - \mathbf{Y}\|_F^2. \quad (10)$$

In the following, we illustrate that minimizing $Q(\mathbf{X}, \mathbf{Y}, \mathbf{w}, \varepsilon)$ will lead to minimizing $P_0(\mathbf{X})$ (see (5) for definition).

First, it can be verified that $\mathbf{w}^*$, whose entries $w_i^* = (\sigma_i(\mathbf{X}) + \varepsilon)^{p-1}$, $i = 1, 2, \dots, m$, minimizes $F(\mathbf{X}, \mathbf{w}, \varepsilon)$ over

¹ Strictly speaking, the Schatten p -norm is not a norm when $0 < p < 1$ since it is nonconvex. However it is called a norm customarily.

² Strictly speaking, this is not a norm. However we preserve the norm notation for conventional reason, since the Schatten p -norm is neither a norm but it is still common to use a norm notation for it in the literature.

Download English Version:

<https://daneshyari.com/en/article/562456>

Download Persian Version:

<https://daneshyari.com/article/562456>

[Daneshyari.com](https://daneshyari.com)