



An efficient method for robust projection matrix design



Tao Hong^{a,*}, Zhihui Zhu^b

^a Department of Computer Science, Technion - Israel Institute of Technology, Haifa, 32000, Israel

^b Department of Electrical Engineering, Colorado School of Mines, Golden, CO 80401 USA

ARTICLE INFO

Article history:

Received 19 January 2017

Revised 3 September 2017

Accepted 6 September 2017

Available online 7 September 2017

Keywords:

Robust projection matrix

Sparse representation error (SRE)

High dimensional dictionary

Mutual coherence

ABSTRACT

Our objective is to efficiently design a robust projection matrix Φ for the Compressive Sensing (CS) systems when applied to the signals that are not exactly sparse. The optimal projection matrix is obtained by mainly minimizing the average coherence of the equivalent dictionary. In order to drop the requirement of the sparse representation error (SRE) for a set of training data as in [15,16], we introduce a novel penalty function independent of a particular SRE matrix. Without requiring of training data, we can efficiently design the robust projection matrix and apply it for most of CS systems, like a CS system for image processing with a conventional wavelet dictionary in which the SRE matrix is generally not available. Simulation results demonstrate the efficiency and effectiveness of the proposed approach compared with the state-of-the-art methods. In addition, we experimentally demonstrate with natural images that under similar compression rate, a CS system with a learned dictionary in high dimensions outperforms the one in low dimensions in terms of reconstruction accuracy. This together with the fact that our proposed method can efficiently work in high dimension suggests that a CS system can be potentially implemented beyond the small patches in sparsity-based image processing.

© 2017 Elsevier B.V. All rights reserved.

1. Introduction

Since the beginning of this century, Compressive Sensing or Compressed Sensing (CS) has received a great deal of attention [1–6]. Generally speaking, CS is a mathematical framework that addresses accurate recovery of a signal vector $\mathbf{x} \in \mathbb{R}^N$ from a set of linear measurements

$$\mathbf{y} = \Phi \mathbf{x} \in \mathbb{R}^M \quad (1)$$

where $M \ll N$ and $\Phi \in \mathbb{R}^{M \times N}$ is referred to as the projection or sensing matrix. CS has found many applications in the areas such as image processing, machine learning, pattern recognition, signal detection/classification etc. We refer the reader to [5,6] and the references therein to find the related topics mentioned above.

Sparsity and coherence are two important concepts in CS theory. We say a signal $\mathbf{x}$ of interest approximately sparse (in some basis or dictionary) if we can approximately express it as a linear combination of few columns (also called atoms) from a well-chosen dictionary:

$$\mathbf{x} = \Psi \boldsymbol{\theta} + \mathbf{e} \quad (2)$$

where $\Psi \in \mathbb{R}^{N \times L}$ is the given or determined dictionary, $\boldsymbol{\theta} \in \mathbb{R}^L$ is a sparse coefficient vector with few non-zero elements, and $\mathbf{e} \in \mathbb{R}^N$

stands for the sparse representation error (SRE). In particular, the vector $\mathbf{x}$ is called (purely or exactly) K -sparse in Ψ if $\|\boldsymbol{\theta}\|_0 = K$ and $\mathbf{e} = \mathbf{0}$ and approximately K -sparse in Ψ if $\|\boldsymbol{\theta}\|_0 = K$ and $\mathbf{e}$ has relatively small energy. Here, $\|\boldsymbol{\theta}\|_0$ denotes the number of non-zero elements in $\boldsymbol{\theta}$ and $\mathbf{0}$ represents a vector whose entries are equivalent to 0. Through this paper, we say $\boldsymbol{\theta}$ is K -sparse if $\|\boldsymbol{\theta}\|_0 = K$ regardless whether $\mathbf{e} = \mathbf{0}$.

Substituting the sparse model (2) of $\mathbf{x}$ into (1) gives

$$\mathbf{y} = \Phi \Psi \boldsymbol{\theta} + \Phi \mathbf{e} \triangleq \mathbf{D} \boldsymbol{\theta} + \boldsymbol{\epsilon} \quad (3)$$

where the matrix $\mathbf{D} = \Phi \Psi$ is referred to as the equivalent dictionary of the CS system and $\boldsymbol{\epsilon} \triangleq \Phi \mathbf{e}$ denotes the projection noise caused by SRE. The goal of a CS system is to retrieve $\boldsymbol{\theta}$ (and hence $\mathbf{x}$) from the measurements $\mathbf{y}$. Due to the fact that $M \ll L$, solving $\mathbf{y} \approx \mathbf{D} \boldsymbol{\theta}$ for $\boldsymbol{\theta}$ is an undetermined problem which has an infinite number of solutions. By utilizing the priori knowledge that $\boldsymbol{\theta}$ is sparse, a CS system typically attempts to recover $\boldsymbol{\theta}$ by solving the following problem:

$$\boldsymbol{\theta} = \arg \min_{\boldsymbol{\theta}} \|\tilde{\boldsymbol{\theta}}\|_0, \quad \text{s.t. } \|\mathbf{y} - \mathbf{D} \tilde{\boldsymbol{\theta}}\|_2 \leq \|\boldsymbol{\epsilon}\|_2 \quad (4)$$

which can be solved by many efficient numerical algorithms including basis pursuit (BP), orthogonal matching pursuit (OMP), least absolute shrinkage and selection operator (LASSO) etc. All of the methods can be found in [5,7] and the references therein.

To ensure exact recovery of $\boldsymbol{\theta}$ through (4), we need certain conditions on the equivalent dictionary $\mathbf{D}$. One of such conditions is

* Corresponding author.

E-mail addresses: hongtao@cs.technion.ac.il (T. Hong), zzhu@mines.edu (Z. Zhu).

related to the concept of *mutual coherence*. The mutual coherence of a matrix $\mathbf{D} \in \mathbb{R}^{M \times L}$ is denoted by

$$\mu(\mathbf{D}) \triangleq \max_{1 \leq i \neq j \leq L} |\tilde{\mathbf{G}}(i, j)| \quad (5)$$

where $\tilde{\mathbf{G}} = \tilde{\mathbf{D}}^T \tilde{\mathbf{D}}$ is called the Gram matrix of $\tilde{\mathbf{D}} = \mathbf{D}\mathbf{S}_{sc}$ with $\mathbf{S}_{sc}$ a diagonal scaling matrix such that each column of $\tilde{\mathbf{D}}$ is of unit length. Here T represents the transpose operator. It is known that $\mu(\mathbf{D})$ is lower bounded by the Welch bound $\underline{\mu}(\mathbf{D}) = \sqrt{\frac{L-M}{M(L-1)}}$, i.e., $\mu(\mathbf{D}) \in [\sqrt{\frac{L-M}{M(L-1)}}, 1]$. The mutual coherence $\mu(\mathbf{D})$ measures the worst-case coherence between any two columns of $\mathbf{D}$ and is one of the fundamental quantities associated with the CS theory. As shown in [5], when there is no projection noise (i.e., $\epsilon = 0$), any K -sparse signal $\boldsymbol{\theta}$ can be exactly recovered by solving the linear system (4) as long as

$$K < \frac{1}{2} \left[1 + \frac{1}{\mu(\mathbf{D})} \right] \quad (6)$$

which indicates that a smaller $\mu(\mathbf{D})$ ensures a CS system to recover the signal with a larger K . Thus, Barchiesi et al. and the following workers [8,9] proposed methods to design a dictionary with small mutual coherence. For a given dictionary Ψ , the mutual coherence of the equivalent dictionary is actually determined or controlled by the projection matrix Φ . So it would be of great interest to design Φ such that $\mu(\mathbf{D})$ is minimized. Another similar indicator used to evaluate the average performance of a CS system is named *average mutual coherence* μ_{av} . The definition of μ_{av} is given as follows:

$$\mu_{av}(\mathbf{D}) \triangleq \frac{\sum_{\forall (i,j) \in S_{av}} |\tilde{\mathbf{G}}(i, j)|}{N_{av}}$$

where $S_{av} \triangleq \{(i, j) : \tilde{\mu} \leq |\tilde{\mathbf{G}}(i, j)|\}$ with $0 \leq \tilde{\mu} < 1$ as a prescribed parameter and N_{av} is the number of components in the index set S_{av} .

There has been much effort [10–14] devoted to designing an optimal Φ that outperforms the widely used random matrix in terms of signal recovery accuracy (SRA). However, all these methods are based on the assumption that the signal is exactly sparse under a given dictionary, which is not true for practical applications. It is experimentally observed that the sensing matrix designed by Elad and the following workers [10–14] based on mutual coherence results in inferior performance for real images (which are generally approximately but not exactly sparse under a well-chosen dictionary). To address this issue, the recent work in [15,16] proposed novel methods to design a robust projection matrix when the SRE exists.¹ Through this paper, similar to what is used in [15,16], a robust projection (or sensing) matrix means it is designed with consideration of possible SRE and hence the corresponding CS system yields superior performance when the SRE $\mathbf{e}$ in (2) is not nil. However, the approaches in [15,16] need the explicit value of the SRE on the training dataset, making them inefficient in several aspects. First, many practical CS systems with predefined analytical dictionaries (e.g., the wavelet dictionary, and the modulated discrete prolate spheroidal sequences (DPSS) dictionary for sampled multiband signals [17]) actually do not involve any training dataset and hence no SRE available. In order to design the robust projection matrix for these CS systems using the framework presented in [15,16], one has to first construct plenty of extra representative dataset for the explicit SRE with the given dictionary, which limits the range of applications. Second, even for the CS system with a dictionary learned typically on a large-scale dataset, we need a lot of memories and computations to store and compute

with the huge dataset as well its corresponding SRE for designing a robust sensing matrix. Moreover, if the CS system is applied to a dynamic dataset, e.g., video stream, it is practically impossible to store all the data and compute its corresponding SRE. Therefore, the requirement of the explicit value of SRE for the training dataset makes the methods in [15,16] limited and inefficient for all the cases discussed above.

In this paper, to drop the requirement of the training dataset as well as its SRE, we propose a novel robust projection matrix framework only involving a predefined dictionary. With this new framework, we can efficiently design projection matrices for the CS systems mentioned above. We stress that by efficient method for robust projection matrix design (which is the title of this paper), we are not providing an efficient method for solving the problems in [15,16]; instead we provide a new framework in which the training dataset and its corresponding SRE are not required any more. Experiments on synthetic data and real images demonstrate the proposed sensing matrix yields a comparable performance in terms of SRA compared with the ones obtained by Li and Hong et al. [15,16].

Before proceeding, we first briefly introduce some notation used throughout the paper. MATLAB notations are adopted in this paper. In this connection, for a vector, $\mathbf{v}(k)$ denotes the k th component of $\mathbf{v}$. For a matrix, $\mathbf{Q}(i, j)$ means the (i, j) th element of matrix $\mathbf{Q}$, while $\mathbf{Q}(k, :)$ and $\mathbf{Q}(:, k)$ indicate the k th row and column vector of $\mathbf{Q}$, respectively. We use $\mathbf{I}$ and $\mathbf{I}_L$ to denote an identity matrix with arbitrary and $L \times L$ dimension, respectively. The k th column of $\mathbf{Q}$ is also denoted by $\mathbf{q}_k$. $\text{trace}(\mathbf{Q})$ denotes the calculation of the trace of $\mathbf{Q}$. The Frobenius norm of a given matrix $\mathbf{Q}$ is $\|\mathbf{Q}\|_F = \sqrt{\sum_{i,j} \|\mathbf{Q}(i, j)\|^2} = \sqrt{\text{trace}(\mathbf{Q}^T \mathbf{Q})}$ where T represents the transpose operator. The definition of l_p norm for a vector $\mathbf{v} \in \mathbb{R}^N$ is $\|\mathbf{v}\|_p \triangleq (\sum_{k=1}^N |\mathbf{v}(k)|^p)^{\frac{1}{p}}$, $p \geq 1$.

The remainder is arranged as follows. Some preliminaries are given in Section 2 to state the motivation of developing such a novel model. The proposed model which does not need the SRE is shown in Section 3 and the corresponding optimal sensing problem is solved in this section. The synthetic and real data experiments are carried out in Section 4 to demonstrate the efficiency and effectiveness of the proposed method. Some conclusions are given in Section 5 to end this paper.

2. Preliminaries

A sparsifying dictionary Ψ for a given dataset $\{\mathbf{x}_k\}_{k=1}^P$ is usually obtained by considering the following problem

$$\{\Psi, \boldsymbol{\theta}_k\} = \arg \min_{\Psi, \boldsymbol{\theta}_k} \sum_{k=1}^P \|\mathbf{x}_k - \tilde{\Psi} \tilde{\boldsymbol{\theta}}_k\|_2^2 \quad \text{s.t.} \quad \|\tilde{\boldsymbol{\theta}}_k\|_0 \leq K \quad (7)$$

which can be addressed by some practical algorithms [18], among which the popularly utilized are the K -singular value decomposition (K -SVD) algorithm [19] and the method of optimal direction (MOD) [20]. As stated in the previous section, the SRE $\mathbf{e}_k = \mathbf{x}_k - \Psi \boldsymbol{\theta}_k$ is generally not nil. We concatenate all the SRE $\{\mathbf{e}_k\}$ into an $N \times P$ matrix:

$$\mathbf{E} \triangleq \mathbf{X} - \Psi \boldsymbol{\Theta}$$

which is referred to as the SRE matrix corresponding to the training dataset $\{\mathbf{x}_k\}$ and the learned dictionary Ψ .

The recent work in [15,16] attempted to design a robust projection matrix with consideration of the SRE matrix $\mathbf{E}$ and proposed to solve

$$\Phi = \arg \min_{\Phi} \|\mathbf{I}_L - \Psi^T \Phi^T \Phi \Psi\|_F^2 + \lambda \|\Phi \mathbf{E}\|_F^2 \quad (8)$$

or

$$\Phi = \arg \min_{\Phi, \mathbf{G} \in \mathcal{H}_g} \|\mathbf{G} - \Psi^T \Phi^T \Phi \Psi\|_F^2 + \lambda \|\Phi \mathbf{E}\|_F^2 \quad (9)$$

¹ We note that the approaches considered in [15,16] share the same framework. The difference is that in [16] Hong et al. utilized an efficient iterative algorithm giving an approximate solution, while a closed form solution is derived in [15].

Download English Version:

<https://daneshyari.com/en/article/4977429>

Download Persian Version:

<https://daneshyari.com/article/4977429>

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