

Unit-Norm Tight Frame-based Sparse Representation With Application to Speech Inpainting

Huang Bai¹, Chuanrong Hong², Sheng Li³, Yimin D. Zhang⁴, Xiumei Li^{1*}

¹*School of Information Science and Technology, Hangzhou Normal University, Hangzhou, 311121, Zhejiang, China*

²*Alibaba Group, Hangzhou, 311121, Zhejiang, China*

³*College of Information Engineering, Zhejiang University of Technology, Hangzhou, 310023, Zhejiang, China*

⁴*Department of Electrical and Computer Engineering, Temple University, Philadelphia, PA 19312, USA*

Abstract

This paper deals with the unit-norm tight frames (UNTFs) that better represent and solve sparse representation problems. Under this sparse representation framework, a novel model termed tight dictionary learning (TDL) is proposed. Unlike those dictionaries that only focus on the signals of interest, TDL intends to learn a more comprehensive dictionary that better represents the signals with the UNTF properties to facilitate sparse recovery. To better describe the physical meanings, the normalization constraints on the dictionary columns are imposed. A gradient-based approach is developed to obtain the tight dictionary where the normalization constraints are embedded into the cost function by a parametrization method. The learned tight dictionary is applied to the speech inpainting problem to restore speech signals corrupted by data missing. In order to compensate the degradation effect of the data missing, a UNTF-based preconditioning technique is developed to reform the frame properties of the degraded system. Parametrization and gradient-based approaches similar to those applied in the TDL are adopted to design a tight preconditioner. Extensive experiments on speech denoising and speech segment restoration are conducted to demonstrate the superiority of the proposed schemes.

Keywords:

Sparse representation, dictionary learning, unit-norm tight frame, speech inpainting, preconditioning.

1. Introduction

Signal representation intends to transform interested signals to a different domain via a basis, where certain signal characteristics become more readily apparent in the transform coefficients for facilitating various signal processing tasks. The notion of basis in finite-dimensional spaces implies that the number of representative vectors is the same as the dimension of the space, which means that the representation is nonredundant and, as a result, corruption or loss of transform coefficients may cause serious consequence. In order to provide a robust transformation, redundant signal representation is usually adopted. The redundant counterpart of a basis is called a *frame* [1], [2].

Frames give stable signal representations and allow modelling for noisy environments [1]. Signal representation using a frame expresses a signal vector $\mathbf{x} \in \mathbb{R}^{N \times 1}$ in the form of

$$\mathbf{x} \approx \sum_{k=1}^K s(k) \mathbf{D}(:, k) \triangleq \mathbf{D} \mathbf{s}, \quad (1)$$

where $\mathbf{D} \in \mathbb{R}^{N \times K}$ ($N < K$) is a frame matrix with its columns $\{\mathbf{D}(:, k)\}$ referred to as atoms, and $\mathbf{s} \in \mathbb{R}^{K \times 1}$ is the corresponding coefficient vector. If $\mathbf{s}$ is sparse, i.e., there are few non-zero elements in $\mathbf{s}$, equation (1) is specified as a *sparse representation* (SR) of $\mathbf{x}$, and the frame matrix $\mathbf{D}$ is named as a *dictionary* [3], [4].

Signal SR theory is an evolving theory with state-of-the-art results applied in many signal processing tasks, such as denoising, video encryption, speech recognition, and compressed sensing [5]-[8]. Frames are less constrained than bases and thus enable more flexible representations of corresponding signals [1], [2]. Hence, investigation of frame-based SR is of great importance.

A fundamental requirement in employing the SR theory is the proper choice of the dictionary and this leads to the well-known *dictionary learning* (DL) problem [9]. A great volume of works have been performed to learn a data-adaptive dictionary so that a particular class of signals can be sparsely represented in this dictionary with a low approximation error [4], [10], [11].

Let $\{\mathbf{x}_l\}_{l=1}^L$ be a set of training signals. The basic problem of DL is to find a dictionary $\mathbf{D}$ such that for each $\mathbf{x}_l$ there exists a sparse representation vector $\mathbf{s}_l$. Denote

$$\mathbf{X} \triangleq \begin{bmatrix} \mathbf{x}_1 & \mathbf{x}_2 & \cdots & \mathbf{x}_L \end{bmatrix} \in \mathbb{R}^{N \times L}, \quad \mathbf{S} \triangleq \begin{bmatrix} \mathbf{s}_1 & \mathbf{s}_2 & \cdots & \mathbf{s}_L \end{bmatrix} \in \mathbb{R}^{K \times L}.$$

The general DL problem is formulated as

$$\min_{\mathbf{D}, \mathbf{S}} \|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2 + f(\mathbf{D}) + f(\mathbf{S}), \quad (2)$$

where $\|\cdot\|_F$ denotes the *Frobenius* norm, and $f(\mathbf{D})$ and $f(\mathbf{S})$ are some regularization functions of $\mathbf{D}$ and $\mathbf{S}$, respectively. In general, $f(\mathbf{D})$ may be subject to the matrix column-normalization constraint, whereas $f(\mathbf{S})$ is set as the sparsity constraint [4]. It is a common assumption that $f(\mathbf{D})$ has nothing to do with $\mathbf{S}$ and $f(\mathbf{S})$ is also independent of $\mathbf{D}$.

Assume $\mathbf{v}(k)$ as the k -th element of a vector $\mathbf{v} \in \mathbb{R}^{K \times 1}$. The ℓ_p -norm of vector $\mathbf{v}$ is defined as

$$\|\mathbf{v}\|_p \triangleq \left(\sum_{k=1}^K |\mathbf{v}(k)|^p \right)^{1/p}, \quad p \geq 1.$$

For convenience, $\|\mathbf{v}\|_0$ is used to denote the number of non-zero elements in $\mathbf{v}$.

A common choice of $f(\mathbf{S})$ in (2) is $f(\mathbf{S}) = \sum_{l=1}^L \tau_l \|\mathbf{s}_l\|_0$ (throughout this paper, this $f(\mathbf{S})$ will be retained) with $\{\tau_l\}$ denoting proper weighting constants [10], [11]. This leads to

$$\min_{\mathbf{D}, \{\mathbf{s}_l\}_{l=1}^L} \|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2 + f(\mathbf{D}) + \sum_{l=1}^L \tau_l \|\mathbf{s}_l\|_0. \quad (3)$$

Such a problem is difficult to solve as it is non-convex in $\mathbf{D}$ and $\mathbf{S}$, and $\|\cdot\|_0$ is non-smooth and highly unstable. A popularly used approach is based on the alternating minimization strategy [4], [10], [11]. A two-stage procedure is usually carried out for solving the above problem and also avoiding the selections of $\{\tau_l\}$. The first stage is referred to as *sparse coding* or *sparse decomposition*, also known as *sparse recovery* in signal reconstruction, aiming at finding the (column) sparse matrix $\mathbf{S}$ with a given dictionary matrix $\mathbf{D}$ through the following inverse problem

$$\hat{\mathbf{s}}_l \triangleq \arg \min_{\mathbf{s}_l} \|\mathbf{x}_l - \mathbf{D}\mathbf{s}_l\|_2^2 \quad \text{s.t.} \quad \|\mathbf{s}_l\|_0 \leq \kappa, \quad \forall l, \quad (4)$$

with κ denoting the sparsity level. The roles of the penalty and the constraints in (4) may be reversed [4]. We can choose to constrain the SR error and obtain the sparsest representation by

$$\hat{\mathbf{s}}_l \triangleq \arg \min_{\mathbf{s}_l} \|\mathbf{s}_l\|_0 \quad \text{s.t.} \quad \|\mathbf{x}_l - \mathbf{D}\mathbf{s}_l\|_2^2 \leq \epsilon, \quad \forall l, \quad (5)$$

where ϵ denotes the error threshold. Such a problem can be solved using the orthogonal matching pursuit (OMP) based methods [12]-[14]. Furthermore, a popular strategy is to replace the ℓ_0 -norm by the ℓ_1 -norm, which is its convex approximation in a natural sense, leading to

$$\hat{\mathbf{s}}_l \triangleq \arg \min_{\mathbf{s}_l} \|\mathbf{s}_l\|_1 \quad \text{s.t.} \quad \|\mathbf{x}_l - \mathbf{D}\mathbf{s}_l\|_2^2 \leq \epsilon, \quad \forall l. \quad (6)$$

Equation (6) is convex and can be solved using algorithms such as the basis pursuit (BP) [15] and the ℓ_1 -based optimization techniques [16].

Many methods that solve (3) differentiate each other mainly in the second stage, namely, *dictionary updating*. In other words, we need to determine what kind of frames should be selected for coding the signals of interest more sparsely and precisely. In classical DL problems with normalization constraints on dictionary columns $\{\mathbf{D}(:, k)\}$, i.e., $f(\mathbf{D}) = \sum_{k=1}^K (\|\mathbf{D}(:, k)\|_2^2 - 1)^2$ in (3), the dictionary is learned adaptively based on the training signals [11]. With such a strategy, some inherent properties of the dictionary, e.g., the tightness property that is concerned in this paper, may lose. It is desirable to learn a dictionary with a desired representation capability for signals as well as possessing certain excellent frame properties to facilitate sparse coding.

For a linear system satisfying expression (1), matrix $\mathbf{D}$ is called the *system matrix*. References [17] and [18] investigated the optimized design of the system matrix that minimizes the mean squared error (MSE) of the oracle estimator [19]. The oracle MSE coincides with the unbiased Cramér-Rao bound for sparse deterministic vectors and represents the best achievable performance for any unbiased estimator [20]. Besides, the oracle MSE performance acts as a benchmark to the performance of various sparse recovery algorithms. For example, it has been demonstrated both theoretically and numerically that the Dantzig selector, the OMP and thresholding algorithms all achieve performance that is proportional to the oracle MSE [20]. Since minimizing the MSE of the oracle estimator with respect to the system matrix is a non-convex problem, Chen *et al.* adopted convex relaxation techniques and proved that a *unit-norm*

tight frame (UNTF) is the closest design in the Frobenius norm sense to the solution of the relaxed problem [17], [18]. UNTFs offer the advantage of redundancy in signal representations and numerical stability of reconstruction and, therefore, have received great interests in recent years in the signal processing community [21]-[25].

A UNTF is viewed as a generalization of an orthonormal basis which brings redundancy and stability [22]. Generally speaking, the properties of a UNTF refer to the *unit-norm* (i.e., normalization) property of the atoms and the *tightness* property of the frame. The unit-norm property avoids the explosion or degeneration of the frame elements and the sparse coefficients, and the tightness property is valuable to ensure fast convergence of signal sparse decompositions [22], [26]. Besides, according to [23], the properties of the UNTF have been proved valid against additive noise and erasures.

The main objective of this paper is to develop a *tight dictionary learning* (TDL) method to learn a dictionary that minimizes the SR error and approximates to a UNTF. In the sequel, tight dictionary implies that both the unit-norm property and the tightness property hold. As speech signals satisfy the SR model with a high probability [27]-[30] and UNTFs show great superiority in sparse signal processing [21]-[25], the learned tight dictionary can be applied to speech inpainting applications. A *preconditioning* technique [31]-[33] based on UNTF is further presented to improve the inpainting performance. Utilizing the properties of the UNTF to design the system matrix to facilitate sparse recovery is termed as *tight sparse representation* (TSR).

The contributions of this paper are four-fold:

- A novel model is proposed for learning a tight dictionary. The inherent properties of the dictionary matrix, including the unit-norm property and the tightness property, are taken into consideration when minimizing the SR error. Because of the UNTF properties that facilitate the sparse recovery, the learned dictionary can well represent the signals of interest.
- An alternating minimization strategy is developed to solve the TDL problem. When updating the dictionary, a parametrization method is carried out to convert the constrained optimization problem to an unconstrained one. A gradient-based approach is developed to solve the parameterized problem and obtain the tight dictionary. Such an updating procedure results in a precise solution of the TDL problem.
- When applied to speech inpainting problems, the impairment of data missing is effectively mitigated. The data missing process is formulated as a degradation matrix that removes samples from the speech signals. Noting the sparse structure of the speech signals in the SR model, the learned tight dictionary is applied to effectively restore the missing data.
- In order to mitigating the degradation effect of data missing, the preconditioning technique based on UNTF is employed. Similar parametrization and gradient-based approaches as in the TDL are adopted to design the tight preconditioner. Such a preconditioner highly reforms the UNTF properties of the degraded system to facilitate the sparse recovery, thereby improving the accuracy of the speech inpainting results.

In addition, extensive experiments on speech data and real speech segments are carried out to demonstrate the superior performance of the proposed TSR schemes.

The remainder of this paper is arranged as follows. In Section 2, some preliminaries are provided and two related works are introduced as the baselines to evaluate our method. A novel model for learning the tight dictionary with desirable representation capability is described in Section 3. In Section 4, application of the SR theory and the preconditioning technique to speech inpainting is discussed in detail. Experiments on speech signals are carried out in Section 5 to examine the performance of the proposed schemes and to compare with existing techniques. Concluding remarks are given in Section 6.

In this paper, we use italic characters to denote scalars, whereas lowercase and uppercase bold italic characters indicate vectors and matrices, respectively. All parameters are real-valued. Besides, MATLAB notations are used to denote the selections of elements in a vector or a matrix.

2. Preliminaries and related works

In this section, we review some preliminaries on frames, summarize two important works related to TDL, and point out a problem that is usually encountered in DL.

2.1. Brief review on frames

A finite frame $\mathbf{D} \in \mathfrak{R}^{N \times K}$ in the Hilbert space $\mathfrak{R}^{N \times 1}$ is a set of $K \geq N$ vectors $\{\mathbf{D}(:, k)\}_{k=1}^K$ that satisfies the following generalized Parseval condition:

$$\alpha \|\mathbf{v}\|_2^2 \leq \sum_{k=1}^K |\langle \mathbf{D}(:, k), \mathbf{v} \rangle|^2 \leq \beta \|\mathbf{v}\|_2^2 \quad (7)$$

for all $\mathbf{v} \in \mathfrak{R}^{N \times 1}$, where α and β are two positive constants referred to as the lower and upper bounds, respectively [1], [2]. If it is possible to have $\alpha = \beta$ in (7), then we achieve an α -tight frame with α being referred to as the tightness constant in this case. A frame $\{\mathbf{D}(:, k)\}_{k=1}^K$ is α -tight if and only if the corresponding matrix $\mathbf{D}$ has a singular value decomposition (SVD) in the following form

$$\mathbf{D} = \mathbf{U} \begin{bmatrix} \sqrt{\alpha} \mathbf{I}_N & \mathbf{0} \end{bmatrix} \mathbf{V}^T, \quad (8)$$

where superscript $\mathcal{T}$ denotes the transpose operator, $\mathbf{I}_N$ is the $N \times N$ identity matrix, and both $\mathbf{U} \in \mathfrak{R}^{N \times N}$ and $\mathbf{V} \in \mathfrak{R}^{K \times K}$ are orthonormal matrices [34], [35].

When designing a frame for a specific application, it is important to control over the spectrum of the corresponding frame matrix [21], [22]. This spectral characteristic is usually studied via the condition number of the frame matrix. The closer the condition number is to 1, the better conditioned the frame matrix will be [36]. A linear system with a

well-conditioned system matrix is said to be stable. It is clear that the condition number of a matrix expressed in (8) is 1, which means that a tight frame forms a well-conditioned frame matrix.

Where all the atoms have the same unit ℓ_2 -norm, those frames are called unit-norm frames, i.e., $\|\mathbf{D}(:, k)\|_2^2 = 1$ for all k . A frame $\mathbf{D}$ is named UNTF when both tightness property and unit-norm property hold [22]-[26]. According to the definitions of the tight frame and the unit-norm frame, it is straightforward to obtain

$$\alpha N = \text{tr}[\mathbf{D}\mathbf{D}^T] = \text{tr}[\mathbf{D}^T\mathbf{D}] = K \quad \Rightarrow \quad \alpha = \frac{K}{N}, \quad (9)$$

by appealing to the cyclic property of the trace operation $\text{tr}[\cdot]$ [36]. It means that the tightness constant of any UNTF is K/N , i.e., the redundancy of corresponding frame.

2.2. Related works of TDL

The two DL schemes proposed in [37] and [38] are closely related to the work presented in this paper, and we briefly introduce them in this part.

Based on the general DL problem (3) and the alternating minimization strategy to solve it, we reformulate (3) as:

$$\min_{\mathbf{D}, \{s_l\}_{l=1}^L} \|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2 + f(\mathbf{D}) \quad \text{s.t.} \quad \|s_l\|_0 \leq \kappa, \quad \forall l. \quad (10)$$

In [37], the regularization term $f(\mathbf{D})$ in (10) is set as

$$f(\mathbf{D}) = \omega \|\mathbf{D}^T\mathbf{D} - \mathbf{I}_K\|_F^2, \quad (11)$$

with ω being the regularization parameter. Substituting (11) into (10) renders the following DL model:

$$\min_{\mathbf{D}, \{s_l\}_{l=1}^L} \|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2 + \omega \|\mathbf{D}^T\mathbf{D} - \mathbf{I}_K\|_F^2 \quad \text{s.t.} \quad \|s_l\|_0 \leq \kappa, \quad \forall l. \quad (12)$$

By employing the alternating minimization strategy, $\{s_l\}_{l=1}^L$ are computed by certain sparse recovery algorithm like OMP, and (12) turns to

$$\min_{\mathbf{D}} \|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2 + \omega \|\mathbf{D}^T\mathbf{D} - \mathbf{I}_K\|_F^2. \quad (13)$$

A gradient-based approach is then carried out for solving (13) [37].

In order to avoid the choice of the regularization parameter ω , [38] modifies the problem (13) to

$$\min_{\mathbf{D}} \|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2 \quad \text{s.t.} \quad \min_{\mathbf{D}} \|\mathbf{D}^T\mathbf{D} - \mathbf{I}_K\|_F^2. \quad (14)$$

That is, they set the regularization term $f(\mathbf{D})$ indicated in (11) as the constraint. The closed-form solution set of $\{\min_{\mathbf{D}} f(\mathbf{D})\}$ was derived in [34] as:

$$\mathbf{D} = \mathbf{U} \begin{bmatrix} \mathbf{I}_N & \mathbf{0} \end{bmatrix} \mathbf{V}^T, \quad (15)$$

where U and V are orthonormal matrices of dimensions $N \times N$ and $K \times K$, respectively. These two matrices U and V are then alternately optimized to minimize the SR error and improve the representation capability of the dictionary [38].

Comparing (15) and (8), it is clear that the result of $\{\min_{\mathbf{D}} f(\mathbf{D})\}$ with $f(\mathbf{D})$ indicated in (11) is exactly the form of the 1-tight frame which is closely related to TDL.

2.3. A shortcoming in general DL

As introduced in Section 1, classical DL problems take the normalization constraints on dictionary columns as the regularization term $f(\mathbf{D})$ because these constraints can avoid the explosion or degeneration of the matrix elements and the sparse coefficients. This is a common choice almost in all DL literatures. For works like [37] and [38], more constraints on the dictionary are made and these models, including TDL, are sometimes called constrained DL.

Most works take two steps to update the dictionary with normalization constraints. First, compute a dictionary without considering the normalization constraints, e.g., (13) and (14), and then execute an extra normalization step to the half-finished dictionary [8]-[10], [37], [38]. However, such a separated strategy compromises the optimality of the solution as those two steps are not compatible. It is important to maintain the normalization throughout the learning of the dictionary.

From the perspective of designing frames, these facts motivate us to develop a novel model with tightness and unit-norm properties taken into consideration simultaneously.

3. Tight dictionary learning model and algorithm

In this section, a novel model for TDL is proposed which minimizes the SR error and, at the same time, approximates a UNTF that maintains the tightness and unit-norm properties. Parametrization and gradient-based approaches are developed to solve the corresponding design problem.

3.1. Tight dictionary learning model

According to *Proposition 2* in [17], any $N \times K$ UNTF is a solution to the optimization problem

$$\min_{\mathbf{D}} \|\mathbf{D}^T \mathbf{D} - \mathbf{I}_K\|_F^2 \quad \text{s.t.} \quad \|\mathbf{D}(:,k)\|_2^2 = 1, \quad \forall k. \quad (16)$$

As indicated in *Theorem 1* of [34], the closed-form solution set of (16) can be derived as

$$\mathbf{D} = \mathbf{U} \begin{bmatrix} \sqrt{\frac{K}{N}} \mathbf{I}_N & \mathbf{0} \end{bmatrix} \mathbf{V}^T, \quad (17)$$

with the orthonormal matrix $\mathbf{V}$ properly designed to ensure the normalization constraints in (16) to hold. As shown in [34] and [39], such a $\mathbf{V}$ can always be achieved. The result (17), i.e., the solution of (16), is a K/N -tight frame with unit-norm atoms, which is the exact definition of an $N \times K$ UNTF [22], [26]. This is consistent with that indicated in

(8) and (9). For convenience, the term tight dictionary is used to represent a dictionary possessing the properties of the UNTF that ensure tightness and unit-norm.

Recalling the definition of DL, the capability of accurately representing the dictionary to signals of interest must be enforced. Therefore, in addition to the inherent properties of the dictionary, a signal-adapted penalty term is also needed. Based on the discussions above, we propose the following TDL model:

$$\min_{\mathbf{D}, \{s_l\}_{l=1}^L} \|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2 + \omega \|\mathbf{D}^T \mathbf{D} - \mathbf{I}_K\|_F^2 \quad \text{s.t.} \quad \|s_l\|_0 \leq \kappa, \forall l, \quad \|\mathbf{D}(:, k)\|_2^2 = 1, \forall k, \quad (18)$$

where ω trades off the significance between the representation accuracy and the frame inherent properties. As the identity matrix has the smallest condition number, model (18) will lead to a well-conditioned dictionary with a flat spectrum. Such a dictionary results in a stable representation system [31], [36], [37].

Remark 1.

- The main difference between (18) and (12) lies in the dictionary column-normalization constraints. These constraints are set for two reasons. On one hand, it gives better description of the physical meaning to hold the unit-norm. As shown in [34], without the normalization constraints, the solution to

$$\arg \min_{\mathbf{D}} \|\mathbf{D}^T \mathbf{D} - \mathbf{I}_K\|_F^2 = \mathbf{U} \begin{bmatrix} \mathbf{I}_N & \mathbf{0} \end{bmatrix} \mathbf{V}^T,$$

is a 1-tight frame but is not necessarily a UNTF. On the other hand, it avoids the degeneration of the sparse coefficients. In equation (1), if some of the elements in $\mathbf{D}$ become very large, small valid values in $\mathbf{s}$ may be easily neglected.

- Frame potential is a famous concept related to the UNTFs and has been defined explicitly in [22] as:

$$\text{FP}(\mathbf{D}) \triangleq \sum_{i=1}^K \sum_{j=1}^K |\langle \mathbf{D}(:, i), \mathbf{D}(:, j) \rangle|^2 = \|\mathbf{D}^T \mathbf{D}\|_F^2.$$

A number of methods were developed to construct UNTFs by minimizing the frame potential [25], [40], [41]. It can be observed that the regularization term $\{\|\mathbf{D}^T \mathbf{D} - \mathbf{I}_K\|_F^2\}$ is closely associated with the frame potential if the unit-norm property holds, i.e., $\text{FP}(\mathbf{D}) = \|\mathbf{D}^T \mathbf{D} - \mathbf{I}_K\|_F^2 + K$ when $\forall k, \|\mathbf{D}(:, k)\|_2^2 = 1$. This fact once again verifies the importance of the normalization constraints.

- Equation (17) is a closed-form expression of the UNTF when a proper $\mathbf{V}$ is set [34]. We do not constrain the form of the dictionary as (17) because there only exists one orthonormal matrix $\mathbf{U}$ to improve the representation capability. The impact of the sole orthonormal matrix on the SR capability is limited. The weighting model described in (18) is flexible to learn a desired dictionary by selecting a proper value of ω .

3.2. Tight dictionary learning algorithm

This part is devoted to the detailed algorithm description for solving the TDL problem. Noting that, because (18) is still non-convex in $\mathbf{D}$ and $\mathbf{S}$, and $\|\cdot\|_0$ is non-smooth and highly unstable, the alternating minimization strategy is employed. The pseudo-code is given in **Algorithm 1**.

Algorithm 1: Tight dictionary learning algorithm

Initialization:

$\mathbf{D}^{[0]} \in \mathbb{R}^{N \times K}$: initial column-normalized dictionary;

$\mathbf{X} = \{\mathbf{x}_l\}_{l=1}^L \in \mathbb{R}^{N \times L}$: training signal set;

κ : sparsity level;

N_{iter} : number of iterations for DL.

Set $i = 1$.

Begin: For $i = 1 : N_{\text{iter}}$, repeat the following two stages:

- *Sparse coding*: With $\mathbf{D}^{[i-1]}$, update the sparse coefficient set $\mathbf{S}^{[i]} = \{\hat{\mathbf{s}}_l\}_{l=1}^L$ by

$$\hat{\mathbf{s}}_l = \arg \min_{\mathbf{s}_l} \|\mathbf{x}_l - \mathbf{D}^{[i-1]} \mathbf{s}_l\|_2^2 \quad \text{s.t.} \quad \|\mathbf{s}_l\|_0 \leq \kappa, \quad \forall l \quad (19)$$

using a certain OMP-based algorithm.

- *Dictionary updating*: For $\mathbf{S}^{[i]}$ fixed, update the dictionary with

$$\mathbf{D}^{[i]} = \arg \min_{\mathbf{D}} \|\mathbf{X} - \mathbf{D} \mathbf{S}^{[i]}\|_F^2 + \omega \|\mathbf{D}^T \mathbf{D} - \mathbf{I}_K\|_F^2 \quad \text{s.t.} \quad \|\mathbf{D}(:, k)\|_2^2 = 1, \quad \forall k. \quad (20)$$

Because of the normalization constraints in (20), the closed-form solution is hard to be derived. A parametrization method will be adopted to embed the normalization constraints on the dictionary columns into the cost function, and then a gradient-based approach is developed to solve the parameterized problem for the dictionary. This will be described in detail below.

End: End the algorithm, and output the tight dictionary.

As (19) is solved by an OMP-based algorithm, the remaining issue is to update the dictionary, i.e., solving (20), which is described below. For notational convenience, the iteration index $[i]$ will be omitted in the following.

Denote $\varrho(\mathbf{D}) \triangleq \|\mathbf{E}_1\|_F^2 + \omega \|\mathbf{E}_2\|_F^2$ with

$$\mathbf{E}_1 \triangleq \mathbf{X} - \mathbf{D} \mathbf{S}, \quad \mathbf{E}_2 \triangleq \mathbf{D}^T \mathbf{D} - \mathbf{I}_K. \quad (21)$$

The minimization problem in (20) turns to

$$\min_{\mathbf{D}} \varrho(\mathbf{D}) \quad \text{s.t.} \quad \|\mathbf{D}(:, k)\|_2^2 = 1, \quad \forall k. \quad (22)$$

It is difficult to derive the closed-form solution of (22) due to the normalization constraints. Referring to [42], we present a gradient-based approach to solve (22). A key step of this approach is the following parametrization of the dictionary

$$\mathbf{D} \triangleq \mathbf{Q} \mathbf{Q}_d, \quad (23)$$

where $\mathbf{Q}$ is an $N \times K$ matrix with a mild condition that it contains no zero columns, and $\mathbf{Q}_d \in \mathbb{R}^{K \times K}$ is a diagonal matrix formed as

$$\mathbf{Q}_d = \text{diag}(\|\mathbf{Q}(:, 1)\|_2^{-1}, \|\mathbf{Q}(:, 2)\|_2^{-1}, \dots, \|\mathbf{Q}(:, K)\|_2^{-1}). \quad (24)$$

It is clear that the columns of $\mathbf{D}$ defined in (23) are inherently ℓ_2 -normalized as $\mathbf{Q}$ has no zero columns.

With the definitions of (23) and (24), the dictionary updating problem (22) is converted to the following unconstrained minimization problem

$$\min_{\mathbf{Q}} \varrho(\mathbf{Q}\mathbf{Q}_d). \quad (25)$$

(25) can be solved using the following gradient descent procedure:

$$\mathbf{Q}_n = \mathbf{Q}_{n-1} - \gamma \left. \frac{\partial \varrho}{\partial \mathbf{Q}} \right|_{\mathbf{Q}=\mathbf{Q}_{n-1}}, \quad (26)$$

where γ is the step size, and $\frac{\partial \varrho}{\partial \mathbf{Q}}$ is the gradient of $\varrho(\mathbf{Q}\mathbf{Q}_d)$ with respect to $\mathbf{Q}$.

Define

$$\varrho_1(\mathbf{Q}) \triangleq \|\mathbf{E}_1\|_F^2, \quad \varrho_2(\mathbf{Q}) \triangleq \|\mathbf{E}_2\|_F^2 \quad (27)$$

with $\mathbf{E}_1$ and $\mathbf{E}_2$ defined in (21). Then, one has

$$\varrho = \varrho_1(\mathbf{Q}) + \omega \varrho_2(\mathbf{Q}). \quad (28)$$

Hence

$$\frac{\partial \varrho}{\partial \mathbf{Q}} = \frac{\partial \varrho_1(\mathbf{Q})}{\partial \mathbf{Q}} + \omega \frac{\partial \varrho_2(\mathbf{Q})}{\partial \mathbf{Q}}. \quad (29)$$

The expressions of $\frac{\partial \varrho_1(\mathbf{Q})}{\partial \mathbf{Q}}$ and $\frac{\partial \varrho_2(\mathbf{Q})}{\partial \mathbf{Q}}$ can be obtained by the following Theorem 1 and Theorem 2, respectively.

Theorem 1. Let $\mathbf{Q} \in \mathbb{R}^{N \times K}$ be a matrix with no zero columns and $\mathbf{Q}_d$ is defined in (24). Denote

$$\mathbf{E}_1 \triangleq \mathbf{X} - (\mathbf{Q}\mathbf{Q}_d)\mathbf{S}. \quad (30)$$

The gradient of $\|\mathbf{E}_1\|_F^2$ with respect to $\mathbf{Q}$ is obtained as

$$\frac{\partial(\|\mathbf{E}_1\|_F^2)}{\partial \mathbf{Q}} = 2\mathbf{Q}\mathbf{Q}_d^3\mathbf{R}_1 - 2\mathbf{E}_1\mathbf{S}^T\mathbf{Q}_d, \quad (31)$$

where

$$\mathbf{R}_1 \triangleq \text{diag}(\mathbf{H}_1(1, 1), \mathbf{H}_1(2, 2), \dots, \mathbf{H}_1(K, K)), \quad \mathbf{H}_1 \triangleq \mathbf{S}\mathbf{E}_1^T\mathbf{Q}. \quad (32)$$

Theorem 2. Let $\mathbf{Q} \in \mathbb{R}^{N \times K}$ be a matrix with no zero columns and $\mathbf{Q}_d$ is defined in (24). Denote

$$\mathbf{E}_2 \triangleq (\mathbf{Q}\mathbf{Q}_d)^T(\mathbf{Q}\mathbf{Q}_d) - \mathbf{I}_K. \quad (33)$$

The gradient of $\|\mathbf{E}_2\|_F^2$ with respect to $\mathbf{Q}$ is obtained as

$$\frac{\partial(\|\mathbf{E}_2\|_F^2)}{\partial \mathbf{Q}} = 2(\mathbf{Q}\mathbf{Q}_d)(\mathbf{E}_2^T + \mathbf{E}_2)\mathbf{Q}_d - 2\mathbf{Q}\mathbf{Q}_d^3\mathbf{R}_2, \quad (34)$$

where

$$\mathbf{R}_2 \triangleq \text{diag}(\mathbf{H}_2(1, 1), \mathbf{H}_2(2, 2), \dots, \mathbf{H}_2(K, K)), \quad \mathbf{H}_2 \triangleq \mathbf{Q}^T(\mathbf{Q}\mathbf{Q}_d)(\mathbf{E}_2^T + \mathbf{E}_2). \quad (35)$$

The proofs of the above two theorems can be found in Appendices A and B, respectively.

Based on Theorems 1 and 2, the expression of $\frac{\partial \mathcal{Q}}{\partial \mathbf{Q}}$ is obtained as

$$\frac{\partial \mathcal{Q}}{\partial \mathbf{Q}} = \frac{\partial \mathcal{Q}_1(\mathbf{Q})}{\partial \mathbf{Q}} + \omega \frac{\partial \mathcal{Q}_2(\mathbf{Q})}{\partial \mathbf{Q}} = 2\mathbf{Q}\mathbf{Q}_d^3\mathbf{R}_1 - 2\mathbf{E}_1\mathbf{S}^\mathcal{T}\mathbf{Q}_d + \omega \left[2(\mathbf{Q}\mathbf{Q}_d)(\mathbf{E}_2^\mathcal{T} + \mathbf{E}_2)\mathbf{Q}_d - 2\mathbf{Q}\mathbf{Q}_d^3\mathbf{R}_2 \right]. \quad (36)$$

Using the result of (36), $\mathbf{Q}$ is updated via (26) by the gradient descent procedure, and the solution of (22) is obtained as $\mathbf{D} = \mathbf{Q}\mathbf{Q}_d$.

Remark 2.

- When $\omega \rightarrow +\infty$, the dictionary updating turns to the UNTF construction problem which is an attractive topic [25], [40], [41]. With the Lagrange multiplier, the problem (20) becomes

$$\min_{\mathbf{D}} \|\mathbf{D}^\mathcal{T}\mathbf{D} - \mathbf{I}_K\|_F^2 + \sum_{k=1}^K \lambda_k (\|\mathbf{D}(:, k)\|_2^2 - 1).$$

The above expression is also known as the famous Paulsen problem [26] with the two penalty terms concerning the tightness property and the unit-norm property, respectively. When dealing with multiple objects, some works alternately optimize these properties until the cost function converges (sometimes the alternating optimization method may not converge) [32], while we search for the tight frame with its atoms inherently normalized.

- For the proposed parametrization and gradient-based strategy, other penalty terms that subject to the unit-norm constraints are convenient to be added to the cost function, e.g., the signal-adapted penalty term $\{\|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2\}$ considered in this work. The computations of the gradients such as (31) and (34) are not complex. This strategy can be conveniently extended to design high dimensional matrices because it does not involve computationally expensive operations such as matrix inversion or SVD.
- It can be seen from the pseudo-code that **Algorithm 1** is carried out with respect to $\mathbf{D}$ and $\mathbf{S}$. The convergence of our algorithm can be easily guaranteed. For sparse coding, the popularly adopted OMP is viewed as an efficient method to solve (19) [12]-[14]. Therefore, in the i -th iteration,

$$\|\mathbf{x}_l - \mathbf{D}^{[i-1]}\hat{\mathbf{s}}_l\|_2^2 \leq \|\mathbf{x}_l - \mathbf{D}^{[i-1]}\mathbf{s}_l\|_2^2, \quad \forall l = 1, 2, \dots, L,$$

for any $\{\mathbf{s}_l\}$. Hence, with $\mathbf{S}^{[i]} = \{\hat{\mathbf{s}}_l\}_{l=1}^L$ and $\mathbf{S}^{[i-1]}$ being the result obtained in the previous iteration,

$$\|\mathbf{X} - \mathbf{D}^{[i-1]}\mathbf{S}^{[i]}\|_F^2 \leq \|\mathbf{X} - \mathbf{D}^{[i-1]}\mathbf{S}^{[i-1]}\|_F^2$$

can be ensured. This leads to

$$\|\mathbf{X} - \mathbf{D}^{[i-1]}\mathbf{S}^{[i]}\|_F^2 + \omega \|(\mathbf{D}^{[i-1]})^\mathcal{T}(\mathbf{D}^{[i-1]}) - \mathbf{I}_K\|_F^2 \leq \|\mathbf{X} - \mathbf{D}^{[i-1]}\mathbf{S}^{[i-1]}\|_F^2 + \omega \|(\mathbf{D}^{[i-1]})^\mathcal{T}(\mathbf{D}^{[i-1]}) - \mathbf{I}_K\|_F^2.$$

Considering dictionary updating, i.e., solving (20), as long as the step size γ in (26) is properly set, the corresponding cost function is non-increasing. Therefore, $\mathbf{D}^{[i]}$ resulted from the gradient-based approach guarantees

$$\|\mathbf{X} - \mathbf{D}^{[i]}\mathbf{S}^{[i]}\|_F^2 + \omega \|(\mathbf{D}^{[i]})^\mathcal{T}(\mathbf{D}^{[i]}) - \mathbf{I}_K\|_F^2 \leq \|\mathbf{X} - \mathbf{D}^{[i-1]}\mathbf{S}^{[i]}\|_F^2 + \omega \|(\mathbf{D}^{[i-1]})^\mathcal{T}(\mathbf{D}^{[i-1]}) - \mathbf{I}_K\|_F^2.$$

Subsequently,

$$\|\mathbf{X} - \mathbf{D}^{[i]}\mathbf{S}^{[i]}\|_F^2 + \omega \|(\mathbf{D}^{[i]})^\mathcal{T}(\mathbf{D}^{[i]}) - \mathbf{I}_K\|_F^2 \leq \|\mathbf{X} - \mathbf{D}^{[i-1]}\mathbf{S}^{[i-1]}\|_F^2 + \omega \|(\mathbf{D}^{[i-1]})^\mathcal{T}(\mathbf{D}^{[i-1]}) - \mathbf{I}_K\|_F^2$$

is guaranteed throughout the iteration indexed by i .

4. Tight sparse representation-based speech inpainting

Speech inpainting, leveraging the concept of image inpainting, refers to filling in missing data whose locations are assumed to be known *a priori* in the speech signals [4], [28]-[30]. In this section, the learned tight dictionary described in the previous section is applied to speech inpainting problems where the speech signals are corrupted by data missing. We formulate the missing process as a degradation matrix that removes samples from the speech signals and employ preconditioning technique to promote the speech restoration.

4.1. Formulation of data missing

Data missing is defined as a general problem that a partial set of reliable speech data is observed while the remaining unreliable data are totally missing. Speech inpainting aims to estimate the unreliable (missing) data from the reliable data portion. This process is also known as missing data imputation [28]-[30].

According to [27], speech signals satisfy the SR model with a high probability. Assume that $\mathbf{x} \in \mathbb{R}^{N \times 1}$ is a speech signal, meaning that there exists a sparse coefficient vector $\mathbf{s} \in \mathbb{R}^{K \times 1}$ with $\|\mathbf{s}\|_0 \leq \kappa$ ($\kappa \ll N$) satisfying $\mathbf{x} \approx \mathbf{D}\mathbf{s}$ with $\mathbf{D} \in \mathbb{R}^{N \times K}$ denoting a pre-specified dictionary.

Further assume that M samples in $\mathbf{x}$ are observed with *a priori* known positions. The degradation matrix $\mathbf{L} \in \mathbb{R}^{M \times N}$ that records the missing process is a matrix which has only a single element of one and other elements are zeros in each row. Such a matrix can be expressed as

$$\mathbf{L}(i, j) = \begin{cases} 1, & \text{when } j = \mathcal{J}(i), \\ 0, & \text{otherwise,} \end{cases} \quad (37)$$

where i is the element index number of the observed signal, and $1 \leq \mathcal{J}(i) \leq N$ is the index number of corresponding element in the original signal $\mathbf{x}$. Matrix $\mathbf{L}$ removes $N - M$ samples from $\mathbf{x}$ and this process is viewed as data missing. This matrix can also be described as taking the $N \times N$ identity matrix and removing the $N - M$ rows corresponding to the missing samples.

Denote the observed signal as $\mathbf{y}$. The speech data missing problem is formulated as

$$\mathbf{y} \triangleq \mathbf{L}\mathbf{x} \approx \mathbf{L}\mathbf{D}\mathbf{s}. \quad (38)$$

When Gaussian noise is added to the observations, speech inpainting should be considered by solving the following inverse problem

$$\hat{\mathbf{s}} \triangleq \arg \min_{\mathbf{s}} \|\mathbf{s}\|_0 \quad \text{s.t.} \quad \|\mathbf{y} - \mathbf{L}\mathbf{D}\mathbf{s}\|_2^2 \leq \epsilon, \quad (39)$$

with ϵ being the error threshold determined by the noise level. Then the restored signal would become $\hat{\mathbf{x}} = \mathbf{D}\hat{\mathbf{s}}$.

Remark 3.

- The above strategy would result in successful missing data imputation based on a presupposed condition that $\mathbf{x}$ can be sparsely represented using dictionary $\mathbf{D}$. This condition can be satisfied by the previously proposed TDL algorithm. As will be seen in the experiment results, the algorithm learns a tight dictionary with promising SR capability. In addition, when the degradation matrix $\mathbf{L} = \mathbf{I}_N$ (i.e., $M = N$, implying that no sample is missing), (39) can be viewed as a speech denoising problem and will also be considered in the experimental studies to demonstrate the performance of the designed tight dictionary.
- Equation (38) is similar to a compressed sensing model, where $\mathbf{L} \in \mathbb{R}^{M \times N}$ is the sensing matrix [8]. In many works, with the dictionary $\mathbf{D}$ fixed, the sensing matrix $\mathbf{L}$ is designed to make the product $\mathbf{LD}$ possess certain properties such as the tightness property to facilitate the sparse recovery process, thereby improving the accuracy of solution indicated in (39) [17], [18], [24], [25]. This is infeasible in our degradation model because $\mathbf{L}$ is determined by the data missing process and cannot be optimized. Depending on the value of $\mathbf{L}$, the property of $\mathbf{LD}$ may be destroyed even when $\mathbf{D}$ is well designed. This fact will significantly compromise the solution of (39). Motivated by this, we introduce the preconditioning technique to modify (39) and ensure the UNTF properties.

4.2. Construction of tight preconditioner

Preconditioning is one of the most critical ingredients in the development of efficient solvers for challenging problems in scientific computation [31]-[33]. The term preconditioning refers to transforming a model like (38) into another form with more favorable properties. A preconditioner is a matrix that affects such a transformation. Generally speaking, preconditioning attempts to improve the spectral properties of the degraded system matrix $\mathbf{LD}$. Hopefully, the transformed (preconditioned) system matrix will have a smaller spectral condition number. A well-conditioned system matrix will lead to a stable sparse coding procedure [31].

Assume $\mathbf{LD}$ to be an $M \times K$ matrix that does not satisfy the necessary conditions required for sparse recovery. Suppose that there exists a nonsingular matrix $\mathbf{P} \in \mathbb{R}^{M \times M}$ such that the product $\mathbf{PLD}$ possesses certain desired properties to facilitate the sparse recovery process. In this case, $\mathbf{P}$ is referred to as a preconditioner. Multiplying both sides of (38) by $\mathbf{P}$, we define

$$\mathbf{z} \triangleq \mathbf{P}\mathbf{y} = \mathbf{P}\mathbf{L}\mathbf{x} \approx \mathbf{A}\mathbf{s}, \quad (40)$$

where $\mathbf{A} \triangleq \mathbf{PLD}$ is the preconditioned (system) matrix. As $\mathbf{P}$ is invertible, the solution $\mathbf{s}$ of (40) is the same as that of (38). The difference between (38) and (40) in calculating $\mathbf{s}$ lies in the properties of $\mathbf{A}$ and $\mathbf{LD}$. The effectiveness of applying algorithms such as OMP to solve (40) for sparse coefficient $\mathbf{s}$ depends strongly on the property of $\mathbf{A}$. The question is how to construct a nonsingular $\mathbf{P}$ such that $\mathbf{s}$ can be precisely obtained via

$$\hat{\mathbf{s}} \triangleq \arg \min_{\mathbf{s}} \|\mathbf{s}\|_0 \quad \text{s.t.} \quad \|\mathbf{z} - \mathbf{A}\mathbf{s}\|_2^2 \leq \epsilon, \quad (41)$$

with $\mathbf{z}$ and $\mathbf{A}$ both known, then $\hat{\mathbf{x}} = \mathbf{D}\hat{\mathbf{s}}$ is recovered.

Equation (41) is similar to the sparse coding procedure in DL. Therefore, the UNTF properties are preferred. The preconditioner construction problem is formulated as

$$\min_{\mathbf{P} \in \mathbb{R}^{M \times M}} \|\mathbf{A}^T \mathbf{A} - \mathbf{I}_K\|_F^2 \quad \text{s.t.} \quad \mathbf{A} = \mathbf{PLD} \quad \text{and} \quad \|\mathbf{A}(:, k)\|_2^2 = 1, \quad \forall k. \quad (42)$$

Note that the column-normalization constraints are included because our goal is to approximate a UNTF in a strict way. The resulting $\mathbf{P}$ obtained from (42) is named as a tight preconditioner.

Once again, due to the normalization constraints in (42), it is difficult to derive a closed-form solution. Denote

$$\tilde{\mathbf{A}} \triangleq \mathbf{L}\mathbf{D} \in \mathfrak{R}^{M \times K}, \quad (43)$$

and

$$\tilde{\mathbf{Q}} \triangleq \mathbf{P}\tilde{\mathbf{A}} \in \mathfrak{R}^{M \times K} \quad (44)$$

is determined by $\mathbf{P}$ with a mild condition that it contains no zero columns. The gradient-based approach will be employed and the key step is the following parametrization

$$\mathbf{A} \triangleq \tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d \quad (45)$$

with $\tilde{\mathbf{Q}}_d \in \mathfrak{R}^{K \times K}$ being a diagonal matrix formed as

$$\tilde{\mathbf{Q}}_d = \text{diag}(\|\tilde{\mathbf{Q}}(:, 1)\|_2^{-1}, \|\tilde{\mathbf{Q}}(:, 2)\|_2^{-1}, \dots, \|\tilde{\mathbf{Q}}(:, K)\|_2^{-1}). \quad (46)$$

It is clear that the columns of $\mathbf{A}$ in (45) are inherently normalized, yielding the following unconstrained problem:

$$\min_{\mathbf{P} \in \mathfrak{R}^{M \times M}} \left\{ \|(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T}(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d) - \mathbf{I}_K\|_F^2 \triangleq \tilde{\varrho}(\mathbf{P}) \right\} \quad \text{s.t.} \quad (44), (46). \quad (47)$$

Similar to (25), equation (47) can be solved using the gradient descent procedure:

$$\mathbf{P}_n = \mathbf{P}_{n-1} - \eta \left. \frac{\partial \tilde{\varrho}(\mathbf{P})}{\partial \mathbf{P}} \right|_{\mathbf{P}=\mathbf{P}_{n-1}}, \quad (48)$$

where η denotes the step size, and $\frac{\partial \tilde{\varrho}(\mathbf{P})}{\partial \mathbf{P}}$ is the gradient of $\tilde{\varrho}(\mathbf{P})$ with respect to $\mathbf{P}$ and can be computed based on Theorem 3 below.

Theorem 3. Given a known matrix $\tilde{\mathbf{A}} \in \mathfrak{R}^{M \times K}$, let $\mathbf{P} \in \mathfrak{R}^{M \times M}$ be a nonsingular matrix such that $\tilde{\mathbf{Q}} = \mathbf{P}\tilde{\mathbf{A}}$ has no zero columns. Denote

$$\tilde{\mathbf{E}} \triangleq (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T}(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d) - \mathbf{I}_K, \quad (49)$$

with $\tilde{\mathbf{Q}}_d$ defined in (46). The gradient of $\|\tilde{\mathbf{E}}\|_F^2$ with respect to $\mathbf{P}$ is calculated as

$$\frac{\partial(\|\tilde{\mathbf{E}}\|_F^2)}{\partial \mathbf{P}} = 2(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)(\tilde{\mathbf{E}}^\mathcal{T} + \tilde{\mathbf{E}})\tilde{\mathbf{Q}}_d\tilde{\mathbf{A}}^\mathcal{T} - 2\mathbf{P}\tilde{\mathbf{A}}\tilde{\mathbf{R}}\tilde{\mathbf{A}}^\mathcal{T}, \quad (50)$$

where

$$\tilde{\mathbf{R}} \triangleq \text{diag}(\tilde{\mathbf{H}}(1, 1), \tilde{\mathbf{H}}(2, 2), \dots, \tilde{\mathbf{H}}(K, K)), \quad \tilde{\mathbf{H}} \triangleq (\tilde{\mathbf{E}}^\mathcal{T} + \tilde{\mathbf{E}})(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T}\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d^3. \quad (51)$$

The proof of Theorem 3 can be found in Appendix C.

With the result of Theorem 3, $\mathbf{P}$ is updated via (48) by the gradient descent procedure, and the designed tight preconditioner can be obtained as $\lim_{n \rightarrow +\infty} \mathbf{P}_n$.

Remark 4.

- In our interested scenarios, the dictionary $\mathbf{D}$ is learned in advance and the degradation matrix $\mathbf{L}$ is decided according to the data missing process. Therefore, the property of $\mathbf{LD}$ is unpredictable and it may not satisfy the necessary conditions for sparse recovery. That is the reason why we introduce the preconditioning.
- In general, a good preconditioner should lead to a model which can be easily handled. The UNTF properties can make the inverse problem easier to be solved and facilitate sparse recovery [17], [18]. We design the preconditioner to make the preconditioned matrix closely approximate the UNTF by enclosing the normalization constraints.
- Looking from a different perspective, a well-designed preconditioned matrix will result in effective sparse coding [31], [37]. Choosing the identity matrix as the target in (42) makes the physical meanings of the model clearer because the identity matrix has the smallest condition number. This is a key point as the purpose of the preconditioning technique is to improve the spectral properties of the system matrix [31].
- A similar preconditioner construction problem has been investigated in [32]. The problem is described as

$$\min_{\mathbf{P} \in \mathbb{R}^{M \times M}} \|\mathbf{A}^T \mathbf{A} - \mathbf{G}_t\|_F^2 \quad \text{s.t.} \quad \mathbf{A} = \mathbf{P} \tilde{\mathbf{A}},$$

where $\mathbf{G}_t$ is the target Gram matrix possessing some desired properties. The authors of [32] first computed a matrix $\mathbf{A}$ for the above problem and then solved a least squares problem to obtain $\mathbf{P} = \arg \min_{\mathbf{P}} \{\|\mathbf{A} - \mathbf{P} \tilde{\mathbf{A}}\|_F^2\} = \mathbf{A} \tilde{\mathbf{A}}^T (\tilde{\mathbf{A}} \tilde{\mathbf{A}}^T)^{-1}$. A similar strategy is adopted in [25] to design the sensing matrix for a compressed sensing system. Such a two-step strategy compromises the optimality of the solution. Using equation (48) and Theorem 3, we can directly construct the preconditioner for problem (42).

- The proposed preconditioner construction method shares some of the ideas in our recent paper [33]. Different from [33], however, the data missing process is taken into consideration in this paper, and the parametrization and gradient-based approaches are described in detail.

5. Experiment results

In this section, we evaluate the performance of the proposed models and algorithms with speech signals that are popularly used for testing the effectiveness of the constrained DL [37], [38].

The speech data are extracted from the GRID corpus [43]. Every recording in the corpus is decomposed into 50% overlapping time-domain frames with a frame length of $N = 256$. A data matrix consisting of 100,000 columns is formed by randomly selecting 100 frames from each of the 1,000 recordings in the corpus. A twice overcomplete dictionary $\mathbf{D}^{[0]}$ (i.e., $K = 512$) is initialized by choosing K columns of the data matrix followed by a column-normalization procedure. The training data matrix $\mathbf{X} = \{\mathbf{x}_l\}_{l=1}^L$ is obtained as a subset of the data matrix with $L = 10,000$ columns to learn the dictionaries. The testing data matrix $\tilde{\mathbf{X}} = \{\tilde{\mathbf{x}}_l\}_{l=1}^{\tilde{L}}$ is obtained in the same way as $\mathbf{X}$ but with $\tilde{L} = 1,000$ columns. It should be noted that the initial dictionary, the training data matrix, and the testing data matrix are all selected randomly without intersection. When learning the dictionaries, the OMP algorithm is employed for sparse coding with sparsity $\kappa = 12$ which corresponds to about 5% of active elements when compared with the length N .

The standard deviation of the Gaussian noise added to the testing data $\tilde{X}$ is denoted as σ . The recovery accuracy is quantified with the output signal-to-noise ratio (SNR) [28]-[30] defined as

$$\rho_{snr} \triangleq 20 \log_{10} \frac{\|\tilde{X}\|_F}{\|\tilde{X} - \mathbf{D}\tilde{\mathbf{S}}\|_F}, \quad (52)$$

where $\mathbf{D}$ is the output dictionary of each learning method and $\tilde{\mathbf{S}} = \{\tilde{s}_l\}_{l=1}^{\tilde{L}}$ is the sparse coefficient matrix calculated from

$$\tilde{s}_l = \arg \min_{s_l} \|s_l\|_1 \quad \text{s.t.} \quad \|\tilde{\mathbf{x}}_l - \mathbf{D}s_l\|_2^2 \leq \epsilon, \quad \forall l$$

by the ℓ_1 -Homotopy algorithm [16] with $\epsilon = 0.5N\sigma^2$. This process can be viewed as the speech denoising problem.

5.1. Performance of the tight dictionary

We now carry out experiments to illustrate the performance of the dictionaries learned with different algorithms. Algorithms developed in [11], [37] and [38] are compared. For convenience, the learning systems are denoted as **Dict**_{NEW}, **Dict**_{KSVD} [11], **Dict**_{SDB} [37], and **Dict**_{BLJ} [38] for the proposed TDL and the algorithms in the respective references. The number of iterations for DL is set to $N_{\text{iter}} = 100$ for all the comparisons. For **Dict**_{NEW}, the gradient descent is executed for 50 times with step size $\gamma = 0.1$. The extra parameter settings for **Dict**_{KSVD}, **Dict**_{SDB} and **Dict**_{BLJ} are kept the same as those in the corresponding references.

5.1.1. Choice of ω in (12) and (18)

Let us consider the effect of factor ω . In order to test the tradeoff between the representation capability and the tightness property roundly, we denote

$$\tilde{\omega} \triangleq \frac{\omega}{1 + \omega}.$$

When ω increases from $0 \rightarrow +\infty$, $\tilde{\omega}$ changes monotonically from $0 \rightarrow 1$.

Gaussian noises with different standard deviations are added to the testing data $\tilde{X}$. Figures 1 (a), (b) and (c) show the results of the recovered SNR with respect to $\tilde{\omega}$. The input SNR is the SNR of the noisy speech determined by the noise standard deviation σ .

For Figures 1 (a), (b) and (c), the significance of the regularization term $\|\mathbf{D}^T \mathbf{D} - \mathbf{I}_K\|_F^2$ is judged by comparing the cases for $\tilde{\omega} = 0$ with the others, while the significance of $\|\mathbf{X} - \mathbf{D}\mathbf{S}\|_F^2$ is examined by the cases with $\tilde{\omega} = 1$. It can be clearly seen that the best results of **Dict**_{NEW} and **Dict**_{SDB} are always achieved when $0 < \tilde{\omega} < 1$, which indicates that the weighting models (12) and (18) are valid. Besides, comparing the results of **Dict**_{NEW} with those of **Dict**_{SDB}, the importance of the column-normalization constraints is revealed, as the satisfactory results of **Dict**_{NEW} are much better than those of **Dict**_{SDB}. Though the optimum $\tilde{\omega}$ values are data-dependent and differ for various σ , hereinafter, $\tilde{\omega} = 0.2$, hence $\omega = 0.25$ is set for **Dict**_{NEW} and **Dict**_{SDB} according to its satisfactory performance for all the tests.

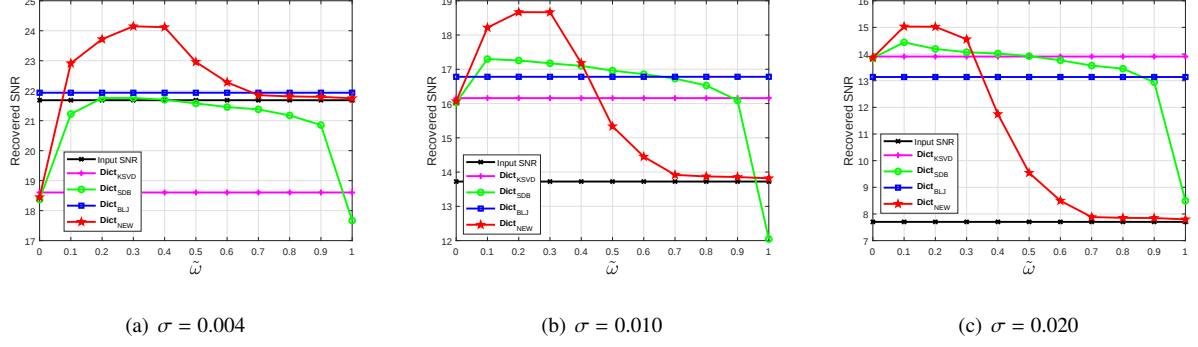


Figure 1: Recovered SNR versus $\tilde{\omega}$.

5.1.2. Theoretical performance of different dictionaries

Before testing the denoising performance, we take a brief look at the singular value distribution of each learned dictionary. As indicated in Section 2, for a specific application, it is important to control over the spectrum of the corresponding system matrix and this spectrum is closely related to the singular value distribution of the system matrix. The more compact the singular value distribution is, the flatter the matrix spectrum will be, and this also implies smaller condition number [21], [22], [36], [37]. So a compact singular value distribution means a well-conditioned dictionary, leading to efficient sparse coding.

A UNTF of dimension $N \times K$ has all its singular values equal to $\sqrt{K/N}$ [34], [37]. Taking this as the benchmark, the singular value distributions of the learned dictionaries are shown in Figure 2.

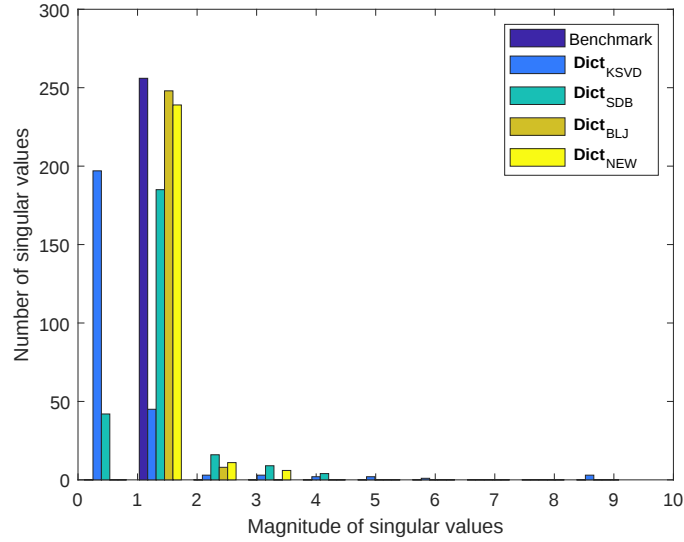


Figure 2: Singular value distributions of the learned dictionaries.

As can be seen, the distribution of $\mathbf{Dict}_{BLJ}$ is closest to the benchmark $\sqrt{K/N}$ because $\mathbf{Dict}_{BLJ}$ is designed under the constraint (15), which is exactly the expression of 1-tight frame. Though extra normalization is enforced, $\mathbf{Dict}_{BLJ}$ is still the best-conditioned one. However, the constraint (15) limits the SR capability of the dictionary to signals since only two orthonormal matrices U and V are optimized to minimize the SR error. Defined as the ratio of the maximum singular value and the minimum one, the condition number of each dictionary is calculated as $\text{cond}(\mathbf{Dict}_{KSVD})=2235.10$, $\text{cond}(\mathbf{Dict}_{SDB})=17.12$, $\text{cond}(\mathbf{Dict}_{BLJ})=1.89$, $\text{cond}(\mathbf{Dict}_{NEW})=2.57$. As indicated in [36], the closer the condition number to 1, the better the matrix stability will be. We can say that $\mathbf{Dict}_{NEW}$ is also well conditioned. It can be expected that, if a large ω is set, $\text{cond}(\mathbf{Dict}_{NEW})$ will be much lower at the cost of the SR capability. So proper tradeoff is necessary in the proposed model (18).

5.1.3. Denoising performance on testing data

Now we test the denoising performance of the compared dictionaries with the testing data. Figure 3 depicts the output SNRs of the recovered signals when the Gaussian noises with standard deviation σ varying between 0.001 and 0.030 are added to the testing data.

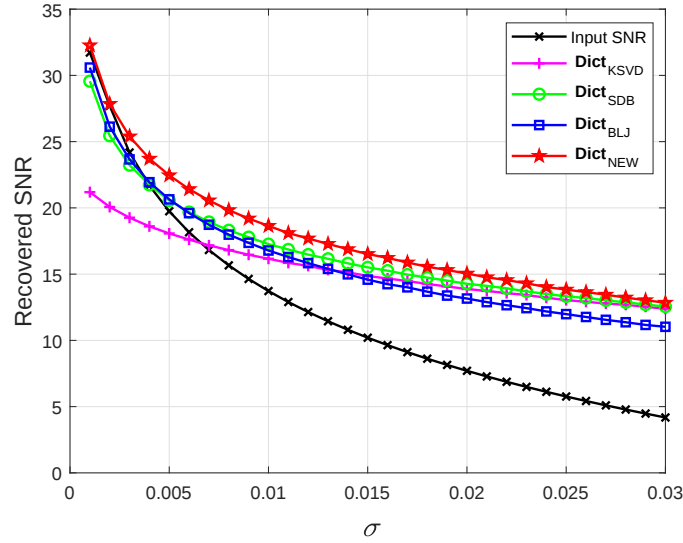


Figure 3: Recovered SNR results versus noise standard deviation σ .

As shown in Figure 3, the recovered SNR values all decrease as σ increases, and $\mathbf{Dict}_{NEW}$ achieves the best denoising results among these tests. The superiority of $\mathbf{Dict}_{NEW}$ over $\mathbf{Dict}_{KSVD}$ is clearly observed in low noise cases. When $\sigma \leq 0.012$, the gap is more than 2 dB. For high noise cases, e.g., $\sigma \geq 0.020$, the performance of $\mathbf{Dict}_{BLJ}$ is poorer than the other three. So we conclude that neither $\mathbf{Dict}_{KSVD}$ nor $\mathbf{Dict}_{BLJ}$ is robust enough. The results of $\mathbf{Dict}_{SDB}$ are close to those of $\mathbf{Dict}_{NEW}$. This verifies the effectiveness of the weighting model which emphasizes the

tightness property of the dictionary. The remaining differences between these two can be explained by the column-normalization constraints that make the physical meanings of our model much clearer.

5.1.4. Denoising performance on speech segment

Besides the speech data, we carry out experiments on real speech segment for testing the denoising performance of different dictionaries. Phrase “place blue at P zero please” is randomly selected for testing. The speech segment is first divided into 143 non-overlapping time-domain frames with each frame of length $N = 256$. Denoising is executed in every frame independently by each of the learned dictionaries, and the restored frames synthesize a speech segment.

Figures 4 (a), (b) and (c) show the time-domain waveforms of the speech segments restored by different dictionaries when the standard deviations of the Gaussian noises are 0.004, 0.010 and 0.020, respectively. In order to clearly see the details, only 500 samples (starting from 5, 501st in 36, 608 samples) of each waveform are drawn for comparison.

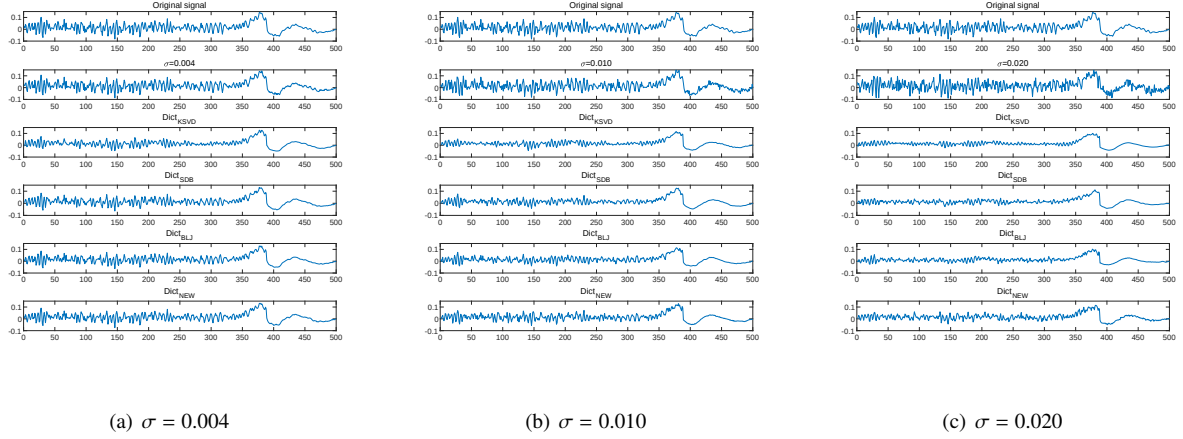
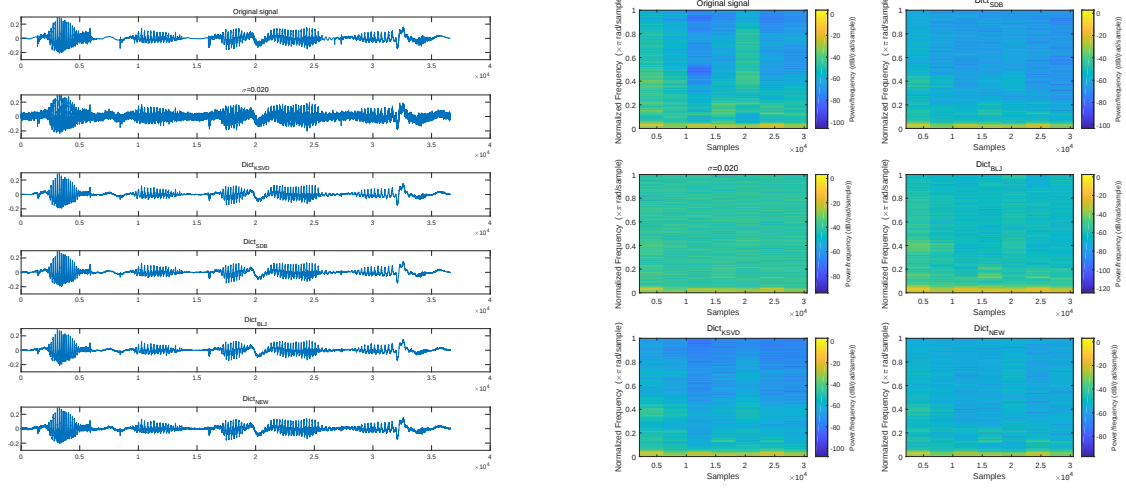


Figure 4: Restored speech waveforms by different dictionaries.

Taking a close look at these waveforms, one notes that **Dict**_{KSVD} loses the details easily when the signals fluctuate acutely, even in the case with a small $\sigma = 0.004$, whereas the other three dictionaries perform favourably. When $\sigma = 0.020$, it can be observed that the waveform resulted from **Dict**_{BLI} is still heavily disturbed by noise. These results confirm that neither **Dict**_{KSVD} nor **Dict**_{BLI} is robust enough.

The spectrograms of the restored speech segments are compared when $\sigma = 0.020$. Figures 5 (a) and (b) depict the complete speech waveforms and the corresponding spectrograms (drawn by MATLAB with default settings), respectively. The dark color parts (closing to blue) and the light color parts (closing to yellow) in Figure 5 (b) refer to low frequency and high frequency components, respectively. As can be observed, **Dict**_{KSVD} loses a lot of high frequency information and **Dict**_{BLI} cannot eliminate noise well. These conclusions are consistent with those obtained from Figure 4.



(a) Restored speech waveforms with $\sigma = 0.020$

(b) Spectrograms of the speech segments

Figure 5: Restored speech waveforms and spectrograms with $\sigma = 0.020$.

We summarize the SNR results of all the restored speech segments in Table 1. The output SNR results also verify the above observations. Although the results of **Dict**_{SDB} are acceptable, it should be emphasized that both the waveforms and the recovered SNR results clearly verify that **Dict**_{NEW} outperforms other methods in all these tests.

Table 1: Recovered SNRs of speech segments by different dictionaries (dB)

	$\sigma = 0.004$	$\sigma = 0.010$	$\sigma = 0.020$
Input SNR	21.93	13.97	7.95
Dict _{KSVD}	19.59	17.02	14.86
Dict _{SDB}	21.83	17.65	14.89
Dict _{BLI}	21.91	16.99	13.40
Dict _{NEW}	23.61	18.85	15.53

5.2. Performance of TSR-based speech inpainting

In this part, the learned tight dictionary **D** is applied to speech inpainting cases where the speech signals are corrupted by data missing and Gaussian noise. In order to understand the inpainting performance intuitively, real speech segments are considered. The inpainting process is carried out on non-overlapping time-domain frames with frame length $N = 256$. For each frame, data missing happens at random and the degradation matrix $\mathbf{L} \in \mathbb{R}^{M \times N}$ that

records the data missing process is structured as (37). It means that $N - M$ samples are lost in every frame. When designing the tight preconditioner $\mathbf{P}$, the gradient descent is executed for 1,000 times with step size $\eta = 0.01$.

5.2.1. Theoretical performance of tight preconditioner

First, the theoretical performance of the constructed preconditioner $\mathbf{P}$ is illustrated in Figure 6 by the singular value distribution of the preconditioned matrix $\mathbf{PLD}$. For comparison, the distribution of $\mathbf{LD}$ is also attached. Taking $\sqrt{K/M}$ (i.e., the spectrum of the $M \times K$ UNTF [34], [37]) as the benchmark, the singular value distributions are shown in Figures 6 (a), (b) and (c) with $M = 60$, $M = 130$ and $M = 190$, respectively.

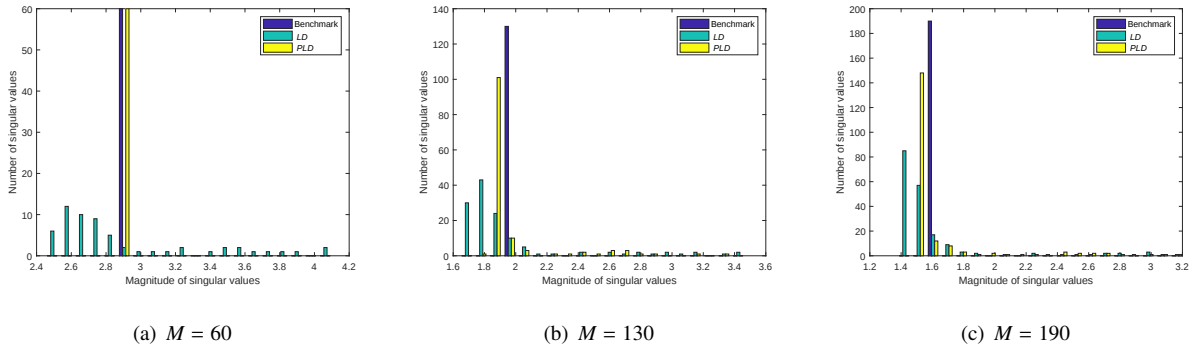


Figure 6: Singular value distributions of corresponding system matrices.

In Figure 6 (a), the preconditioner $\mathbf{P}$ makes the preconditioned matrix $\mathbf{PLD}$ an exact UNTF when $M = 60$. This verifies the validity of our parametrization process (45) and the gradient-based approach for solving (42). Though similar phenomenon is not observed in Figures 6 (b) and (c), a more compact singular value distribution is expected if the step size η in (48) is chosen more carefully and intelligently. This would be explored in the future work.

For these three values of M , the condition numbers of the corresponding system matrices are summarized in Table 2. The results further confirm the effectiveness of the designed preconditioner. For these three values of M , the condition numbers of $\mathbf{PLD}$ are much closer to 1 than those of $\mathbf{LD}$. They lead to stable preconditioned matrices that facilitate the sparse recovery process, i.e., solving the inverse problem (41).

Table 2: Condition numbers of corresponding system matrices

	$M = 60$	$M = 130$	$M = 190$
$\text{cond}(\mathbf{LD})$	1.68	2.11	2.34
$\text{cond}(\mathbf{PLD})$	1.00	1.89	2.18

5.2.2. Inpainting performance evaluation

In this part, the speech inpainting problem is investigated and three processes are compared. The speech inpainting process executed by (39) is a TSR-based strategy that is denoted as $\mathbf{SI}_{\text{TSR}}$. (41) is a preconditioned sparse recovery process and we name it as $\mathbf{SI}_{\text{PTSR}}$. In addition, $\mathbf{Dict}_{\text{NEW}}$ is applied to eliminate noise without data missing just as in the denoising part in Section 5.1.4 for comparison. In each of the tests, the input SNR is calculated prior to data missing.

The same phrase “place blue at P zero please” is chosen for testing as in the denoising cases. When $M = 130$, the recovered SNR values are depicted in Figure 7 for σ varying between 0.001 and 0.030. The results verify the validity of the TSR-based inpainting. Figure 7 shows that, when $\sigma > 0.008$, the results of $\mathbf{SI}_{\text{TSR}}$ are higher than the input SNR (without data missing taken into account) values. It means that the inpainting strategy not only restores missing data, but also eliminates additive noises.

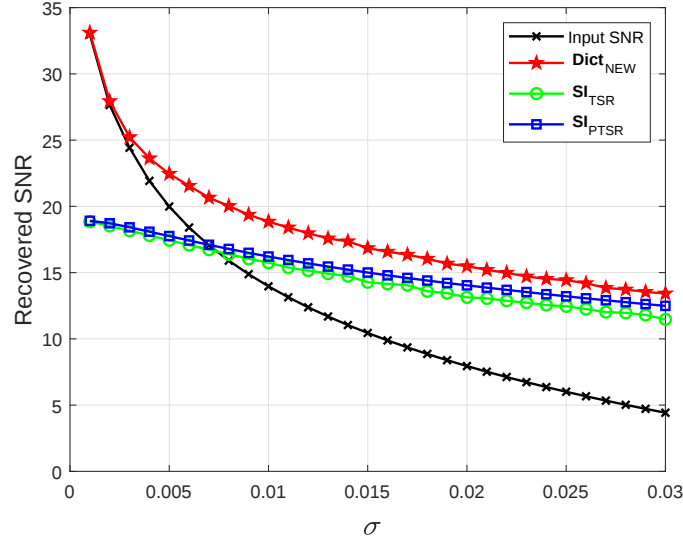


Figure 7: Recovered SNR results versus σ with $M = 130$.

Next, we evaluate the effects of the number of observations. For $\sigma = 0.010$ fixed, when M changes between 60 and 190, the SNR results are shown in Figure 8. As can be seen, even when only $M = 80$ observations are reserved, $\mathbf{SI}_{\text{TSR}}$ restores the signals and leads to an SNR result that is comparable with the input SNR. More observations will result in remarkable improvements.

With two more speech segments tested, the statistics of the output SNR for these three phrases in different σ and M cases are summarized in Table 3. In all the tests, $\mathbf{SI}_{\text{PTSR}}$ achieves higher SNR results than those of $\mathbf{SI}_{\text{TSR}}$. This indicates that the preconditioning technique employed to facilitate the sparse recovery process is effective, and the tight preconditioner constructed by our method is superior.

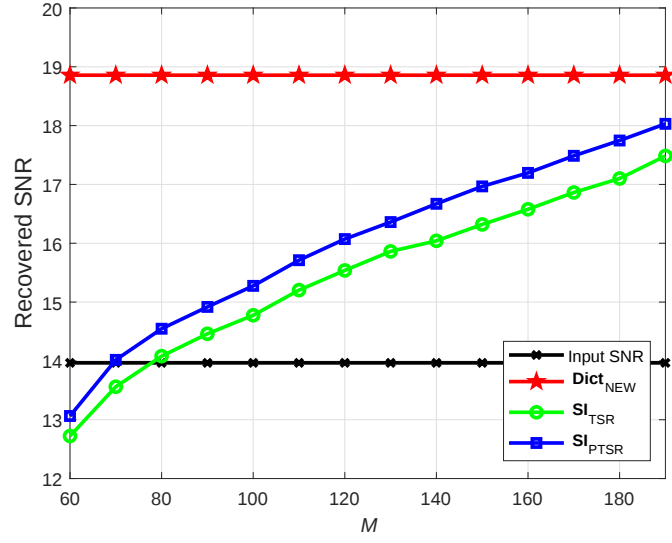


Figure 8: Recovered SNR results versus M with $\sigma = 0.010$.

Table 3: Statistics of SNR for three phrases in different σ and M cases (dB)

Phrases	Processes	$\sigma = 0.004$			$\sigma = 0.010$			$\sigma = 0.020$		
		$M = 60$	$M = 130$	$M = 190$	$M = 60$	$M = 130$	$M = 190$	$M = 60$	$M = 130$	$M = 190$
place blue at P zero please	Input SNR	21.93			13.97			7.95		
	Dict _{NEW}	23.66			18.95			15.60		
	SI _{TSR}	13.76	17.59	20.39	12.73	15.89	17.48	10.32	13.48	14.74
	SI _{PTSR}	14.07	17.93	20.86	13.08	16.38	18.04	10.70	14.16	15.59
bin blue at L two please	Input SNR	21.94			13.98			7.96		
	Dict _{NEW}	24.39			19.54			16.03		
	SI _{TSR}	16.04	19.08	21.07	13.95	16.59	18.15	10.47	14.12	15.26
	SI _{PTSR}	16.34	19.39	21.43	14.25	17.08	18.73	10.82	14.81	16.34
set white with I nine soon	Input SNR	21.92			13.97			7.95		
	Dict _{NEW}	23.73			18.31			14.53		
	SI _{TSR}	12.31	16.68	19.83	8.78	14.26	16.35	7.46	11.52	13.19
	SI _{PTSR}	12.57	17.07	20.38	8.99	14.87	17.20	7.80	12.23	14.15

6. Conclusion

In this paper, we have investigated the TSR problem with application to speech inpainting. A novel DL model, named TDL, was proposed. Under this model, the dictionary is learned to minimize the SR error adapted to the training signals and, at the same time, approximate the UNTF to gain the desired frame properties. The parametrization

method was adopted to embed the normalization constraints on the dictionary columns into the cost function. A gradient-based approach was developed to solve the parameterized problem for the tight dictionary. Moreover, we have applied the tight dictionary to the speech inpainting problem where data missing is considered. Finally, in order to maintain the frame properties of the system, the preconditioning technique based on UNTF was employed, and the parametrization and gradient-based approaches were carried out once again to solve the normalization constrained tight preconditioner design problem. Extensive experiments have been conducted to demonstrate the validity and effectiveness of the proposed TSR schemes on speech denoising and speech segments restoration.

Acknowledgement

This work is supported in part by the Natural Science Foundation of China under Grants 61801159 and 61571174, in part by the Key R&D Program Projects in Zhejiang Province under Grant 2020C03074, in part by the Start-up Research Program from Hangzhou Normal University under Grant 2018QDL018, and in part by the China-Montenegro Bilateral Science & Technology Cooperation Project.

Declaration of interest

The authors declare that they do not have any financial or non-financial conflict of interests.

Appendix A. Proof of Theorem 1

First, one notes that

$$\|\mathbf{E}_1\|_F^2 = \text{tr}[\mathbf{E}_1 \mathbf{E}_1^\mathcal{T}]$$

with $\text{tr}[\cdot]$ denoting the trace operator.

Denote $q_{ij} = \mathbf{Q}(i, j)$, then

$$\frac{\partial(\|\mathbf{E}_1\|_F^2)}{\partial q_{ij}} = 2\text{tr}\left[\frac{\partial \mathbf{E}_1}{\partial q_{ij}} \mathbf{E}_1^\mathcal{T}\right],$$

and

$$\frac{\partial \mathbf{E}_1}{\partial q_{ij}} = -\frac{\partial(\mathbf{Q}\mathbf{Q}_d)}{\partial q_{ij}} \mathbf{S}.$$

Given the definition of $\mathbf{Q}_d$ in (24), it can be shown that

$$\frac{\partial(\mathbf{Q}\mathbf{Q}_d)}{\partial q_{ij}} = \mathbf{e}_i \mathbf{e}_j^\mathcal{T} \mathbf{Q}_d - \mathbf{Q} q_{ij} d_j^3 \mathbf{e}_j \mathbf{e}_j^\mathcal{T},$$

where $\mathbf{e}_k$ denotes the k -th elementary vector with a proper dimension, whose elements are all zeros except the k -th one that is 1. For example, in the above expression, the dimensions of $\mathbf{e}_i$ and $\mathbf{e}_j$ are respectively $N \times 1$ and $K \times 1$.

$d_j \triangleq (\sum_i q_{ij}^2)^{-1/2} = \|\mathbf{Q}(:, j)\|_2^{-1}$ is the j -th diagonal element of $\mathbf{Q}_d$. Then, we get

$$\frac{\partial \mathbf{E}_1}{\partial q_{ij}} = -\mathbf{e}_i \mathbf{e}_j^\mathcal{T} \mathbf{Q}_d \mathbf{S} + q_{ij} d_j^3 \mathbf{Q} \mathbf{e}_j \mathbf{e}_j^\mathcal{T} \mathbf{S}$$

and

$$\frac{\partial(\|\mathbf{E}_1\|_F^2)}{\partial q_{ij}} = 2(-\mathbf{e}_j^\mathcal{T} \mathbf{Q}_d \mathbf{S} \mathbf{E}_1^\mathcal{T} \mathbf{e}_i + q_{ij} d_j^3 \mathbf{e}_j^\mathcal{T} \mathbf{S} \mathbf{E}_1^\mathcal{T} \mathbf{Q} \mathbf{e}_j).$$

As $q_{ij} = \mathbf{Q}(i, j)$, i.e., $q_{ij} = \mathbf{e}_i^\mathcal{T} \mathbf{Q} \mathbf{e}_j$, we finally obtain

$$\frac{\partial(\|\mathbf{E}_1\|_F^2)}{\partial \mathbf{Q}} = 2\mathbf{Q}\mathbf{Q}_d^3 \mathbf{R}_1 - 2\mathbf{E}_1 \mathbf{S}^\mathcal{T} \mathbf{Q}_d$$

with $\mathbf{R}_1$ defined in (32). This completes the proof.

Appendix B. Proof of Theorem 2

Similar to Appendix A, we have

$$\frac{\partial(\|E_2\|_F^2)}{\partial q_{ij}} = 2\text{tr}\left[\frac{\partial E_2}{\partial q_{ij}}E_2^\mathcal{T}\right]$$

and

$$\frac{\partial E_2}{\partial q_{ij}} = \frac{\partial(QQ_d)^\mathcal{T}}{\partial q_{ij}}(QQ_d) + (QQ_d)^\mathcal{T} \frac{\partial(QQ_d)}{\partial q_{ij}} = Q_d e_j e_i^\mathcal{T} (QQ_d) - q_{ij} d_j^3 e_j e_j^\mathcal{T} Q^\mathcal{T} (QQ_d) + (QQ_d)^\mathcal{T} e_i e_j^\mathcal{T} Q_d - q_{ij} d_j^3 (QQ_d)^\mathcal{T} Q e_j e_j^\mathcal{T}$$

as Q_d is diagonal. Hence

$$\frac{\partial(\|E_2\|_F^2)}{\partial q_{ij}} = 2 \left[e_i^\mathcal{T} (QQ_d) E_2^\mathcal{T} Q_d e_j - q_{ij} d_j^3 e_j^\mathcal{T} Q^\mathcal{T} (QQ_d) E_2^\mathcal{T} e_j + e_j^\mathcal{T} Q_d E_2^\mathcal{T} (QQ_d)^\mathcal{T} e_i - q_{ij} d_j^3 e_j^\mathcal{T} E_2^\mathcal{T} (QQ_d)^\mathcal{T} Q e_j \right].$$

Then, we obtain

$$\frac{\partial(\|E_2\|_F^2)}{\partial Q} = 2(QQ_d)(E_2^\mathcal{T} + E_2)Q_d - 2QQ_d^3 R_2$$

with R_2 defined in (35). This completes the proof.

Appendix C. Proof of Theorem 3

Because

$$\|\tilde{\mathbf{E}}\|_F^2 = \text{tr}[\tilde{\mathbf{E}}\tilde{\mathbf{E}}^\mathcal{T}],$$

we obtain

$$\frac{\partial(\|\tilde{\mathbf{E}}\|_F^2)}{\partial p_{ij}} = 2\text{tr}\left[\frac{\partial\tilde{\mathbf{E}}}{\partial p_{ij}}\tilde{\mathbf{E}}^\mathcal{T}\right]$$

and

$$\frac{\partial\tilde{\mathbf{E}}}{\partial p_{ij}} = \frac{\partial(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T}}{\partial p_{ij}}(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d) + (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T}\frac{\partial(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)}{\partial p_{ij}}.$$

Let $K \times K$ matrix $\Xi_k \triangleq \mathbf{e}_k \mathbf{e}_k^\mathcal{T}$. Denote

$$\tilde{d}_k = \|\tilde{\mathbf{Q}}(:, k)\|_2^{-1} = \|\mathbf{P}\tilde{\mathbf{A}}(:, k)\|_2^{-1}$$

as the k -th diagonal element of $\tilde{\mathbf{Q}}_d$. Then $\tilde{\mathbf{Q}}_d$ can be expressed as

$$\tilde{\mathbf{Q}}_d = \sum_{k=1}^K \tilde{d}_k \Xi_k.$$

Denote $\mathbf{a}_k = \tilde{\mathbf{A}}(:, k)$ as the k -th column of $\tilde{\mathbf{A}}$. One can derive that

$$\frac{\partial\tilde{\mathbf{Q}}_d}{\partial p_{ij}} = -\sum_{k=1}^K \tilde{d}_k^3 \mathbf{a}_k^\mathcal{T} \mathbf{e}_j \mathbf{e}_i^\mathcal{T} \mathbf{P} \mathbf{a}_k \Xi_k$$

with $\mathbf{e}_i$ and $\mathbf{e}_j$ both elementary vectors of dimension $M \times 1$. Hence

$$\frac{\partial(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)}{\partial p_{ij}} = \frac{\partial\tilde{\mathbf{Q}}}{\partial p_{ij}}\tilde{\mathbf{Q}}_d + \tilde{\mathbf{Q}}\frac{\partial\tilde{\mathbf{Q}}_d}{\partial p_{ij}} = \mathbf{e}_i \mathbf{e}_j^\mathcal{T} \tilde{\mathbf{A}}\tilde{\mathbf{Q}}_d - \tilde{\mathbf{Q}} \sum_{k=1}^K \tilde{d}_k^3 \mathbf{a}_k^\mathcal{T} \mathbf{e}_j \mathbf{e}_i^\mathcal{T} \mathbf{P} \mathbf{a}_k \Xi_k.$$

Therefore,

$$\begin{aligned} & \text{tr}\left[(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T} \frac{\partial(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)}{\partial p_{ij}} \tilde{\mathbf{E}}^\mathcal{T}\right] \\ &= \text{tr}\left[(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T} (\mathbf{e}_i \mathbf{e}_j^\mathcal{T} \tilde{\mathbf{A}}\tilde{\mathbf{Q}}_d - \tilde{\mathbf{Q}} \sum_{k=1}^K \tilde{d}_k^3 \mathbf{a}_k^\mathcal{T} \mathbf{e}_j \mathbf{e}_i^\mathcal{T} \mathbf{P} \mathbf{a}_k \Xi_k) \tilde{\mathbf{E}}^\mathcal{T}\right] \\ &= \mathbf{e}_j^\mathcal{T} \tilde{\mathbf{A}}\tilde{\mathbf{Q}}_d \tilde{\mathbf{E}}^\mathcal{T} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T} \mathbf{e}_i - \sum_{k=1}^K \mathbf{e}_i^\mathcal{T} \mathbf{P} \mathbf{a}_k \tilde{d}_k^3 \text{tr}\left[(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T} \tilde{\mathbf{Q}} \Xi_k \tilde{\mathbf{E}}^\mathcal{T}\right] \mathbf{a}_k^\mathcal{T} \mathbf{e}_j \\ &= \mathbf{e}_j^\mathcal{T} \tilde{\mathbf{A}}\tilde{\mathbf{Q}}_d \tilde{\mathbf{E}}^\mathcal{T} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T} \mathbf{e}_i - \sum_{k=1}^K \mathbf{e}_i^\mathcal{T} \mathbf{P} \mathbf{a}_k \tilde{d}_k^3 \left[\mathbf{e}_k^\mathcal{T} \tilde{\mathbf{E}}^\mathcal{T} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T} \tilde{\mathbf{Q}} \mathbf{e}_k\right] \mathbf{a}_k^\mathcal{T} \mathbf{e}_j \\ &\triangleq \mathbf{e}_j^\mathcal{T} \tilde{\mathbf{A}}\tilde{\mathbf{Q}}_d \tilde{\mathbf{E}}^\mathcal{T} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T} \mathbf{e}_i - \mathbf{e}_i^\mathcal{T} \mathbf{P} \sum_{k=1}^K \mathbf{a}_k \tilde{d}_k^3 \tilde{\mathbf{H}}_1(k, k) \mathbf{a}_k^\mathcal{T} \mathbf{e}_j \\ &\triangleq \mathbf{e}_j^\mathcal{T} \tilde{\mathbf{A}}\tilde{\mathbf{Q}}_d \tilde{\mathbf{E}}^\mathcal{T} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^\mathcal{T} \mathbf{e}_i - \mathbf{e}_i^\mathcal{T} \mathbf{P} \mathbf{G}_1 \mathbf{e}_j, \end{aligned}$$

where

$$\tilde{\mathbf{H}}_1 \triangleq \tilde{\mathbf{E}}^T (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^T \tilde{\mathbf{Q}}, \quad \mathbf{G}_1 \triangleq \sum_{k=1}^K \tilde{d}_k^3 \tilde{\mathbf{H}}_1(k, k) \mathbf{a}_k \mathbf{a}_k^T.$$

According to the characteristics of trace operation, one has

$$\text{tr} \left[\frac{\partial(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^T}{\partial p_{ij}} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d) \tilde{\mathbf{E}}^T \right] = \text{tr} \left[(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^T \frac{\partial(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)}{\partial p_{ij}} \tilde{\mathbf{E}} \right].$$

Therefore,

$$\text{tr} \left[\frac{\partial(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^T}{\partial p_{ij}} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d) \tilde{\mathbf{E}}^T \right] = \mathbf{e}_j^T \tilde{\mathbf{A}} \tilde{\mathbf{Q}}_d \tilde{\mathbf{E}} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^T \mathbf{e}_i - \mathbf{e}_i^T \mathbf{P} \mathbf{G}_2 \mathbf{e}_j,$$

where

$$\mathbf{G}_2 \triangleq \sum_{k=1}^K \tilde{d}_k^3 \tilde{\mathbf{H}}_2(k, k) \mathbf{a}_k \mathbf{a}_k^T$$

with

$$\tilde{\mathbf{H}}_2 \triangleq \tilde{\mathbf{E}} (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^T \tilde{\mathbf{Q}}.$$

Subsequently,

$$\frac{\partial(\|\tilde{\mathbf{E}}\|_F^2)}{\partial p_{ij}} = 2 \text{tr} \left[\frac{\partial \tilde{\mathbf{E}}}{\partial p_{ij}} \tilde{\mathbf{E}}^T \right] = 2 \left[\mathbf{e}_j^T \tilde{\mathbf{A}} \tilde{\mathbf{Q}}_d (\tilde{\mathbf{E}}^T + \tilde{\mathbf{E}}) (\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)^T \mathbf{e}_i - \mathbf{e}_i^T \mathbf{P} (\mathbf{G}_1 + \mathbf{G}_2) \mathbf{e}_j \right],$$

which leads to

$$\frac{\partial(\|\tilde{\mathbf{E}}\|_F^2)}{\partial \mathbf{P}} = 2(\tilde{\mathbf{Q}}\tilde{\mathbf{Q}}_d)(\tilde{\mathbf{E}}^T + \tilde{\mathbf{E}})\tilde{\mathbf{Q}}_d \tilde{\mathbf{A}}^T - 2\mathbf{P}\mathbf{G},$$

where

$$\mathbf{G} \triangleq \mathbf{G}_1 + \mathbf{G}_2 = \sum_{k=1}^K \tilde{d}_k^3 \left[\tilde{\mathbf{H}}_1(k, k) + \tilde{\mathbf{H}}_2(k, k) \right] \mathbf{a}_k \mathbf{a}_k^T.$$

Using the definition of $\tilde{\mathbf{R}}$ in (51), $\mathbf{G}$ can be calculated as $\mathbf{G} = \tilde{\mathbf{A}} \tilde{\mathbf{R}} \tilde{\mathbf{A}}^T$. This completes the proof.

References

- [1] P. G. Casazza, G. Kutyniok, and F. Philipp, "Introduction to finite frame theory." In: P. G. Casazza, G. Kutyniok (eds.), *Finite Frames: Theory and Applications*, Applied and Numerical Harmonic Analysis, pp. 1-53, Birkhäuser, 2013.
- [2] O. Christensen, *An Introduction to Frames and Riesz Bases (Second Edition)*, Birkhäuser, 2016.
- [3] A. M. Bruckstein, D. L. Donoho, and M. Elad, "From sparse solutions of systems of equations to sparse modeling of signals and images," *SIAM Rev.*, vol. 51, no. 1, pp. 34-81, Mar. 2009.
- [4] M. Elad, *Sparse and Redundant Representations: From Theory to Applications in Signal and Image Processing*, Springer, 2010.
- [5] M. Elad and M. Aharon, "Image denoising via sparse and redundant representations over learned dictionaries," *IEEE Trans. Image Process.*, vol. 15, no. 12, pp. 3736-3745, Dec. 2006.
- [6] J. Karmakar, A. Pathak, D. Nandi, and M. K. Mandal, "Sparse representation based compressive video encryption using hyper-chaos and DNA coding," *Digit. Signal Process.*, vol. 117, Article 103143, Oct. 2021.
- [7] S. Ding, G. Zhao, C. Liberatore, and R. Gutierrez-Osuna, "Learning structured sparse representations for voice conversion," *IEEE-ACM Trans. Audio Speech Lang.*, vol. 28, pp. 343-354, 2020.
- [8] H. Bai, G. Li, S. Li, Q. Li, Q. Jiang, and L. Chang, "Alternating optimization of sensing matrix and sparsifying dictionary for compressed sensing," *IEEE Trans. Signal Process.*, vol. 63, no. 6, pp. 1581-1594, Mar. 2015.
- [9] I. Tošić and P. Frossard, "Dictionary learning: What is the right representation for my signal?," *IEEE Signal Process. Mag.*, vol. 28, no. 2, pp. 27-38, Mar. 2011.
- [10] K. Engan, S. O. Aase, and J. Håkon-Husøy, "Method of optimal directions for frame design," in *1999 IEEE Int. Conf. Acoust., Speech, Signal Process. (ICASSP99)*, Phoenix, USA, pp. 2443-2446, Mar. 1999.
- [11] M. Aharon, M. Elad, and A. Bruckstein, "K-SVD: An algorithm for designing overcomplete dictionaries for sparse representation," *IEEE Trans. Signal Process.*, vol. 54, no. 11, pp. 4311-4322, Nov. 2006.
- [12] J. A. Tropp, "Greed is good: Algorithmic results for sparse approximation," *IEEE Trans. Inf. Theory*, vol. 50, no. 10, pp. 2231-2242, Oct. 2004.
- [13] J. A. Tropp and A. Gilbert, "Signal recovery from random measurements via orthogonal matching pursuit," *IEEE Trans. Inf. Theory*, vol. 53, no. 12, pp. 4655-4666, Dec. 2007.
- [14] Q. Jiang, S. Li, Z. Zhu, H. Bai, X. He, and R. C. de Lamare, "Design of compressed sensing system with probability-based prior information," *IEEE Trans. Multimedia*, vol. 22, no. 3, pp. 594-609, Mar. 2020.
- [15] E. J. Candès and M. B. Wakin, "An introduction to compressive sampling," *IEEE Signal Process. Mag.*, vol. 25, no. 2, pp. 21-30, Mar. 2008.
- [16] M. S. Asif and J. Romberg, "Sparse recovery of streaming signals using ℓ_1 -Homotopy," *IEEE Trans. Signal Process.*, vol. 62, no. 16, pp. 4209-4223, Aug. 2014.
- [17] W. Chen, M. R. D. Rodrigues, and I. J. Wassell, "On the use of unit-norm tight frames to improve the average MSE performance in compressive sensing applications," *IEEE Signal Process. Lett.*, vol. 19, no. 1, pp. 8-11, Jan. 2012.
- [18] W. Chen, M. R. D. Rodrigues, and I. J. Wassell, "Projection design for statistical compressive sensing: A tight frame based approach," *IEEE Trans. Signal Process.*, vol. 61, no. 8, pp. 2016-2029, Apr. 2013.
- [19] E. J. Candès and T. Tao, "The Dantzig selector: Statistical estimation when p is much larger than n ," *Ann. Stat.*, vol. 35, no. 6, pp. 2313-2351, 2007.
- [20] Z. Ben-Haim and Y. C. Eldar, "The Cramér-Rao bound for estimating a sparse parameter vector," *IEEE Trans. Signal Process.*, vol. 58, no. 6, pp. 3384-3389, Jun. 2010.
- [21] J. Kovačević and A. Chebira, "Life beyond bases: The advent of frames (Part I)," *IEEE Signal Process. Mag.*, vol. 24, no. 4, pp. 86-104, Jul. 2007.

- [22] J. J. Benedetto and M. Fickus, "Finite normalized tight frames," *Adv. Comput. Math.*, vol. 18, pp. 357-385, Feb. 2003.
- [23] P. G. Casazza and J. Kovačević, "Equal-norm tight frames with erasures," *Adv. Comput. Math.*, vol. 18, pp. 387-430, Feb. 2003.
- [24] P. Zhang, L. Gan, S. Sun, and C. Ling, "Modulated unit-norm tight frames for compressed sensing," *IEEE Trans. Signal Process.*, vol. 63, no. 15, pp. 3974-3985, Aug. 2015.
- [25] R. R. Naidu and C. R. Murthy, "Construction of unimodular tight frames for compressed sensing using majorization-minimization," *Signal Process.*, vol. 172, Article 107516, Jul. 2020.
- [26] D. G. Mixon, "Unit norm tight frames in finite-dimensional spaces." In: K. A. Okoudjou (eds.), *Finite Frame Theory: A Complete Introduction to Overcompleteness*, Proceedings of Symposia in Applied Mathematics, AMS Short Course, San Antonio, Texas, Jan. 2015.
- [27] M. D. Plumbley, T. Blumensath, L. Daudet, R. Gribonval, and M. E. Davies, "Sparse representations in audio and music: From coding to source separation," *Proc. IEEE*, vol. 98, no. 6, pp. 995-1005, Jun. 2010.
- [28] A. Adler, V. Emiya, M. G. Jafari, M. Elad, R. Gribonval, and M. D. Plumbley, "Audio inpainting," *IEEE Trans. Audio Speech Lang. Process.*, vol. 20, no. 3, pp. 922-932, Mar. 2012.
- [29] F. Lieb and H.-G. Stark, "Audio inpainting: Evaluation of time-frequency representations and structured sparsity approaches," *Signal Process.*, vol. 153, pp. 291-299, Dec. 2018.
- [30] O. Mokry and P. Rajmic, "Audio inpainting: Revisited and reweighted," *IEEE-ACM Trans. Audio Speech Lang.*, vol. 28, pp. 2906-2918, Oct. 2020.
- [31] M. Benzi, "Preconditioning techniques for large linear systems: A survey," *J. Comput. Phys.*, vol. 182, no. 2, pp. 418-477, Nov. 2002.
- [32] E. Tsiligianni, L. P. Kondi, and A. K. Katsaggelos, "Preconditioning for underdetermined linear systems with sparse solutions," *IEEE Signal Process. Lett.*, vol. 22, no. 9, pp. 1239-1243, Sept. 2015.
- [33] H. Bai, C. Hong, and X. Li, "Construction of unit-norm tight frame based preconditioner for sparse coding," in *2021 IEEE Int. Conf. Acoust., Speech, Signal Process. (ICASSP 2021)*, Toronto, Canada, pp. 5400-5404, Jun. 2021.
- [34] G. Li, Z. Zhu, D. Yang, L. Chang, and H. Bai, "On projection matrix optimization for compressive sensing systems," *IEEE Trans. Signal Process.*, vol. 61, no. 11, pp. 2887-2898, Jun. 2013.
- [35] H. Bai, S. Li, and X. He, "Sensing matrix optimization based on equiangular tight frames with consideration of sparse representation error," *IEEE Trans. Multimedia*, vol. 18, no. 10, pp. 2040-2053, Oct. 2016.
- [36] R. A. Horn and C. R. Johnson, *Matrix Analysis (second edition)*, Cambridge University Press, 2013.
- [37] C. D. Sigg, T. Dikk, and J. M. Buhmann, "Learning dictionaries with bounded self-coherence," *IEEE Signal Process. Lett.*, vol. 19, no. 12, pp. 861-864, Dec. 2012.
- [38] H. Bai, S. Li, and Q. Jiang, "An efficient algorithm for learning dictionary under coherence constraint," *Math. Probl. Eng.*, Article ID 5737381, 11 pages, 2016.
- [39] M. Gevers and G. Li, *Parametrizations in control, estimation and filtering problems: Accuracy aspects, communication and control engineering series*, Springer, 1993.
- [40] B. G. Bodmann and P. G. Casazza, "The road to equal-norm Parseval frames," *J. Funct. Anal.*, vol. 258, no. 2, pp. 397-420, Jan. 2010.
- [41] P. G. Casazza, M. Fickus, and D. G. Mixon, "Auto-tuning unit norm frames," *Appl. Comput. Harmon. Anal.*, vol. 32, no. 1, pp. 1-15, Jan. 2012.
- [42] X. Li, H. Bai, and B. Hou, "A gradient-based approach to optimization of compressed sensing systems," *Signal Process.*, vol. 139, pp. 49-61, Oct. 2017.
- [43] M. Cooke, J. Barker, S. Cunningham, and X. Shao, "An audio-visual corpus for speech perception and automatic speech recognition," *J. Acoust. Soc. Am.*, vol. 120, no. 5, pp. 2421-2424, Nov. 2006.