
Refining Dual Spectral Sparsity in Transformed Tensor Singular Values

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 The Tensor Nuclear Norm (TNN), derived from the tensor Singular Value Decom-
2 position, is a central low-rank modeling tool that enforces *element-wise sparsity*
3 on frequency-domain singular values and has been widely used in multi-way data
4 recovery for machine learning and computer vision. However, as a direct extension
5 of the matrix nuclear norm, it inherits the assumption of *single-level spectral spar-*
6 *sity*, which strictly limits its ability to capture the *multi-level spectral structures*
7 inherent in real-world data—particularly the coexistence of low-rankness within
8 and sparsity across frequency components. To address this, we propose the tensor
9 ℓ_p -Schatten- q quasi-norm ($p, q \in (0, 1]$), a new metric that enables *dual spectral*
10 *sparsity control* by jointly regularizing both types of structure. While this formula-
11 tion generalizes TNN and unifies existing methods such as the tensor Schatten- p
12 norm and tensor average rank, it differs fundamentally in modeling principle by
13 coupling global frequency sparsity with local spectral low-rankness. This coupling
14 introduces significant theoretical and algorithmic challenges. To tackle these chal-
15 lenges, we provide a theoretical characterization by establishing the first minimax
16 error bounds under dual spectral sparsity, and an algorithmic solution by designing
17 an efficient reweighted optimization scheme tailored to the resulting nonconvex
18 structure. Numerical experiments demonstrate the effectiveness of our method in
19 modeling complex multi-way data.

20 1 Introduction

21 Modeling latent structural patterns in high-dimensional signals is a fundamental challenge across
22 domains such as machine learning and signal processing [17, 38, 19]. Real-world datasets are often
23 inherently multi-modal and high-dimensional (tensor-form), containing intricate dependencies that
24 cannot be adequately captured by naïve modeling or vector/matrix-based representations [4]. A
25 common strategy to uncover these relationships is to impose a *low-rank* prior, which isolates essential
26 information and reduces the degrees of freedom, focusing on the principal components of the signal
27 [21, 1]. Traditional tensor decomposition methods, such as CANDECOMP/PARAFAC (CP) [3],
28 Tucker [27], and Tensor Train [23], have been widely used to model tensor signals [4, 16, 8, 34].
29 While effective in certain scenarios, these methods rely on the assumption of intrinsic low-rankness
30 in the *original domain*, which may fail to hold in complex, real-world applications. This limitation
31 has led to the development of *transformed-domain* modeling, where linear transformations like
32 the Discrete Fourier Transform (DFT) are applied to reveal more pronounced low-rank patterns.
33 Within this paradigm, the tensor Singular Value Decomposition (t-SVD) has emerged as a powerful
34 framework with notable success in applications such as image and video analysis [17, 38, 32, 30].

35 Building on the t-SVD framework, the Tensor Nuclear Norm (TNN) has become an extensively
36 adopted regularizer for low-rank tensor modeling [20, 38, 25, 6, 36, 18, 39]. By extending the matrix

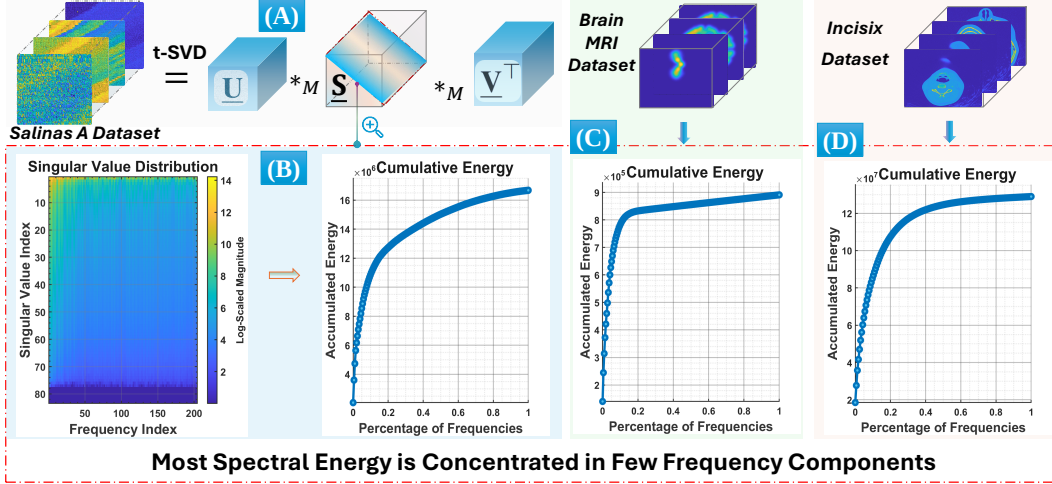


Figure 1: Empirical illustration of dual spectral sparsity patterns in the transformed (DCT) domain via t-SVD. (A) The t-SVD framework decomposes a tensor into frequency-domain singular structures. (B)-Left: Singular value heatmap of the *Salinas A* dataset under t-SVD—each column represents one frequency slice. Vertical decay reveals intra-frequency low-rankness, while horizontal variation indicates sparsity across frequencies. (B)-Right, (C), (D): Cumulative energy curves for *Salinas A*, *Brain MRI*, and *Incisix* datasets show that over 80% of total spectral energy is concentrated in the top 15%–30% frequency components, confirming frequency-wise sparsity. These observations support the presence of a dual-level spectral structure and motivate regularization schemes that go beyond uniform norms like TNN [38, 19] to jointly model frequency sparsity and low-rankness.

37 nuclear norm to the tensor setting, TNN promotes low-rankness by enforcing *element-wise sparsity*
 38 on singular values in the transformed domain [13, 38]. This formulation effectively captures low-rank
 39 dependencies within individual frequency components.

40 However, a long-overlooked limitation of TNN lies in its assumption of *uniform spectral regulariza-*
 41 *tion*, which treats all frequency components equally regardless of their relative importance. From a
 42 signal processing perspective, this *single-level sparsity* design fails to account for the dual-level struc-
 43 ture often observed in transformed tensor data. In particular, real-world tensors may exhibit strong
 44 low-rankness within each frequency component along with sparsity across the frequency domain.
 45 As illustrated in Fig. 1 and further discussed in §3, empirical analyses of several datasets, including
 46 hyperspectral images and medical imaging volumes, indicate that a small subset of frequency slices
 47 contributes the majority of spectral energy. In addition, these dominant components often exhibit
 48 pronounced low-rank structures within each frequency slice. These observations suggest the need
 49 for a more flexible regularization framework that can separately characterize both intra-frequency
 50 low-rankness and inter-frequency sparsity, instead of relying on a uniform scheme like TNN.

51 These limitations necessitate a new method capable of modeling both levels of sparsity. This raises
 52 three interconnected questions:

53 **RQ1 (Modeling):** *how to effectively model both intra-frequency and inter-frequency dependencies in*
 54 *tensor data?*

55 **RQ2 (Theory):** *can we establish rigorous theoretical guarantees to validate such a framework, given*
 56 *the challenges of analyzing coupled sparsity?*

57 **RQ3 (Algorithm):** *can efficient algorithms be designed to tackle the optimization challenges*
 58 *introduced by the coupled sparsity structure?*

59 To address these questions, we propose the *tensor ℓ_p -Schatten- q quasi-norm*, a novel framework
 60 introducing *dual spectral sparsity control* to simultaneously model both within-frequency and across-
 61 frequency dependencies. Specifically, parameter p governs sparsity among different frequency
 62 components (RQ1), while parameter q controls low-rankness within each frequency component. This
 63 framework generalizes and extends TNN, unifying existing methods such as the tensor Schatten- p
 64 quasi-norm [12] and tensor average rank [31] into a single, versatile framework.

While our framework offers promising modeling capabilities, the *coupled nature of this dual spectral sparsity* introduces significant theoretical and computational challenges. Our main contributions in developing and validating this framework are as follows:

- **Structural Modeling (RQ1):** To the best of our knowledge, this work is the first to rigorously formalize and explicitly model a coupled spectral structure within the t-SVD framework, where inter-frequency sparsity coexists with intra-frequency low-rankness (Section 3). The proposed ℓ_p -Schatten- q quasi-norm jointly models both inter-frequency sparsity and intra-frequency low-rankness, while allowing separate control over each via parameters p and q . This formulation captures hierarchical spectral structure beyond uniform regularizers such as TNN.
- **Theoretical Guarantees (RQ2):** We establish sharp minimax lower and upper bounds for tensor estimation under dual spectral sparsity, covering both hard and soft regimes (Section 4). The analysis introduces new techniques to characterize the complexity of coupled parameter spaces, extending classical tools such as covering numbers and metric entropy to the tensor spectral setting.
- **Optimization and Empirical Validation (RQ3):** We develop a scalable proximal algorithm tailored to the proposed quasi-norm (Section 5). It employs a reweighted $\ell_{1/2}$ approximation and frequency-wise singular value updates in the transform domain, effectively handling the nonconvexity and structural coupling induced by dual spectral sparsity. Experiments on real-world tensor recovery tasks demonstrate the potential applicability of our method (Section 6).

The remainder of the paper is organized as follows. Section 2 reviews basic preliminaries. Section 3 introduces the proposed quasi-norm. Sections 4 and 5 present the theoretical analysis and optimization algorithm, respectively. Experimental results are reported in Section 6, followed by the conclusion in Section 7. Details on related work, proofs, algorithms, and experiments are provided in the appendix.

2 Notations and Preliminaries

Notations. For any positive integer d , let $[d] := \{1, \dots, d\}$. We denote vectors by lowercase bold letters (e.g., $\mathbf{a}$), matrices by uppercase bold letters (e.g., $\mathbf{A}$), and 3-way tensors by underlined uppercase letters (e.g., $\underline{\mathbf{A}}$). Constants, represented as c and its variants (e.g., c_1 , C), may vary in value across contexts. For a 3-way tensor of size $d_1 \times d_2 \times m$, we assume $d_1 \geq d_2$ without loss of generality.

For a matrix $\mathbf{A} \in \mathbb{R}^{d_1 \times d_2}$, we define $\sigma(\mathbf{A})$ as the vector of its singular values, arranged in descending order. The spectral norm $\|\mathbf{A}\|_{\text{spec}}$ and nuclear norm $\|\mathbf{A}\|_*$ of $\mathbf{A}$ are defined as the largest and the sum of its singular values, respectively. For any tensor $\underline{\mathbf{A}}$, we define its ℓ_p -norm as $\|\underline{\mathbf{A}}\|_p := \|\text{vec}(\underline{\mathbf{A}})\|_p$ and its Frobenius norm as $\|\underline{\mathbf{A}}\|_F := \|\text{vec}(\underline{\mathbf{A}})\|_2$, where $\text{vec}(\cdot)$ denotes the vectorization operation [11]. The inner product of two tensors $\underline{\mathbf{A}}$ and $\underline{\mathbf{B}}$ is given by $\langle \underline{\mathbf{A}}, \underline{\mathbf{B}} \rangle := \text{vec}(\underline{\mathbf{A}})^\top \text{vec}(\underline{\mathbf{B}})$. For a tensor $\underline{\mathbf{A}} \in \mathbb{R}^{d_1 \times d_2 \times m}$, we denote its i -th frontal slice as $\underline{\mathbf{A}}_{:, :, i}$ or simply $\underline{\mathbf{A}}_i$ when clear from context.

The t-SVD Framework. The t-SVD framework is based on the t-product operation, a generalization of matrix multiplication to tensors, which operates under an invertible linear transform M [9]. By enhancing low-rank properties through specific linear transformations, this approach effectively exploits intrinsic correlations within the data [36, 29]. This paper adopts the convention of using orthogonal matrices for M due to their stability and computational advantages [18, 28]. Specifically, for an orthogonal matrix $\mathbf{M} \in \mathbb{R}^{m \times m}$, we define the M -linear transform and its inverse on a tensor $\underline{\mathbf{T}} \in \mathbb{R}^{d_1 \times d_2 \times m}$ as:

$$M(\underline{\mathbf{T}}) := \underline{\mathbf{T}} \times_3 \mathbf{M}, \quad \text{and} \quad M^{-1}(\underline{\mathbf{T}}) := \underline{\mathbf{T}} \times_3 \mathbf{M}^{-1}, \quad (1)$$

where $\times_3$ denotes the mode-3 tensor-matrix product [9]. Using this transform, we introduce the basic notions in the t-SVD framework.

Definition 2.1 (t-product [9]). The t-product of two tensors $\underline{\mathbf{A}} \in \mathbb{R}^{d_1 \times d_2 \times m}$ and $\underline{\mathbf{B}} \in \mathbb{R}^{d_2 \times d_3 \times m}$ under the transform M in (1) is denoted by $\underline{\mathbf{A}} *_M \underline{\mathbf{B}} = \underline{\mathbf{C}} \in \mathbb{R}^{d_1 \times d_3 \times m}$, where $M(\underline{\mathbf{C}}) = M(\underline{\mathbf{A}}) \odot M(\underline{\mathbf{B}})$ in the transformed domain, and $\odot$ denotes the frontal-slice-wise product of the tensors.

Definition 2.2 (M -block-diagonal matrix [28]). For a tensor $\underline{\mathbf{T}} \in \mathbb{R}^{d_1 \times d_2 \times m}$, its M -block-diagonal matrix $\bar{\mathbf{T}} \in \mathbb{R}^{d_1 m \times d_2 m}$ is defined as

$$\bar{\mathbf{T}} := \text{bdiag}(M(\underline{\mathbf{T}})) = \text{diag}(M(\underline{\mathbf{T}})_{:, :, 1}, \dots, M(\underline{\mathbf{T}})_{:, :, m}),$$

where $M(\underline{\mathbf{T}})$ is the mode-3 transform of $\underline{\mathbf{T}}$, and the operator $\text{bdiag}(\cdot)$ stacks the frontal slices as diagonal blocks.

115 We now formally introduce the t-SVD, as illustrated in Fig. 1-(A).

116 **Definition 2.3** (t-SVD and tensor tubal rank [9]). The tensor Singular Value Decomposition (t-SVD)
 117 of a tensor $\underline{\mathbf{T}} \in \mathbb{R}^{d_1 \times d_2 \times m}$ under the invertible linear transform M in (1) is:

$$\underline{\mathbf{T}} = \underline{\mathbf{U}} *_{\mathbf{M}} \underline{\mathbf{S}} *_{\mathbf{M}} \underline{\mathbf{V}}^{\top}, \quad (2)$$

118 where $\underline{\mathbf{U}} \in \mathbb{R}^{d_1 \times d_1 \times m}$ and $\underline{\mathbf{V}} \in \mathbb{R}^{d_2 \times d_2 \times m}$ are t-orthogonal tensors, and $\underline{\mathbf{S}} \in \mathbb{R}^{d_1 \times d_2 \times m}$ is an
 119 f-diagonal tensor. The tubal rank of $\underline{\mathbf{T}}$ is defined as the number of non-zero tubes in $\underline{\mathbf{S}}$ in the t-SVD,
 120 i.e., $r_{\text{tb}}(\underline{\mathbf{T}}) := \#\{i \mid \underline{\mathbf{S}}_{i,i,:} \neq \mathbf{0}, i \leq \min\{d_1, d_2\}\}$.

121 To further model the low-rank structure of tensors in the transformed domain, the tensor nuclear norm
 122 (TNN) is proposed as a key regularizer in low-rank tensor learning:

123 **Definition 2.4** (Tensor nuclear norm [20]). The tensor nuclear norm (TNN) of a tensor $\underline{\mathbf{T}} \in$
 124 $\mathbb{R}^{d_1 \times d_2 \times m}$ under the transform M are defined as $\|\underline{\mathbf{T}}\|_* := \|\bar{\mathbf{T}}\|_* = \|\sigma(\bar{\mathbf{T}})\|_1$.

125 In this definition, TNN captures *the element-wise sparsity of the transformed spectrum* $\sigma(\bar{\mathbf{T}}) \in$
 126 $\mathbb{R}^{m \times \min\{d_1, d_2\}}$, allowing it to promote low-rank characteristics in the spectral domain. This property
 127 has made TNN a foundational tool in tensor analysis, particularly for low-rank tensor recovery in
 128 various applications such as image inpainting [19].

129 3 Dual Spectral Sparsity in the t-SVD Framework

130 Effectively capturing both intra-frequency low-rankness and inter-frequency sparsity (**RQ1**) is es-
 131 sential for modeling structured tensor data. While methods like TNN emphasize within-frequency
 132 low-rankness, they overlook sparsity across frequencies, limiting their ability to represent hierar-
 133 chical dependencies. To overcome this, we introduce the ℓ_p -Schatten- q quasi-norm, a dual-sparsity
 134 regularization framework designed to capture both levels of structure in a unified way.

135 **Limitations of TNN from a Group Sparsity Perspective.** According to Definition 2.4, the tensor
 136 nuclear norm (TNN) promotes low-rankness by enforcing element-wise sparsity on singular values in
 137 the transformed domain, effectively capturing intra-frequency low-rank structures. However, it applies
 138 uniform regularization across all frequency components, regardless of their spectral importance. This
 139 design fails to exploit the potential sparsity across frequency slices that is often present in real-world
 140 tensors. Fig. 1 presents empirical evidence from three representative datasets—*Salinas A*¹, *Brain MRI*
 141 [33], and *Incisix* [5]—demonstrating that only a small portion of frequency components accounts
 142 for the majority of spectral energy. Specifically, more than 80% of the energy is concentrated in the
 143 top 15%–30% of frequency bands. Meanwhile, the singular value heatmap (Fig. 1(B)-Left) reveals
 144 pronounced horizontal sparsity, indicating that many frequency slices contribute minimally. Within
 145 each active frequency slice, singular values decay rapidly, confirming low-rankness.

146 These observations suggest a dual-level structure comprising inter-frequency sparsity and intra-
 147 frequency low-rankness. From a group sparsity perspective, the spectrum $\sigma(\bar{\mathbf{T}})$ can be partitioned
 148 into groups, where each group corresponds to the singular values $\sigma(M(\underline{\mathbf{T}})_{:,i,:})$ of a specific frequency
 149 slice. TNN enforces uniform regularization across these groups, overlooking their heterogeneous
 150 importance. As a result, it may underperform when modeling data with hierarchical spectral structures.
 151 These limitations motivate a more expressive framework that separately accounts for both levels of
 152 structure.

153 **Hard Dual Spectral Sparsity.** To address the limitations of TNN, we first define a hard dual spectral
 154 sparsity structure, where the tensor is assumed to satisfy exact sparsity constraints across and within
 155 frequency components. This serves as an idealized formulation that captures the extreme case of dual
 156 spectral sparsity and provides a clean theoretical foundation for later analysis.

157 **Definition 3.1** (Hard Dual Spectral Sparsity). A tensor $\underline{\mathbf{T}} \in \mathbb{R}^{d_1 \times d_2 \times m}$ is said to exhibit (s, r) -dual
 158 sparsity under a linear transform M if it satisfies two constraints:

159 **I. Inter-frequency sparsity:** The number of active frequency components is limited to at most s .
 160 Specifically, only s out of the m frequency components can have non-zero singular value vectors:
 161 $\sum_{i=1}^m \mathbb{I}(\sigma(M(\underline{\mathbf{T}})_{:,i,:}) \neq \mathbf{0}) \leq s$, where $\sigma(M(\underline{\mathbf{T}})_{:,i,:})$ denotes the singular value vector of the i -th
 162 frontal slice in the transformed domain.

¹https://www.ehu.eus/ccwintco/index.php?title=Hyperspectral_Remote_Sensing_Scenes

163 **II. Intra-frequency low-rankness:** Within each active frequency component, the number of non-zero
 164 singular values is constrained to at most r . This condition ensures a low-rank structure for
 165 each frequency slice ($\forall i \in [m]$): $\sum_{j=1}^{\min\{d_1, d_2\}} \mathbb{I}(\sigma_j(M(\mathbf{T})_{::,i}) \neq 0) \leq r$, where $\sigma_j(M(\mathbf{T})_{::,i})$
 166 denotes the j -th singular value of the i -th frontal slice of $M(\mathbf{T})$.

167 This definition captures a strict form of dual-level structure by simultaneously enforcing sparsity
 168 across frequencies and low-rankness within each active frequency slice. While such hard constraints
 169 may be too restrictive in practical scenarios, especially where spectral contributions decay grad-
 170 ually, they provide a clear conceptual framework to motivate and analyze the more flexible soft
 171 regularization.

172 **Soft Dual Spectral Sparsity.** While the hard dual spectral sparsity model provides a clean conceptual
 173 foundation, its strict assumption of exact sparsity and fixed-rank constraints is often impractical in
 174 real-world scenarios. In many cases, singular values decay gradually rather than drop abruptly to
 175 zero, and the true number of active frequency components may be ambiguous or noise-sensitive. To
 176 overcome these limitations, we introduce a soft relaxation that allows for approximate sparsity and
 177 low-rankness in a continuous manner. Specifically, we propose the ℓ_p -Schatten- q quasi-norm, which
 178 relaxes the hard dual-sparsity constraints into a soft dual spectral sparsity framework.

179 **Definition 3.2** (Tensor ℓ_p -Schatten- q quasi-norm). For a tensor $\mathbf{T} \in \mathbb{R}^{d_1 \times d_2 \times m}$, we define its tensor
 180 ℓ_p -Schatten- q quasi-norm (abbreviated as $\ell_p(S_q)$ -norm) as:

$$\|\mathbf{T}\|_{\ell_p(S_q)} := \left(\sum_{i=1}^m \left(\sum_{j=1}^{d_1 \wedge d_2} \sigma_j(M(\mathbf{T})_{::,i})^q \right)^{\frac{p}{q}} \right)^{\frac{1}{p}}, \quad (3)$$

181 where the exponents $(p, q) \in (0, 1]^2$.

182 In this quasi-norm, p governs the inter-frequency sparsity by promoting a group-wise regularization
 183 across frequency components, effectively highlighting significant groups while suppressing others.
 184 Simultaneously, q controls the intra-frequency low-rankness by encouraging sparsity in the singular
 185 values within each frequency slice, thereby modeling the intrinsic low-rank structure of the data.
 186 This soft dual spectral sparsity framework provides a unified yet versatile approach to address the
 187 hierarchical complexity of tensor data.

188 The ℓ_p -Schatten- q quasi-norm encompasses several existing regularization methods: it recovers TNN
 189 when $(p, q) = (1, 1)$ [20], approximates the average rank as $(p, q) \rightarrow (1, 0)$ [31], and reduces to the
 190 tensor Schatten- q norm when $p = q$ [12], thereby offering greater modeling flexibility. *Despite*
 191 *generalizing these regularizers, it fundamentally differs by jointly enforcing global frequency sparsity*
 192 *and local spectral low-rankness.*

193 While TNN applies uniform regularization across all singular values, the ℓ_p -Schatten- q quasi-norm
 194 introduces dual spectral sparsity control, modeling both inter-frequency sparsity through the ℓ_p -quasi-
 195 norm and intra-frequency low-rankness via the Schatten- q quasi-norm. This dual-level flexibility
 196 makes the proposed framework particularly well-suited for hierarchical and multi-scale data, where
 197 dependencies and sparsity exhibit layered structures. By bridging the gap between element-wise
 198 sparsity (as in TNN) and structured group sparsity, the ℓ_p -Schatten- q quasi-norm offers a more
 199 expressive and adaptable approach, enabling precise control over structural patterns in modern
 200 tensor-based analysis and recovery tasks.

201 4 Theory of Dual Spectral Sparse Tensor Estimation

202 This section develops the theoretical foundations of tensor estimation with dual spectral sparsity
 203 structures (**RQ2**).

204 **Challenges.** Dual spectral sparsity, combining inter-frequency sparsity with intra-frequency low-
 205 rankness, leads to a *globally coupled structure* that fundamentally differs from classical decoupled
 206 models like TNN. The ℓ_p -Schatten- q quasi-norm imposes interdependent constraints across frequency
 207 slices, resulting in a highly non-convex parameter space with nested sparsity patterns. This coupling
 208 prohibits slice-wise decomposition and complicates the use of standard tools. Accurately characteriz-
 209 ing the estimation complexity demands novel extensions of covering numbers and metric entropy
 210 that jointly capture discrete sparsity and continuous low-rank structure.

To understand the statistical limits of learning under dual spectral sparsity, we analyze a simplified but representative model: the Gaussian location model, where the observed tensor is corrupted by additive noise. This setting preserves the core structural properties—inter-frequency sparsity and intra-frequency low-rankness—while avoiding complications unrelated to sparsity itself. Within this framework, we define structured parameter spaces that capture hard and soft variants of dual spectral sparsity, and establish sharp minimax lower and upper bounds under each. These results reveal how the joint effects of frequency selection and within-slice spectral decay determine the fundamental estimation limits, and provide theoretical justification for our proposed regularization.

4.1 Gaussian Location Model

Consider the Gaussian location model (GLM) [14], where n independent noisy realizations of the target tensor $\mathbf{L}^* \in \mathbb{R}^{d_1 \times d_2 \times m}$ are observed as:

$$\mathbf{Y}_i = \mathbf{L}^* + \mathbf{E}_i, \quad i \in [n], \quad (4)$$

where $\mathbf{Y}_i \in \mathbb{R}^{d_1 \times d_2 \times m}$ is the observed tensor, $\mathbf{L}^*$ represents the ground truth tensor of interest, and $\mathbf{E}_i \in \mathbb{R}^{d_1 \times d_2 \times m}$ denotes the noise tensor with entries independently drawn from $\mathcal{N}(0, \sigma^2)$. The parameter σ characterizes the noise level. To simplify the analysis, we consider the sample mean of observations $\bar{\mathbf{Y}} = n^{-1} \sum_{i=1}^n \mathbf{Y}_i = \mathbf{L}^* + \bar{\mathbf{E}}$, where $\bar{\mathbf{E}} = n^{-1} \sum_{i=1}^n \mathbf{E}_i$ is the aggregated noise tensor with entries independently distributed as $\mathcal{N}(0, \sigma^2/n)$. The goal is to estimate the ground truth tensor $\mathbf{L}^*$ based on the noisy observations $\{\mathbf{Y}_i\}_{i=1}^n$. In particular, we aim to recover $\mathbf{L}^*$ under dual spectral sparsity assumptions.

Remark 4.1. We adopt the Gaussian location model to isolate the core effects of dual spectral sparsity and the ℓ_p -Schatten- q regularization, avoiding additional complications from design tensors or sampling operators in tensor regression [35, 29, 24]. This simplified setting enables cleaner analysis and yields insights that extend naturally to regression problems under standard conditions such as RIP [35] or RSC [29, 24, 22].

Dual Spectral Sparsity Assumptions. We consider three distinct sparsity models for $\mathbf{L}^*$:

A1. Hard dual spectral sparsity: Let $\mathbf{L}^*$ belong to the parameter space

$$\mathbf{T}_{0,0}(s, r) = \{\mathbf{L} : \text{at most } s \text{ active frequency slices, each of rank at most } r\}. \quad (5)$$

This model enforces exact inter-frequency sparsity and intra-frequency low-rankness.

A2. Hard frequency sparsity and soft rank constraint (hard-soft sparsity): Let $\mathbf{L}^*$ lie in

$$\mathbf{T}_{0,q}(s, R) = \left\{ \mathbf{L} : |\{i : M(\mathbf{L})_{:, :, i} \neq \mathbf{0}\}| \leq s, \|M(\mathbf{L})_{:, :, i}\|_{S_q}^q \leq R, \forall i \in [m] \right\}. \quad (6)$$

This space imposes hard inter-frequency sparsity and soft Schatten- q constraints within each active slice.

A3. Soft dual spectral sparsity: Let $\mathbf{L}^*$ belong to the parameter space

$$\mathbf{T}_{p,q}(R) = \left\{ \mathbf{L} : \|\mathbf{L}\|_{\ell_p(S_q)}^p \leq R \right\}. \quad (7)$$

Here, p promotes inter-frequency sparsity and q controls intra-frequency low-rankness via spectral decay; R specifies the quasi-norm ball radius.

These parameter spaces offer different views on structured tensor estimation: the *hard sparsity* model enforces strict thresholds, the *hard-soft model* balances structure with adaptability, and the *fully soft model* captures gradual spectral decay. Our goal is to estimate $\mathbf{L}^*$ and derive minimax bounds under these assumptions.

4.2 Minimax Risk over Dual-level Sparse Structures

A key theoretical question in high-dimensional tensor estimation is: *What are the fundamental limits for recovering a tensor with dual spectral sparsity from noisy observations?* To address this, we establish minimax lower and upper bounds that characterize the best possible estimation accuracy achievable by any estimator under dual spectral sparsity assumptions.

$$\mathfrak{M}(\mathbf{T}) = \inf_{\hat{\mathbf{L}}} \sup_{\mathbf{L}^* \in \mathbf{T}} \mathbb{E} \left[\|\hat{\mathbf{L}} - \mathbf{L}^*\|_{\mathbb{F}}^2 \right], \quad (8)$$

where $\mathbf{T}$ is the parameter space. Following [18, 19], we consider $d_1 = d_2 = d$ for simplicity.

Theorem 4.2 (Minimax Bounds). *The minimax risk under dual spectral sparsity satisfies the following bounds under certain conditions²:*

I. *Hard constraints on both frequency sparsity and per-slice low-rankness:*

$$\mathfrak{M}(\mathbf{T}_{0,0}(s, r)) \asymp \frac{\sigma^2}{n} \left(s \log \frac{em}{s} + srd \right).$$

II. *Hard frequency sparsity with soft intra-slice Schatten- q constraints:*

$$\mathfrak{M}(\mathbf{T}_{0,q}(s, R)) \asymp \frac{\sigma^2}{n} s \log \frac{em}{s} + sR \left(\frac{\sigma^2}{n} d \right)^{1-\frac{q}{2}}.$$

III. *Soft $\ell_p(S_q)$ constraints over both frequency and rank dimensions:*

$$\mathfrak{M}(\mathbf{T}_{p,q}(R)) \asymp \begin{cases} R \left(\frac{\sigma^2 n}{d} \right)^{\frac{p-2}{2}} + R \left(\frac{\sigma^2 n}{\log m} \right)^{\frac{p-2}{2}}, & p > q, \\ R^{\frac{q}{p}} \left(\frac{\sigma^2 n}{d} \right)^{\frac{q-2}{2}} + R \left(\frac{\sigma^2 n}{\log m} \right)^{\frac{p-2}{2}}, & p \leq q, m > d^2, \\ R^{\frac{q}{p}} \left(\frac{\sigma^2 n}{d} \right)^{\frac{q-2}{2}}, & p \leq q, m \leq d^2. \end{cases}$$

Theorem 4.2 establishes the fundamental limits of estimation accuracy under different dual spectral sparsity structures. The minimax risk quantifies the worst-case squared Frobenius norm error that any estimator must incur when recovering a structured tensor from noisy observations. The results reveal the intricate balance between inter-frequency sparsity and intra-frequency low-rankness, showing how these factors jointly govern estimation complexity:

I. In the *hard sparsity* case, the estimation error consists of two terms: (i) $s \log(em/s)$, which reflects the difficulty of selecting s active frequency components, and (ii) srd , which characterizes the challenge of estimating rank- r matrices within each component.

II. In the *hard-soft sparsity* setting, the second term adapts to $sR(n^{-1}d)^{1-q/2}$, incorporating a smoother spectral decay controlled by q . Smaller q values impose stronger low-rank constraints, effectively reducing estimation complexity by promoting more aggressive rank sparsity.

III. In the *fully soft sparsity* scenario, where both inter-frequency sparsity and intra-frequency rank constraints are relaxed, the minimax risk follows distinct scaling behaviors across regimes. When $p > q$, the error rate is dominated by ℓ_p sparsity, with S_q low-rankness playing a minor role. For $p \leq q$ and $m \geq d^2$, both the ℓ_p -ball and S_q -ball influence the estimation error, demonstrating an interplay between structured sparsity and low-rank regularization. When $m \leq d^2$, the error rate is dictated by S_q , making it independent of m , emphasizing the fundamental role of rank constraints in this regime.

5 Optimization for Dual Spectral Sparse Tensor Estimation

Efficiently solving tensor estimation problems with dual spectral sparsity (**RQ3**) is key to leveraging the proposed ℓ_p -Schatten- q quasi-norm in practice. However, this task presents substantial challenges due to the non-convexity and coupled structure of this regularization.

Challenges. Even in the vector setting, optimizing dual-level sparse structures is notoriously difficult due to the combination of *non-convexity* and *structural coupling* [7, 15]. In our tensor case, these challenges are further compounded by the need to simultaneously enforce inter-frequency sparsity and intra-frequency low-rankness. Most existing tensor optimization methods either treat frequency components independently or impose low-rank constraints without spectral sparsity considerations, making them ill-suited for the proposed dual-spectral regularization. The ℓ_p -Schatten- q quasi-norm is non-convex whenever $p, q \in (0, 1]$, ruling out standard convex optimization techniques and necessitating a structure-aware, non-convex optimization strategy.

To address these difficulties, our approach is naturally motivated by the structural properties of the problem. We adopt a *proximal update scheme* that takes advantage of the separability of the

²The conditions in each setting are provided in the appendix.

transform-domain representation $M(\underline{\mathbf{L}})$, allowing frequency-wise updates, along with an iterative reweighting strategy that facilitates optimization in the presence of non-convex regularization.

Proximal Operator Formulation. To handle the non-convex ℓ_p -Schatten- q regularization, we adopt a proximal update scheme that enforces dual spectral sparsity while remaining computationally efficient. Specifically, at iteration t , the update is given by solving:

$$\underline{\mathbf{L}}^{t+1} \in \arg \min_{\underline{\mathbf{L}}} \frac{1}{2} \|\underline{\mathbf{L}} - \underline{\mathbf{Z}}\|_{\text{F}}^2 + \lambda \sum_{k=1}^m \|M(\underline{\mathbf{L}})_{::,k}\|_{S_q}^{p/q}, \quad (9)$$

where $\underline{\mathbf{Z}}$ denotes the intermediate variable aggregating previous updates and gradient information. Since the transform $M(\cdot)$ allows slice-wise decomposition [10], Problem (9) reduces to m subproblems over frequency components $k \in [m]$:

$$\min_{\mathbf{A}_k} \frac{1}{2} \|\mathbf{A}_k - M(\underline{\mathbf{Z}})_{::,k}\|_{\text{F}}^2 + \lambda \|\mathbf{A}_k\|_{S_q}^{p/q}, \quad (10)$$

where $\mathbf{A}_k := M(\underline{\mathbf{L}})_{::,k}$ denotes the k -th frontal slice of the transformed tensor $M(\underline{\mathbf{L}})$. Problem (10) is difficult due to the non-convexity and lack of smoothness of the Schatten- q quasi-norm, which admits no closed-form or standard proximal solution in general.

To efficiently approximate Problem (10), we adopt a reweighted $\ell_{1/2}$ -surrogate for $\|\mathbf{A}_k\|_{S_q}^{p/q}$ based on singular values:

$$\sum_{i=1}^d w_{i,k} \cdot \sigma_i(\mathbf{A}_k)^{1/2}, \quad (11)$$

with weights defined as $w_{i,k} = (\sum_{j=1}^d \varsigma_{j,k}^q + \epsilon)^{p/q-1} \cdot (\varsigma_{i,k}^{1/2} + \epsilon)^{2q-1}$, where ϵ is a small regularization constant and $\varsigma_{j,k} := \sigma_j(M(\underline{\mathbf{L}}^t)_{::,k})$ are the singular values from the previous iterate. The update for each singular value then becomes a soft-thresholding step:

$$\sigma_i^{(t+1)}(M(\underline{\mathbf{L}})_{::,k}) = \mathcal{S}_{\lambda w_{i,k}}^{\ell_{1/2}}(\sigma_i(M(\underline{\mathbf{Z}})_{::,k})), \quad (12)$$

where $\mathcal{S}^{\ell_{1/2}}$ is the proximal operator for the $\ell_{1/2}$ -norm (see Appendix for closed-form expression).

After singular value shrinkage, we reconstruct each slice $M(\underline{\mathbf{L}}^{t+1})_{::,k} = \mathbf{U}_k \cdot \text{diag}(\boldsymbol{\sigma}^{(t+1)}) \cdot \mathbf{V}_k^\top$, where $\mathbf{U}_k$ and $\mathbf{V}_k$ are from the SVD of $M(\underline{\mathbf{Z}})_{::,k}$. Finally, applying the inverse transform yields the updated tensor $\underline{\mathbf{L}}^{t+1}$ in the original domain.

6 Experiments

Having established the theoretical foundations and algorithmic framework, we now evaluate the empirical performance of the proposed ℓ_p -Schatten- q quasi-norm in tensor estimation tasks. We conduct extensive experiments on three types of remote sensing data to demonstrate its effectiveness in noisy tensor completion tasks.

Experimental Setup. We consider the noisy tensor completion which involves reconstructing a tensor from noisy incomplete observations. Given a clean tensor $\underline{\mathbf{L}}$ of size $d_1 \times d_2 \times d_3$, we introduce *i.i.d.* Gaussian noise with standard deviation $\sigma = c\sigma_0$, where $c = 0.05$ and $\sigma_0 = \|\underline{\mathbf{L}}\|_{\text{F}} / \sqrt{d_1 d_2 d_3}$. A uniform sampling strategy is applied with sampling ratios $p \in \{0.05, 0.1, 0.15\}$, meaning that 95%, 90%, and 85% of the entries are missing, respectively. Each setting is tested over 10 trials, and the averaged PSNR (dB) and SSIM values are reported. To benchmark our method, we compare the proposed $\ell_p(S_q)$ -quasi-norm against several low-rank regularizers, including matrix nuclear norm (NN) [2], Tucker-based tensor nuclear norm (SNN) [16], TNN-DFT [37], TNN-DCT [20], tensor k -Support norm (k -Supp) ($k = 2$) [29], tensor ℓ_{1-2} -norm (ℓ_{1-2}) [26], tensor Schatten- p -norm ($p = 1/2$) [12]. In our implementation, we set the sparsity parameters³ to $(p, q) = (0.8961, 0.8966)$ and employ the Discrete Cosine Transform (DCT) as the transform operator $M(\cdot)$. Details of the experiments are given in the appendix.

³We first performed a coarse grid search over $p, q \in \{0.1, 0.2, \dots, 1.0\}$ and observed consistent performance peaks near $p = q = 0.9$. We then manually fine-tuned within $[0.88, 0.92]$ based on PSNR, selecting $(p, q) = (0.8961, 0.8966)$ as the best-performing pair.

Table 1: Results for noisy tensor completion on remote sensing datasets are shown below. The best result in each case is highlighted in **bold**, while the second-best is underlined.

Dataset	SR	Metric	NN	SNN	TNN-DFT	TNN-DCT	k-Supp	ℓ_{1-2}	Schatten-1/2	$\ell_p(S_q)$ (proposed)
SalinasA	5%	PSNR	15.21	20.79	22.55	26.52	22.58	22.21	22.45	28.43
		SSIM	0.2594	0.7547	0.5667	0.7384	0.5689	0.5524	0.4474	<u>0.7374</u>
		PSNR	20.62	25.56	25.72	29.61	25.89	26.14	25.86	31.81
	10%	SSIM	0.4775	0.8284	0.7027	<u>0.8403</u>	0.7231	0.7197	0.6058	0.8484
		PSNR	23.09	27.99	28.06	<u>31.32</u>	28.09	28.13	26.98	33.23
		SSIM	0.5643	0.8622	0.7804	<u>0.8798</u>	0.7810	0.7795	0.6505	0.8830
IndianPines	5%	PSNR	20.44	22.01	25.68	<u>26.26</u>	25.70	25.73	24.68	27.05
		SSIM	0.3895	0.6359	0.6293	<u>0.6727</u>	0.6289	0.6316	0.5361	0.6740
		PSNR	22.23	24.94	27.45	<u>28.40</u>	27.48	27.52	25.72	28.92
	10%	SSIM	0.4836	0.7171	0.7226	0.7744	0.7219	0.7249	0.5991	<u>0.7617</u>
		PSNR	23.52	26.61	28.54	<u>29.52</u>	28.53	28.63	26.24	29.89
		SSIM	0.5438	0.7668	0.7713	0.8177	0.7709	0.7741	0.6258	<u>0.7997</u>
Cloth	5%	PSNR	20.10	20.95	25.00	<u>26.09</u>	25.08	25.09	24.96	26.99
		SSIM	0.3762	0.5096	0.6773	<u>0.7283</u>	0.6792	0.6793	0.6305	0.7422
		PSNR	21.14	22.72	28.00	<u>29.24</u>	28.12	28.14	27.98	30.63
	10%	SSIM	0.4341	0.5983	0.8132	<u>0.8540</u>	0.8143	0.8163	0.7668	0.8658
		PSNR	22.05	24.18	30.03	<u>31.36</u>	30.08	30.11	29.50	32.71
		SSIM	0.4889	0.6783	0.8722	<u>0.9054</u>	0.8727	0.8733	0.8153	0.9090
Hair	5%	PSNR	25.33	30.09	33.16	<u>35.31</u>	33.19	33.27	33.43	36.95
		SSIM	0.7147	0.8631	0.8917	0.9248	0.8921	0.8919	0.8240	<u>0.9196</u>
		PSNR	29.52	33.35	36.22	<u>38.18</u>	36.17	36.30	35.69	39.91
	10%	SSIM	0.8008	0.9122	0.9292	0.9535	0.9286	0.9296	0.8640	<u>0.9517</u>
		PSNR	31.12	35.24	38.00	<u>39.88</u>	37.91	38.07	36.46	41.52
		SSIM	0.8364	0.9336	0.9449	0.9650	0.9442	0.9448	0.8735	<u>0.9641</u>
JellyBeans	5%	PSNR	16.33	18.21	25.43	<u>26.47</u>	25.38	25.62	25.39	27.91
		SSIM	0.2397	0.4942	0.6726	0.7223	0.6714	0.6733	0.5504	<u>0.7115</u>
		PSNR	18.12	22.11	28.50	<u>30.14</u>	28.47	28.67	28.41	31.95
	10%	SSIM	0.3169	0.6629	0.7900	0.8518	0.7902	0.7932	0.6905	<u>0.8486</u>
		PSNR	19.92	24.67	30.51	<u>32.33</u>	30.52	30.61	29.96	33.97
		SSIM	0.4053	0.7592	0.8489	0.9030	0.8504	0.8499	0.7516	<u>0.8980</u>
OSU Thermal	5%	PSNR	13.19	15.83	28.06	<u>27.99</u>	28.01	28.19	28.11	30.06
		SSIM	0.1848	0.4759	0.8584	<u>0.8707</u>	0.8579	0.8603	0.7928	0.8759
		PSNR	14.67	19.75	31.30	<u>31.62</u>	31.28	31.60	30.51	33.67
	10%	SSIM	0.2509	0.6594	0.9151	0.9326	0.9147	0.9168	0.8358	<u>0.9272</u>
		PSNR	16.27	22.52	33.02	<u>33.51</u>	33.05	33.11	30.99	35.09
		SSIM	0.3273	0.7621	0.9315	0.9509	0.9321	0.9318	0.8373	<u>0.9404</u>

Datasets. We validate our approach on three categories of remote sensing data. First, for hyperspectral images, we employ the corrected *Indian Pines* and *Salinas A* datasets from the AVIRIS sensor, containing 200 and 204 spectral bands respectively. Due to computational considerations, we utilize the first 30 bands in our experiments. Second, we evaluate on multispectral images from the Columbia MSI Database, including *Cloth*, *Hair*, and *Jelly Beans*, each with dimensions $512 \times 512 \times 31$ and normalized intensity values in $[0,1]$. Finally, for thermal imaging, we use sequences from the *OSU Thermal Database*, specifically the first 30 frames of Sequence 1, forming a tensor of size $320 \times 240 \times 30$.

Results and Analysis. Table 1 summarizes the PSNR and SSIM results across different missing rates. The proposed $\ell_p(S_q)$ -quasi-norm achieves the highest PSNR, demonstrating its effectiveness in preserving spectral information. Its SSIM results rank among the top two, indicating that our approach better retains structural integrity compared to competing methods. These experimental results demonstrate the effectiveness of the proposed ℓ_p -Schatten- q quasi-norm in robust tensor recovery, showing how characterizing dual spectral sparsity structures in transformed domains benefits tensor reconstruction performance.

7 Conclusion

This paper identifies and formalizes a coupled spectral structure within the t-SVD framework, where inter-frequency sparsity coexists with intra-frequency low-rankness. To capture this structure, we propose a unified modeling approach based on the ℓ_p -Schatten- q quasi-norm, which enables separate control over spectral sparsity at different levels and generalizes existing tensor norms. We provide sharp minimax guarantees under both hard and soft sparsity regimes, and develop an efficient proximal algorithm tailored to this setting. Experimental results demonstrate the practical potential of the proposed approach for structured tensor recovery.

Limitation. To highlight the fundamental properties of the proposed ℓ_p -Schatten- q quasi-norm, our analysis employs several simplifications, including Gaussian location model and idealized sparsity patterns. While our optimization algorithm shows promising empirical performance, its theoretical convergence properties remain to be established. These theoretical and algorithmic limitations suggest important directions for future research.

References

- [1] G. Bergqvist and E. G. Larsson. The higher-order singular value decomposition: Theory and an application [lecture notes]. *IEEE signal processing magazine*, 27(3):151–154, 2010.
- [2] E. J. Candès and T. Tao. The power of convex relaxation: Near-optimal matrix completion. *IEEE Transactions on Information Theory*, 56(5):2053–2080, 2010.
- [3] J. D. Carroll and J. Chang. Analysis of individual differences in multidimensional scaling via an n -way generalization of “Eckart-Young” decomposition. *Psychometrika*, 35(3):283–319, 1970.
- [4] A. Cichocki, N. Lee, I. Oseledets, A. H. Phan, Q. Zhao, and D. P. Mandic. Tensor networks for dimensionality reduction and large-scale optimization: Part 1 low-rank tensor decompositions. *Foundations & Trends® in Machine Learning*, 9(4-5):249–429, 2016.
- [5] S. Gandy, B. Recht, and I. Yamada. Tensor completion and low-n-rank tensor recovery via convex optimization. *Inverse Problems*, 27(2):025010, 2011.
- [6] J. Hou, F. Zhang, H. Qiu, J. Wang, Y. Wang, and D. Meng. Robust low-tubal-rank tensor recovery from binary measurements. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2021.
- [7] Y. Hu, C. Li, K. Meng, J. Qin, and X. Yang. Group sparse optimization via $\ell_{p,q}$ regularization. *Journal of Machine Learning Research*, 18(30):1–52, 2017.
- [8] M. Imaizumi, T. Maehara, and K. Hayashi. On tensor train rank minimization: Statistical efficiency and scalable algorithm. In *Advances in Neural Information Processing Systems*, pages 3930–3939, 2017.
- [9] E. Kernfeld, M. Kilmer, and S. Aeron. Tensor–tensor products with invertible linear transforms. *Linear Algebra and its Applications*, 485:545–570, 2015.
- [10] M. E. Kilmer, K. Braman, et al. Third-order tensors as operators on matrices: A theoretical and computational framework with applications in imaging. *SIAM J MATRIX ANAL A*, 34(1):148–172, 2013.
- [11] T. G. Kolda and B. W. Bader. Tensor decompositions and applications. *SIAM Review*, 51(3):455–500, 2009.
- [12] H. Kong, X. Xie, and Z. Lin. t-Schatten- p norm for low-rank tensor recovery. *IEEE Journal of Selected Topics in Signal Processing*, 12(6):1405–1419, 2018.
- [13] C. Li, W. He, L. Yuan, Z. Sun, and Q. Zhao. Guaranteed matrix completion under multiple linear transformations. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 11136–11145, 2019.
- [14] Z. Li, Y. Zhang, and J. Yin. Estimating double sparse structures over $\ell_u(\ell_q)$ -balls: Minimax rates and phase transition. *IEEE Transactions on Information Theory*, 2024.
- [15] R. Lin, S. Chen, H. Feng, and Y. Liu. Computing the proximal operator of the $\ell_{p,q}$ -norm for group sparsity. *arXiv preprint arXiv:2409.14156*, 2024.
- [16] J. Liu, P. Musialski, P. Wonka, and J. Ye. Tensor completion for estimating missing values in visual data. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 35(1):208–220, 2013.
- [17] X. Liu, S. Aeron, V. Aggarwal, and X. Wang. Low-tubal-rank tensor completion using alternating minimization. *IEEE Transactions on Information Theory*, 66(3):1714–1737, 2020.
- [18] C. Lu. Transforms based tensor robust PCA: Corrupted low-rank tensors recovery via convex optimization. In *ICCV*, pages 1145–1152, 2021.
- [19] C. Lu, J. Feng, W. Liu, Z. Lin, S. Yan, et al. Tensor robust principal component analysis with a new tensor nuclear norm. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2019.
- [20] C. Lu, X. Peng, and Y. Wei. Low-rank tensor completion with a new tensor nuclear norm induced by invertible linear transforms. In *CVPR*, pages 5996–6004, 2019.
- [21] C. D. Martin, R. Shafer, and B. Larue. An order- p tensor factorization with applications in imaging. *SIAM Journal on Scientific Computing*, 35(1), 2013.
- [22] S. Negahban and M. J. Wainwright. Restricted strong convexity and weighted matrix completion: Optimal bounds with noise. *The Journal of Machine Learning Research*, 13:1665–1697, 2012.

- 402 [23] I. V. Oseledets. Tensor-train decomposition. *SIAM Journal on Scientific Computing*, 33(5):2295–2317,
403 2011.
- 404 [24] Y. Qiu, G. Zhou, A. Wang, Q. Zhao, and S. Xie. Balanced unfolding induced tensor nuclear norms for
405 high-order tensor completion. *IEEE Transactions on Neural Networks and Learning Systems*, 2024.
- 406 [25] G. Song, M. K. Ng, and X. Zhang. Robust tensor completion using transformed tensor singular value
407 decomposition. *NUMER LINEAR ALGEBR*, 27(3):e2299, 2020.
- 408 [26] Z. Tan, L. Huang, H. Cai, and Y. Lou. Non-convex approaches for low-rank tensor completion under tubal
409 sampling. In *ICASSP*, pages 1–5. IEEE, 2023.
- 410 [27] L. R. Tucker. Some mathematical notes on three-mode factor analysis. *Psychometrika*, 31(3):279–311,
411 1966.
- 412 [28] A. Wang, C. Li, M. Bai, Z. Jin, G. Zhou, and Q. Zhao. Transformed low-rank parameterization can help
413 robust generalization for tensor neural networks. *NeurIPS*, 36, 2023.
- 414 [29] A. Wang, G. Zhou, Z. Jin, and Q. Zhao. Tensor recovery via $*_l$ -spectral k -support norm. *IEEE Journal of*
415 *Selected Topics in Signal Processing*, 15(3):522–534, 2021.
- 416 [30] H. Wang, J. Yang, X. Yu, Y. Zhang, J. Qian, and J. Wang. Tensor-flamingo unravels the complexity of
417 single-cell spatial architectures of genomes at high-resolution. *Nature Communications*, 16(1):3435, 2025.
- 418 [31] Z. Wang, J. Dong, X. Liu, and X. Zeng. Low-rank tensor completion by approximating the tensor average
419 rank. In *ICCV*, pages 4612–4620, 2021.
- 420 [32] Y. Xie, D. Tao, W. Zhang, Y. Liu, L. Zhang, and Y. Qu. On unifying multi-view self-representations for
421 clustering by tensor multi-rank minimization. *International Journal of Computer Vision*, 126(11):1157–
422 1179, 2018.
- 423 [33] Y. Xu, R. Hao, W. Yin, and Z. Su. Parallel matrix factorization for low-rank tensor completion. *Inverse*
424 *Problems and Imaging*, 9(2):601–624, 2015.
- 425 [34] M. Yuan and C. H. Zhang. On tensor completion via nuclear norm minimization. *Foundations of*
426 *Computational Mathematics*, 16(4):1–38, 2016.
- 427 [35] F. Zhang, W. Wang, J. Huang, J. Wang, and Y. Wang. RIP-based performance guarantee for low-tubal-rank
428 tensor recovery. *Journal of Computational and Applied Mathematics*, 374:112767, 2020.
- 429 [36] X. Zhang and M. K.-P. Ng. Low rank tensor completion with poisson observations. *IEEE Transactions on*
430 *Pattern Analysis and Machine Intelligence*, 2021.
- 431 [37] Z. Zhang and S. Aeron. Exact tensor completion using t-svd. *IEEE Transactions on Signal Processing*,
432 65(6):1511–1526, 2017.
- 433 [38] Z. Zhang, G. Ely, S. Aeron, et al. Novel methods for multilinear data completion and de-noising based on
434 tensor-SVD. In *CVPR*, pages 3842–3849, 2014.
- 435 [39] P. Zhou and J. Feng. Outlier-robust tensor PCA. In *Proceedings of the IEEE Conference on Computer*
436 *Vision and Pattern Recognition*, 2017.

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes], [No], or [NA].
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]", it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

IMPORTANT, please:

- **Delete this instruction block, but keep the section heading “NeurIPS Paper Checklist”.**
- **Keep the checklist subsection headings, questions/answers and guidelines below.**
- **Do not modify the questions and only use the provided macros for your answers.**

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [Yes]

Justification: The abstract and introduction clearly state the central focus of the paper—the modeling, theoretical analysis, and algorithmic solution for dual spectral sparsity in tensors—which directly correspond to the three core contributions developed in the main body.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: The paper includes a dedicated discussion on limitations in the conclusion section, noting simplifications such as the use of the Gaussian location model and the lack of theoretical convergence guarantees for the proposed non-convex optimization algorithm.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: For each theoretical result, the paper states all necessary assumptions explicitly and provides complete proofs in the appendix. The analysis includes both lower and upper minimax bounds under different dual-sparsity regimes, supported by standard and extended techniques such as entropy numbers and packing arguments.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: The paper provides sufficient information to reproduce the main experimental results, including dataset descriptions, sampling settings, noise levels, baseline configurations, and implementation details of the proposed algorithm. Additional algorithmic and parameter details are included in the appendix to support reproducibility.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: The authors have uploaded the code as part of the supplementary material, along with sufficient implementation details and instructions to reproduce the main experimental results. All datasets used are publicly available and clearly specified.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).

- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The paper specifies all necessary details, including sampling ratios, noise levels, hyperparameter settings, evaluation metrics (PSNR and SSIM), and implementation choices. Additional settings such as initialization and convergence criteria are provided in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: The paper reports averaged performance over multiple random trials for each experimental setting. While explicit error bars are not shown in tables, standard practice is followed by reporting stable metrics (PSNR and SSIM) under fixed noise and sampling ratios, ensuring reliable comparative evaluation.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).

- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: The appendix provides information on the computational environment used for the experiments, including hardware specifications such as CPU type and memory.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The research does not involve human subjects, personally identifiable information, sensitive data, or potentially harmful applications. It focuses on theoretical modeling and algorithmic development for tensor estimation, aligning fully with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: The paper focuses solely on the theoretical modeling and optimization of the tensor recovery algorithm, without involving any societal issues, and therefore does not require discussion of societal impacts.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to

generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.

- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This work does not involve the release of high-risk models or datasets. It focuses on a general-purpose tensor estimation framework using publicly available data, and does not include pretrained models or components with significant misuse potential.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: All external datasets and baseline implementations used in the paper are publicly available and properly cited. Their licenses and terms of use have been respected, and relevant references are provided in the main text and appendix.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

750 • If this information is not available online, the authors are encouraged to reach out to
751 the asset’s creators.

752 **13. New assets**

753 Question: Are new assets introduced in the paper well documented and is the documentation
754 provided alongside the assets?

755 Answer: [\[Yes\]](#)

756 Justification: The paper introduces a new optimization algorithm and associated code,
757 which are included in the supplementary material. The code is well documented with
758 clear instructions, comments, and reproducibility guidelines to support independent use and
759 verification.

760 Guidelines:

- 761 • The answer NA means that the paper does not release new assets.
- 762 • Researchers should communicate the details of the dataset/code/model as part of their
763 submissions via structured templates. This includes details about training, license,
764 limitations, etc.
- 765 • The paper should discuss whether and how consent was obtained from people whose
766 asset is used.
- 767 • At submission time, remember to anonymize your assets (if applicable). You can either
768 create an anonymized URL or include an anonymized zip file.

769 **14. Crowdsourcing and research with human subjects**

770 Question: For crowdsourcing experiments and research with human subjects, does the paper
771 include the full text of instructions given to participants and screenshots, if applicable, as
772 well as details about compensation (if any)?

773 Answer: [\[NA\]](#)

774 Justification: This work does not involve any human subjects or crowdsourcing experiments.

775 Guidelines:

- 776 • The answer NA means that the paper does not involve crowdsourcing nor research with
777 human subjects.
- 778 • Including this information in the supplemental material is fine, but if the main contribu-
779 tion of the paper involves human subjects, then as much detail as possible should be
780 included in the main paper.
- 781 • According to the NeurIPS Code of Ethics, workers involved in data collection, curation,
782 or other labor should be paid at least the minimum wage in the country of the data
783 collector.

784 **15. Institutional review board (IRB) approvals or equivalent for research with human
785 subjects**

786 Question: Does the paper describe potential risks incurred by study participants, whether
787 such risks were disclosed to the subjects, and whether Institutional Review Board (IRB)
788 approvals (or an equivalent approval/review based on the requirements of your country or
789 institution) were obtained?

790 Answer: [\[NA\]](#)

791 Justification: This study does not involve human participants and therefore does not require
792 IRB approval or risk disclosure.

793 Guidelines:

- 794 • The answer NA means that the paper does not involve crowdsourcing nor research with
795 human subjects.
- 796 • Depending on the country in which research is conducted, IRB approval (or equivalent)
797 may be required for any human subjects research. If you obtained IRB approval, you
798 should clearly state this in the paper.
- 799 • We recognize that the procedures for this may vary significantly between institutions
800 and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the
801 guidelines for their institution.

802 • For initial submissions, do not include any information that would break anonymity (if
803 applicable), such as the institution conducting the review.

804 **16. Declaration of LLM usage**

805 Question: Does the paper describe the usage of LLMs if it is an important, original, or
806 non-standard component of the core methods in this research? Note that if the LLM is used
807 only for writing, editing, or formatting purposes and does not impact the core methodology,
808 scientific rigorousness, or originality of the research, declaration is not required.

809 Answer: [NA]

810 Justification: This work does not use large language models (LLMs) as part of its core
811 methodology or experimental pipeline.

812 Guidelines:

813 • The answer NA means that the core method development in this research does not
814 involve LLMs as any important, original, or non-standard components.

815 • Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>)
816 for what should or should not be described.