

Degradation Accordant Plug-and-Play for Low-Rank Tensor Completion

Yexun Hu¹, Tai-Xiang Jiang^{1*} and Xi-Le Zhao²

¹School of Computing and Artificial Intelligence, Southwestern University of Finance and Economics

²School of Mathematical Sciences, University of Electronic Science and Technology of China
220081202001@smail.swufe.edu.cn, taixiangjiang@gmail.com, xlzhao122003@163.com

Abstract

Tensor completion aims at estimating missing values from an incomplete observation, playing a fundamental role in many applications. This work proposes a novel low-rank tensor completion model, in which the inherent low-rank prior and external degradation accordant data-driven prior are simultaneously utilized. Specifically, the tensor nuclear norm (TNN) is adopted to characterize the overall low-dimensionality of the tensor data. Meanwhile, an implicit regularizer is formulated and its related subproblem is solved via a deep convolutional neural network (CNN) under the plug-and-play framework. This CNN, pretrained for the *inpainting* task on a mass of natural images, is expected to express the external data-driven prior and this plugged inpainter is consistent with the original degradation process. Then, an efficient alternating direction method of multipliers (ADMM) is designed to solve the proposed optimization model. Extensive experiments are conducted on different types of tensor imaging data with the comparison with state-of-the-art methods, illustrating the effectiveness and the remarkable generalization ability of our method.

1 Introduction

The main aim of tensor completion is to fill the missing values of the data in the tensor format when the observation is incomplete owing to objective conditions. It has been widely studied for many real-world applications, such as, color image and video processing [Liu *et al.*, 2013], personalized web search [Sun *et al.*, 2005], high-order web link analysis [Kolda *et al.*, 2005], and fine-grained indoor localization [Liu *et al.*, 2015]. As many types of real-world high-dimensional tensor data, e.g., imaging data, seismic data, and traffic data, maintain low-dimensional structures, the tensor low-rankness is widely utilized for the tensor completion [Liu *et al.*, 2013], and low-rank tensor completion (LRTC) methods have achieved great successes in the past decade.

In this work, we mainly focus on the completion of multidimensional imaging data, including color images, videos, magnetic resonance imaging (MRI) data, multispectral images (MSIs). Meanwhile, although the rank of the tensor is not uniquely defined, we fix our attention on the tensor tubal rank defined based on the tensor singular value decomposition (t-SVD) [Kilmer and Martin, 2011] framework since that the t-SVD is specialized for third-order (namely, cubic) tensors, being suitable for multidimensional imaging data, and the data would be processed integrally within the t-SVD algebraic framework avoiding the loss of information inherent in matricization [Kilmer *et al.*, 2013]. Zhang *et al.* minimize the tensor tubal nuclear norm (TNN) [Zhang and Aeron, 2017], which is a convex surrogate of the tensor tubal rank, for LRTC and given the theoretical guarantee for exact recovery.

As we know, multidimensional images reserve abundant spatial details, making themselves not strictly low-rank. Thus, to better address this inverse problem, additional prior knowledge is needed to better preserve those details. For example, the total variation (TV) regularizer is introduced to characterize the piece-wise local smoothness in [Jiang *et al.*, 2018]. More recently, Zhao *et al.* [Zhao *et al.*, 2020] formulate an implicit regularizer and adopt a deep convolution neural network (CNN), which is pretrained on a large number of natural images for denoising, to express the deep denoising prior under the plug-and-play (PnP) framework [Venkatakrishnan *et al.*, 2013].

Within the flexible PnP framework, the degradation and the prior knowledge can be well decoupled after variable splitting via half-quadratic splitting [Nikolova and Ng, 2005] or the alternating direction method of multipliers (ADMM) [Boyd *et al.*, 2011], and one can directly employ off-the-shelf denoising algorithms to express the prior knowledge for various image inverse problems, such as deblurring and superresolution. When the well-known BM3D [Dabov *et al.*, 2007] is plugged in, it is believed to enhance the nonlocal self-similarity. With the rapid development of deep learning, we can also use deep CNN denoisers to utilize the data-driven prior learned from a mass of training images. Thanks to the high model capacity of the CNN, the method in [Zhao *et al.*, 2020] achieves promising performance for tensor completion.

We would like to take a further step to plug in a CNN inpainter for LRTC in this work. Two facts motivate us to conduct this. First, the plugged-in inpainter is designed for the

*Corresponding author.

inpainting task, which is naturally consistent with the degradation in tensor completion, being more suitable than denoisers. Second, as the model capacity of the CNN is quite high, it is reasonable to resort to the plugged-in CNN for undertaking parts of the degradation as well as expressing the image prior. Thus, our model is given by combining the internal low-rank prior with the external degradation accordant data-driven prior together. On the one hand, the TNN regularizer characterizes the global low-dimensional structure and inherent connection of the multi-dimensional data. On the other hand, an off-the-shelf CNN, which is pretrained for the image inpainting task on images, is employed to preserve the spatial details. Our contributions are summarised as follows.

- A novel degradation accordant PnP LRTC model is proposed for multidimensional images. In our model, the TNN is utilized to depict the inherent global low-rank structure of underlying multidimensional images. An implicit regularizer is formulated to better preserve the abundant details in multidimensional images via plugging in a CNN image inpainters. Those two terms are organically complementary to each other and our model is therefore expected to be effective.

- We customize the ADMM for the proposed model and the subproblem corresponding to the implicit regularizer is solved via the CNN inpainter, which is in accord with the original degradation process. This inpainter is pretrained on natural images, which are readily accessible, to faithfully express the data-driven image prior. Numerical experiments are conducted on various types of multidimensional images. Comparisons with state-of-the-art methods illustrate the excellent performance and generalization ability to different types of data of our method.

The outline of this paper is given as follows. Sec. 2 gives the basic preliminaries of the tensor. Sec. 3 presents main results. Sec. 4 illustrates experimental results. Finally, conclusions are drawn in Sec. 5.

2 Preliminaries

Throughout this paper, lowercase letters, e.g., x , boldface lowercase letters, e.g., $\mathbf{x}$, boldface uppercase letters, e.g., $\mathbf{X}$, and boldface calligraphic letters, e.g., $\mathcal{X}$, are used to denote scalars, vectors, matrices, and tensors, respectively. Given a third-order tensor $\mathcal{X} \in \mathbb{R}^{n_1 \times n_2 \times n_3}$, we use $\mathcal{X}_{ijk}$ to denote its (i, j, k) -th element. The k -th frontal slice of $\mathcal{X}$ is denoted as $\mathcal{X}^{(k)}$ (or $\mathcal{X}(:, :, k)$, $\mathbf{X}^k$). We use the notation $\hat{\mathcal{A}}$ to denote the Fourier transformed (along the third mode) tensor of $\mathcal{A}$. The conjugate transpose of a tensor $\mathcal{A} \in \mathbb{C}^{n_1 \times n_2 \times n_3}$ is tensor $\mathcal{A}^H \in \mathbb{C}^{n_1 \times n_2 \times n_3}$ obtained by conjugate transposing each of the frontal slice and then reversing the order of transposed frontal slices 2 through n_3 .

Definition 1 (T-prod [Kilmer and Martin, 2011]). *The tensor-tensor-product (t-prod) $\mathcal{C} = \mathcal{A} * \mathcal{B}$ of $\mathcal{A} \in \mathbb{R}^{n_1 \times n_2 \times n_3}$ and $\mathcal{B} \in \mathbb{R}^{n_2 \times n_4 \times n_3}$ is a tensor of the size $n_1 \times n_4 \times n_3$, where the (i, j) -th tube $\mathbf{c}_{ij}$ is given by $\mathbf{c}_{ij} = \sum_{k=1}^{n_2} \mathcal{A}(i, k, :) \otimes \mathcal{B}(k, j, :)$, where $\otimes$ denotes the circular convolution between two tubes of the same size.*

Definition 2 (Special tensors [Kilmer and Martin, 2011]). *The identity tensor $\mathcal{I} \in \mathbb{R}^{n_1 \times n_1 \times n_3}$ is the tensor whose first frontal slice is the $n_1 \times n_1$ identity matrix, and whose other*

*frontal slices are all zeros. A tensor $\mathcal{Q} \in \mathbb{C}^{n_1 \times n_1 \times n_3}$ is **orthogonal** if it satisfies $\mathcal{Q}^H * \mathcal{Q} = \mathcal{Q} * \mathcal{Q}^H = \mathcal{I}$. A tensor $\mathcal{A}$ is called **f-diagonal** if each frontal slice $\mathcal{A}^{(i)}$ is a diagonal matrix.*

Theorem 1 (T-SVD [Kilmer and Martin, 2011]). *For $\mathcal{A} \in \mathbb{R}^{n_1 \times n_2 \times n_3}$, the t-SVD of $\mathcal{A}$ is given by*

$$\mathcal{A} = \mathcal{U} * \mathcal{S} * \mathcal{V}^H \quad (1)$$

where $\mathcal{U} \in \mathbb{R}^{n_1 \times n_1 \times n_3}$ and $\mathcal{V} \in \mathbb{R}^{n_2 \times n_2 \times n_3}$ are orthogonal tensors, and $\mathcal{S} \in \mathbb{R}^{n_1 \times n_2 \times n_3}$ is an f-diagonal tensor.

Definition 3 (Tensor tubal-rank [Zhang and Aeron, 2017]). *The tubal-rank of a tensor $\mathcal{A} \in \mathbb{R}^{n_1 \times n_2 \times n_3}$, denoted as $\text{rank}_t(\mathcal{A})$, is defined to be the number of non-zero singular tubes of $\mathcal{S}$, where $\mathcal{S}$ comes from the t-SVD of $\mathcal{A}$: $\mathcal{A} = \mathcal{U} * \mathcal{S} * \mathcal{V}^T$.*

Definition 4 (Tubal-nuclear-norm (TNN) [Zhang and Aeron, 2017]). *The tensor nuclear norm of a tensor $\mathcal{A} \in \mathbb{R}^{n_1 \times n_2 \times n_3}$, denoted as $\|\mathcal{A}\|_{\text{TNN}}$, is defined as*

$$\|\mathcal{A}\|_{\text{TNN}} \triangleq \sum_{i=1}^{n_3} \|\hat{\mathcal{A}}(:, :, i)\|_*, \quad (2)$$

where $\|\cdot\|_$ is the matrix nuclear norm.*

3 Main Results

Our LRTC model is formulated as

$$\begin{aligned} \min_{\mathcal{X}} \quad & \lambda \|\mathcal{X}\|_{\text{TNN}} + \Phi(\mathcal{X}) \\ \text{s.t.} \quad & \mathcal{A}(\mathcal{X}) = \mathcal{O}. \end{aligned} \quad (3)$$

where $\mathcal{A}$ is a linear degradation operator, $\mathcal{O} \in \mathbb{R}^{n_1 \times n_2 \times n_3}$ is the observed multidimensional image¹ and $\Phi(\cdot)$ is an implicit regularizer. We introduce two auxiliary variables and reformulate (3) as

$$\begin{aligned} \min_{\mathcal{X}, \mathcal{Y}, \mathcal{Z}} \quad & \lambda \|\mathcal{X}\|_{\text{TNN}} + \Phi(\mathcal{X}) \\ \text{s.t.} \quad & \mathcal{A}(\mathcal{X}) = \mathcal{O}, \mathcal{A}(\mathcal{X}) = \mathcal{A}(\mathcal{Y}), \\ & \mathcal{Z} = \mathcal{Y}. \end{aligned} \quad (4)$$

Generally, one auxiliary variable is enough to decouple the low-rank regularizer and the implicit regularizer. We adopt two to formulate the implicit regularizer related subproblem in the image restoration format instead of denosing. This makes our method different from previous PnP methods.

Then, the augmented Lagrangian function is

$$\begin{aligned} L(\mathcal{X}, \mathcal{Y}, \mathcal{Z}, \Lambda_i) \\ = \lambda \|\mathcal{X}\|_{\text{TNN}} + \Phi(\mathcal{X}) + \frac{\beta_2}{2} \|\mathcal{A}(\mathcal{X}) - \mathcal{A}(\mathcal{Y})\|_F^2 + \frac{\Lambda_2}{\beta_2} \|\mathcal{A}(\mathcal{X}) - \mathcal{O}\|_F^2 \\ + \frac{\beta_1}{2} \|\mathcal{A}(\mathcal{X}) - \mathcal{O}\|_F^2 + \frac{\Lambda_1}{\beta_1} \|\mathcal{Y} - \mathcal{Z}\|_F^2 + \frac{\Lambda_3}{\beta_3} \|\mathcal{Y} - \mathcal{Z}\|_F^2, \end{aligned}$$

where Λ_i s ($i = 1, 2, 3$) are the Lagrangian multipliers, and β_i s ($i = 1, 2, 3$) are nonnegative parameters. Next, we set up ADMM iterations via solving following subproblems.

¹The spatial resolution is $n_1 \times n_2$ and the spectral (or temporal) length of the multidimensional image is n_3 .

Z-subproblem: At the k -th iteration, the subproblem with respect to $\mathcal{Z}$ is

$$\min_{\mathcal{Z}} \lambda \|\mathcal{Z}\|_{\text{TNN}} + \frac{\beta_3}{2} \|\mathcal{Y}^k - \mathcal{Z} + \frac{\Lambda_3^k}{\beta_3}\|_F^2. \quad (5)$$

Let the t-SVD of $(\mathcal{Y}^k + \frac{\Lambda_3^k}{\beta_3})$ be $\mathcal{U} * \mathcal{S} * \mathcal{V}^H$, the closed-form solution of (5) is

$$\mathcal{Z}^{k+1} = \mathcal{U} * \mathcal{D} * \mathcal{V}^H, \quad (6)$$

where $\mathcal{D}$ is an f-diagonal tensor obtained satisfying $\hat{D}(i, i, k) = \max\{\hat{S}(i, i, k) - \frac{\lambda}{n_3\beta_3}, 0\}$.

Y-subproblem: At the k -th iteration, the subproblem with respect to $\mathcal{Y}$ is

$$\begin{aligned} \min_{\mathcal{Y}} \frac{\beta_2}{2} \|\mathcal{A}(\mathcal{Y}) - \mathcal{A}(\mathcal{X}^k) - \frac{\Lambda_2^k}{\beta_2}\|_F^2 \\ + \frac{\beta_3}{2} \|\mathcal{Y} - \mathcal{Z}^{k+1} + \frac{\Lambda_3^k}{\beta_3}\|_F^2. \end{aligned} \quad (7)$$

Denoting the adjoint operator of $\mathcal{A}$ as $\mathcal{A}^*$, we can obtain the solution of (7) as

$$\begin{aligned} \mathcal{Y}^{k+1} = (\beta_2 \mathcal{A}^* \mathcal{A} + \beta_3 \mathcal{I})^{-1} \\ \cdot \left(\beta_2 \mathcal{A}^* \left(\mathcal{A}(\mathcal{X}^k) + \frac{\Lambda_2^k}{\beta_2} \right) + \beta_3 \mathcal{Z}^{k+1} - \Lambda_3^k \right), \end{aligned}$$

where $\mathcal{I}$ denotes the identity mapping.

X-subproblem: At the k -th iteration, the subproblem with respect to $\mathcal{X}$ is

$$\min_{\mathcal{X}} \Phi'(\mathcal{X}) + \frac{1}{2} \|\mathcal{A}(\mathcal{X}) - \mathcal{B}\|_F^2, \quad (8)$$

where $\mathcal{B} = \frac{(\beta_1 \mathcal{O} - \Lambda_1^k + \beta_2 \mathcal{A}(\mathcal{Y}^{k+1}) - \Lambda_2^k)}{\beta_1 + \beta_2}$ and $\Phi'(\mathcal{X}) = \frac{\Phi(\mathcal{X})}{\beta_1 + \beta_2}$. We further decouple (8) into n_3 subproblems as

$$\min_{\mathcal{X}(:, :, i)} \Phi'(\mathcal{X}(:, :, i)) + \frac{1}{2} \|\mathcal{A}(\mathcal{X}(:, :, i)) - \mathcal{B}(:, :, i)\|_F^2,$$

for $i = 1, 2, \dots, n_3$. Thus, it becomes n_3 image restoration problems and can be solved via off-the-shelf algorithms under the PnP framework. For the tensor completion, the degradation operator is indeed the projection operator $\mathcal{P}_{\Omega}(\cdot)$, which keeps the entries in Ω and set remaining entries as 0. We can solve it using inpainting algorithms as

$$\mathcal{X}^{k+1}(:, :, i) = \text{Inpainting}(\mathcal{B}(:, :, i)), \quad (9)$$

for $i = 1, 2, \dots, n_3$.

Finally, the update of multipliers follows the standard ADMM and can be seen in Supplementary Material².

Remark: Our method is not limited to the tensor completion. Other multidimensional image restoration tasks, such as video superresolution or deblurring, can also be addressed via replacing the degradation operator $\mathcal{A}(\cdot)$ and adopting corresponding image restoration algorithms in (9). We would like to utilize CNNs in (9). As analyzed in [Ryu *et al.*, 2019], the convergence of our algorithm would be guaranteed if the CNN is properly trained. Meanwhile, we want to also emphasize the generalization ability of our method as our method can be applied for different types of multidimensional images and only needs the inpainting CNN trained on grayscale or color images.

²Can be found at <https://github.com/TaiXiangJiang/>.

Mask Type	Type-1		Type-2		Type-3		Time
Method	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	(s)
Observed	18.31	0.948	12.22	0.612	17.38	0.874	—
HaLRTC	30.48	0.967	12.22	0.612	30.85	0.956	21.2
TNN	22.88	0.960	18.26	0.732	25.60	0.928	8.6
DCTNN	30.81	0.969	26.29	0.894	29.40	0.949	3.9
FTNN	19.71	0.953	26.84	0.920	27.87	0.944	34.6
DP3LRTC	31.18	0.971	26.46	0.911	35.12	0.980	9.6
OPN	29.54	0.966	26.20	0.911	34.84	0.981	0.8
Deepfillv2	29.93	0.969	24.77	0.897	34.53	0.981	34.5
DAP-LRTC	31.59	0.972	27.12	0.929	36.57	0.987	131.7

Table 1: Quantitative results by different methods on **color images** with different types of structural missing. The **best** and the **second best** values are respectively highlighted by **red** and **blue** colors.

4 Experiments

In this section, we evaluate the effectiveness of our degradation accordant PnP LRTC (DAP-LRTC) method and compare it with other state-of-the-art methods on color images, videos, MSIs, and MRI data. Compared methods consist of: the Tucker-rank based method HaLRTC³ [Liu *et al.*, 2013], a t-SVD based method (TNN)⁴ [Zhang and Aeron, 2017], a DCT induced TNN minimization method (DCTNN)⁵ [Lu *et al.*, 2019], a framelet represented TNN minimization method (FTNN)⁶ [Jiang *et al.*, 2020], a deep denoiser regularized TNN minimization method (DP3LRTC)⁷ [Zhao *et al.*, 2020], and a deep video inpainter called Onion-peel networks (OPN)⁸ [Oh *et al.*, 2019].

For all experiments, two numerical metrics are employed, including the Peak signal-to-noise ratio (PSNR), the structural similarity index (SSIM) [Wang *et al.*, 2004]. Higher PSNR and SSIM values mean better performance. Additionally, we introduce the mean spectral angle mapper (SAM) for MSIs, and lower SAM indicates better results. All experiments were conducted on the platform of Window 10 with an AMD Ryzen9 3950X CPU and RTX 2080Ti GPU and 32RAM. 2 CNN inpainters are considered in our method: i) Deepfillv2⁹ [Yu *et al.*, 2019] for structural missing, ii) CRUNet (we borrow the network structure of DRUNet¹⁰ in [Zhang *et al.*, 2021] and train it for grey-scale image completion) for random missing. The training details of CRUNet are given in Supplementary Materials. The results by applying Deepfillv2 and CRUNet on each band of the incomplete data would also be posted for reference.

4.1 Color Image

In this subsection, 6 color images¹¹ of the size $512 \times 512 \times 3$ are selected. We consider the structural missing in all

³<https://www.cs.rochester.edu/~jliu/code/TensorCompletion.zip>

⁴https://github.com/jamiezeminzhang/Tensor_Completion_and_Tensor_RPCA

⁵Implemented by ourselves based on the code of TNN

⁶<https://github.com/TaiXiangJiang/Framelet-TNN>

⁷<https://taixiangjiang.github.io/>

⁸<https://github.com/seoungwugoh/opn-demo>

⁹https://github.com/JiahuiYu/generative_inpainting

¹⁰<https://github.com/csnn/DPIR>

¹¹Available at <http://sipi.usc.edu/database/database.php>.

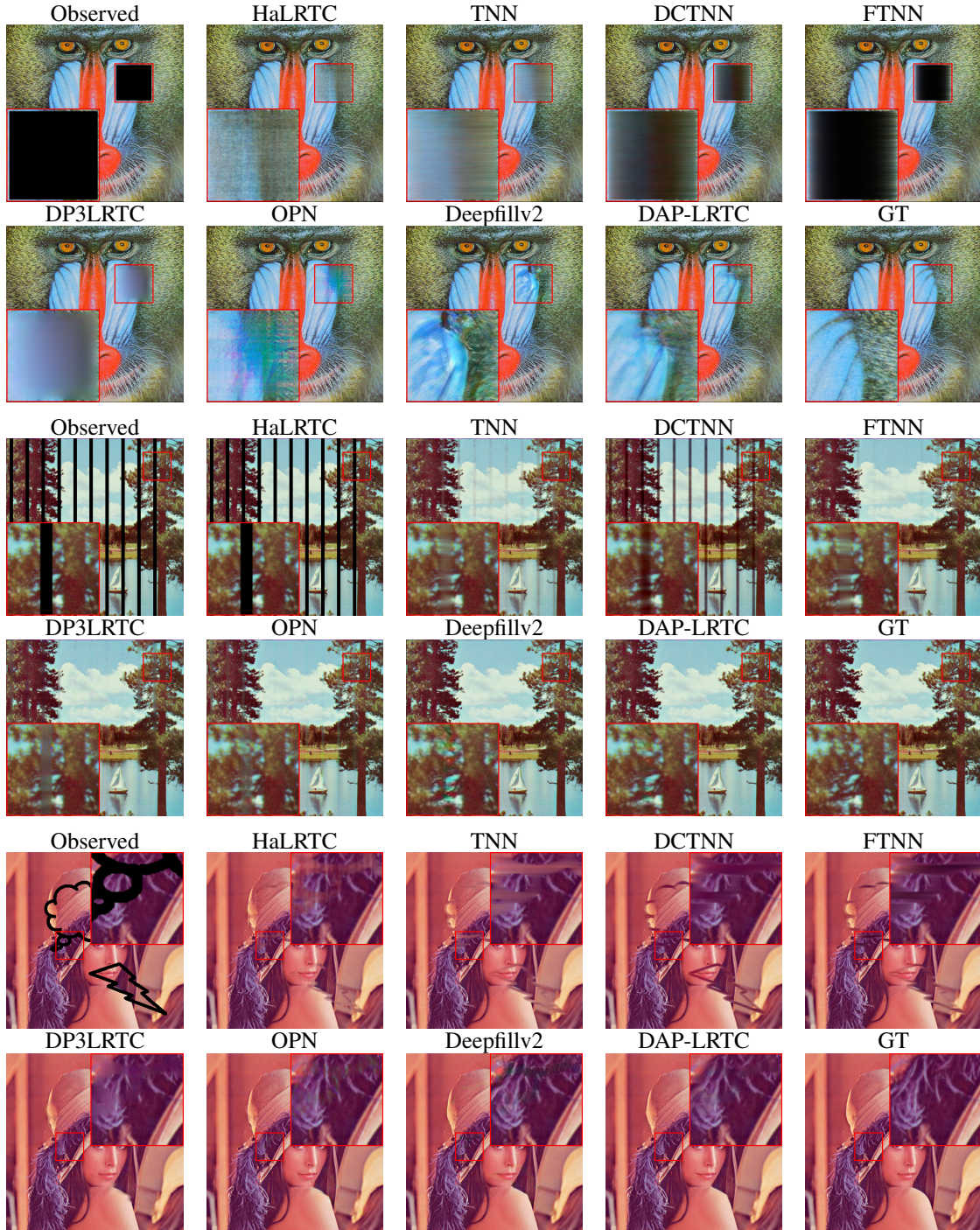


Figure 1: From top to bottom: The inpainting results on the different color images with different structural missing types.

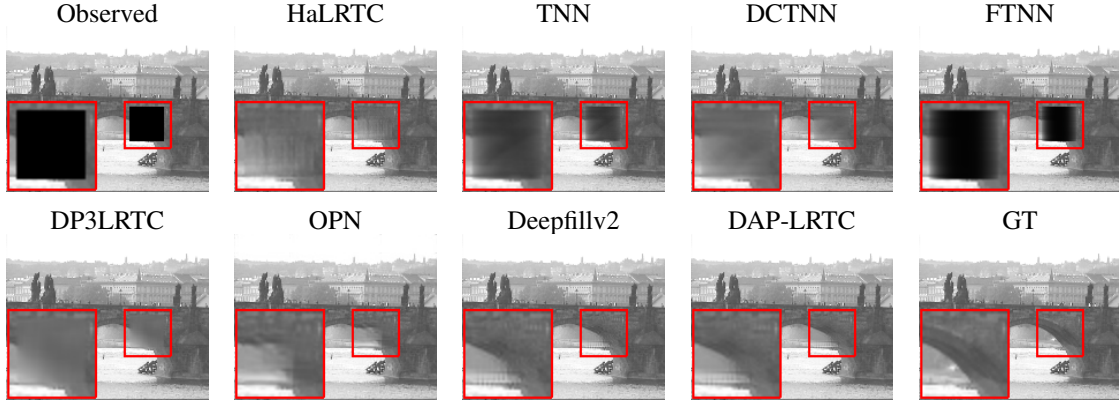
RGB channels. We adopt deepfillv2 to solve (9). Tab.1 reports quantitative metrics and the running time in seconds. We can see that our method outperforms compared methods. Fig.1 shows the inpainting results by different methods on color images (*Baboon*, *Painting*, *Lena*) with three different masks, respectively. From the enlarged areas, the superior of our method is obvious. More results on color images are

shown in Supplementary Material.

4.2 Video

In this subsection, we choose four videos¹² to test the capacity of the proposed method. The size of all videos is $144 \times 176 \times 50$. For videos, we consider both the random missing and

¹²Videos available at <http://trace.eas.asu.edu/yuv/>.


 Figure 2: The 27-th frame of recovered results of different methods on video *Bridge-close* with a block of the size 30 by 30 missing.

Mask Type	Bridge-far		Bridge-close		Time (s)
Method	PSNR	SSIM	PSNR	SSIM	
Observed	19.95	0.938	20.76	0.945	—
HaLRTC	47.87	0.993	35.20	0.975	4.3
TNN	31.45	0.975	28.06	0.968	19.0
DCTNN	47.65	0.994	34.92	0.979	17.6
FTNN	24.11	0.954	22.68	0.955	57.9
DP3LRTC	48.34	0.994	36.70	0.982	47.2
OPN	45.11	0.990	34.24	0.971	8.7
Deepfillv2	47.29	0.987	37.54	0.981	3.4
DAP-LRTC	51.89	0.995	37.93	0.983	83.6

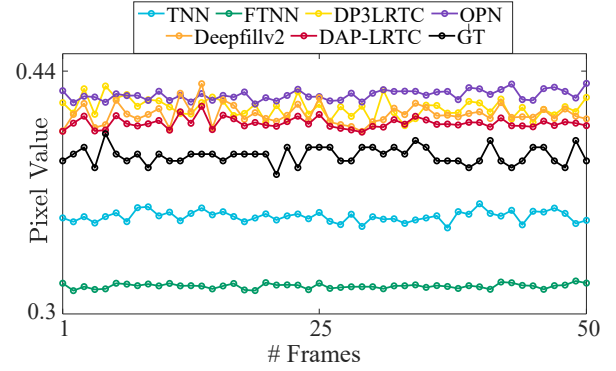
 Table 2: Quantitative results by different methods on **videos** with a random block loss of the size 30 by 30. The **best** and the **second best** values are respectively highlighted by **red** and **blue** colors.

structural missing. Deepfillv2 and CRUNet are respectively employed for these two cases. Tab.2 exhibits the quantitative metrics and the running time on the videos *Bridge-close* and *Bridge-far* for structural missing. It can be observed that our method gets the highest metrics. The remainder of the experiments on video data is also given in the supplementary material. The visual results on the video *Bridge-close* are shown in Fig.2. We can see that our method and Deepfillv2 preserve the structure of the bridge while other methods failed. We further plot one temporal vector of the missing area in Fig.3. From Fig.3, it can be seen that the temporal vector of our result is the closest to the original data. This shows the effectiveness brought by the TNN regularizer and validates that the TNN and the CNN inpainter are complementary to each other. Please see Supplementary Material for more results on videos.

4.3 MSI

In this part, we select 4 MSIs¹³ of the size $512 \times 512 \times 31$. The observation is obtained via random sampling with the sampling rates (SRs) 0.03 and 0.05, respectively. We use CRUNet in our method for this random missing case. In Tab.3, we list quantitative metrics of results by different methods on MSI *Thread spool* with different SRs. We can see that the performance of FTNN, DP3LRTC, and CRUNet are

¹³<http://www.cs.columbia.edu/CAVE/databases/multispectral/>.


 Figure 3: The temporal vectors in the missing area of recovered results on the video *Bridge-close* by different methods.

good as they occasionally achieve the second best. Although the inpainter is not trained for MSIs, our method obtains the best values for different metrics and sampling rates, showing remarkable generalization ability. The pseudo-color images (composed of three bands of the MSI) of the reconstructed results by different methods are displayed in Fig.4. We can see that OPN, which is trained for videos could not handle the MSI. Our method, DP3LRTC, and CRUNet obtain good results while we can see over-smoothness in the result by DP3LRTC and the color distortion in the result by CRUNet. More experimental results are in Supplementary Material.

4.4 Parameter Analysis and Convergence

There are four parameters, i.e., λ , β_1 , β_2 , and β_3 , in our method. To test the effects from different values of them, we conduct experiments on the MRI¹⁴ data with the random sampling rate of 10%. When testing one parameter, the other three are fixed as default values. We illustrate the PSNR and SSIM values with respect to different values of those parameters in Fig. 5. From Fig. 5, we can see that the performance of our method is more sensitive to λ and β_3 .

¹⁴Available at <https://brainweb.bic.mni.mcgill.ca/brainweb>. Please refer to Supplementary Material for more results on MRI data.

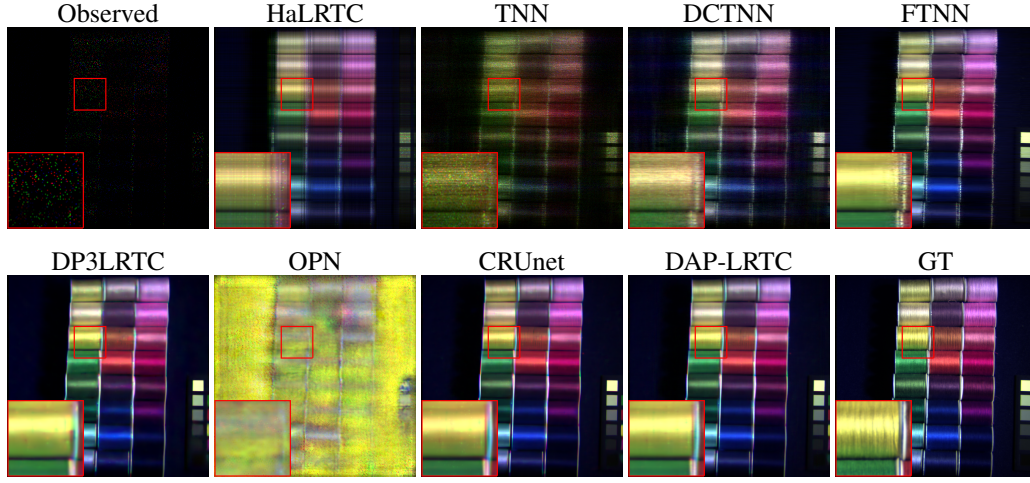


Figure 4: Pseudo color images (composed of the 25-th, 15-th, and the 1st bands) of recovered results by different methods on the MSI *Thread spools* with SR= 0.03.

SR	3%			5%			Time (s)
Method	PSNR	SSIM	SAM	PSNR	SSIM	SAM	
Observed	16.55	0.273	—	16.64	0.291	—	—
HaLRTC	25.30	0.835	11.543	27.98	0.895	8.798	135.1
TNN	21.79	0.680	34.656	25.60	0.798	25.316	64.4
DCTNN	26.67	0.770	19.996	31.28	0.895	12.966	41.5
FTNN	31.20	0.951	6.264	34.38	0.977	4.707	582.7
DP3LRTC	32.64	0.912	8.408	35.88	0.968	5.479	697.8
OPN	11.64	0.205	45.503	12.32	0.223	44.527	3.7
CRUnet	28.91	0.925	8.060	35.43	0.975	4.497	4.0
DAP-LRTC	33.52	0.963	5.506	36.17	0.978	4.410	208.8

Table 3: Quantitative results by different methods on the MSI *Thread spools* with different sample rates. The **best** and the **second best** values are respectively highlighted by **red** and **blue** colors.

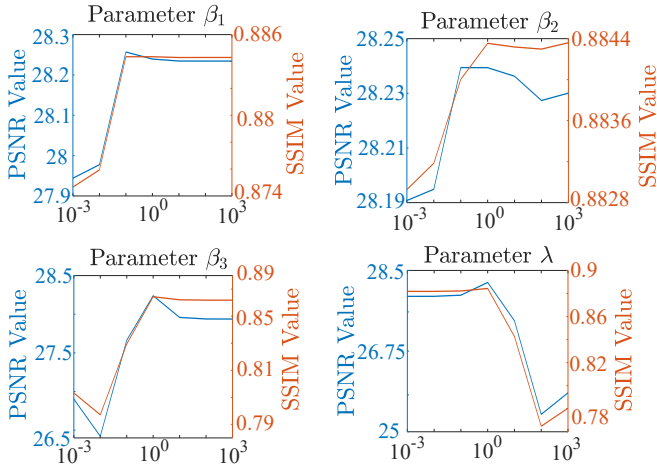


Figure 5: The PSNR and SSIM values of results by our method with different β_1 , β_2 , β_3 , and λ on the MRI data (SR = 10%).

Also on the MRI data with random sampling rate 10%, we report the relative change of each variable with respect to iteration numbers in Fig. 6. Meanwhile, the result on the color

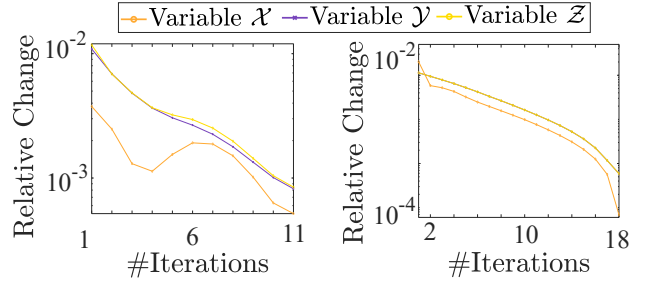


Figure 6: The relative changes of the variables. Left: MRI data with SR=10%. Right: color image “baboon” with structural missing.

image “baboon” with structural missing (first row of Fig. 1) is also plotted. We plug in the CRUnet for the MRI data with random missing and Deepfillv2 for the color image with structural missing. We can see the empirical convergence of our method for different settings.

5 Conclusion

We propose a novel degradation accordant PnP LRTC approach for multidimensional images. In our model, the TNN is utilized to depict the global low-rank structure of underlying multidimensional images. An implicit regularizer is formulated and its corresponding subproblem is equivalent to a series of single image restoration problems. The CNN pre-trained for image inpainting is plugged in as the solution. It is consistent with the original degradation process and faithfully expresses the image prior learned from a large number of training images. An ADMM algorithm is tailored for our model. Numerical experiments are conducted on various types of multidimensional images to compare our method with state-of-the-art methods, illustrating the excellent performance of our method. Although our method obtains the best performance and shows remarkable generalization ability for different types of data, it is still time-consuming. This would be considered in future work.

Acknowledgements

We would like thank the authors of [Liu *et al.*, 2013; Zhang and Aeron, 2017; Lu *et al.*, 2019; Zhang *et al.*, 2021] for their generous sharing of their codes and data. This research is supported by National Natural Science Foundation of China (Nos. 12001446, 61876203, 61772003), the Key Project of Applied Basic Research in Sichuan Province (No. 2020YJ0216), the Applied Basic Research Project of Sichuan Province (No. 2021YJ0107), National Key Research and Development Program of China (No. 2020YFA0714001), the Fundamental Research Funds for the Central Universities under Grant (Nos. JBK2202049, JBK2102001).

References

- [Boyd *et al.*, 2011] Stephen Boyd, Neal Parikh, Eric Chu, Borja Peleato, and Jonathan Eckstein. Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends® in Machine Learning*, 3(1):1–122, 2011.
- [Dabov *et al.*, 2007] Kostadin Dabov, Alessandro Foi, Vladimir Katkovnik, and Karen Egiazarian. Image denoising by sparse 3-D transform-domain collaborative filtering. *IEEE Transactions on image processing*, 16(8):2080–2095, 2007.
- [Jiang *et al.*, 2018] Fei Jiang, Xiao-Yang Liu, Hongtao Lu, and Ruimin Shen. Anisotropic total variation regularized low-rank tensor completion based on tensor nuclear norm for color image inpainting. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing*, pages 1363–1367, 2018.
- [Jiang *et al.*, 2020] Tai-Xiang Jiang, Michael K Ng, Xi-Le Zhao, and Ting-Zhu Huang. Framelet representation of tensor nuclear norm for third-order tensor completion. *IEEE Transactions on Image Processing*, 29:7233–7244, 2020.
- [Kilmer and Martin, 2011] Misha E Kilmer and Carla D Martin. Factorization strategies for third-order tensors. *Linear Algebra and its Applications*, 435(3):641–658, 2011.
- [Kilmer *et al.*, 2013] Misha E Kilmer, Karen Braman, Ning Hao, and Randy C Hoover. Third-order tensors as operators on matrices: A theoretical and computational framework with applications in imaging. *SIAM Journal on Matrix Analysis and Applications*, 34(1):148–172, 2013.
- [Kolda *et al.*, 2005] Tamara G Kolda, Brett W Bader, and Joseph P Kenny. Higher-order web link analysis using multilinear algebra. In *Proceedings of the IEEE International Conference on Data Mining*, pages 242–249, 2005.
- [Liu *et al.*, 2013] Ji Liu, Przemyslaw Musialski, Peter Wonka, and Jieping Ye. Tensor completion for estimating missing values in visual data. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 35(1):208–220, 2013.
- [Liu *et al.*, 2015] Xiao-Yang Liu, Shuchin Aeron, Vaneet Aggarwal, Xiaodong Wang, and Min-You Wu. Adaptive sampling of RF fingerprints for fine-grained indoor localization. *IEEE Transactions on Mobile Computing*, 15(10):2411–2423, 2015.
- [Lu *et al.*, 2019] Canyi Lu, Xi Peng, and Yunchao Wei. Low-rank tensor completion with a new tensor nuclear norm induced by invertible linear transforms. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 5996–6004, 2019.
- [Nikolova and Ng, 2005] Mila Nikolova and Michael K Ng. Analysis of half-quadratic minimization methods for signal and image recovery. *SIAM Journal on Scientific computing*, 27(3):937–966, 2005.
- [Oh *et al.*, 2019] Seoung Wug Oh, Sungho Lee, Joon-Young Lee, and Seon Joo Kim. Onion-peel networks for deep video completion. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 4403–4412, 2019.
- [Ryu *et al.*, 2019] Ernest Ryu, Jialin Liu, Sicheng Wang, Xiaohan Chen, Zhangyang Wang, and Wotao Yin. Plug-and-play methods provably converge with properly trained denoisers. In *Proceedings of the International Conference on Machine Learning*, pages 5546–5557, 2019.
- [Sun *et al.*, 2005] Jian-Tao Sun, Hua-Jun Zeng, Huan Liu, Yuchang Lu, and Zheng Chen. CubeSVD: a novel approach to personalized web search. In *Proceedings of the International Conference on World Wide Web*, pages 382–390, 2005.
- [Venkatakrishnan *et al.*, 2013] Singanallur V Venkatakrishnan, Charles A Bouman, and Brendt Wohlberg. Plug-and-play priors for model based reconstruction. In *Proceedings of the IEEE Global Conference on Signal and Information Processing*, pages 945–948, 2013.
- [Wang *et al.*, 2004] Zhou Wang, Alan C Bovik, Hamid R Sheikh, and Eero P Simoncelli. Image quality assessment: from error visibility to structural similarity. *IEEE Transactions on Image Processing*, 13(4):600–612, 2004.
- [Yu *et al.*, 2019] Jiahui Yu, Zhe Lin, Jimei Yang, Xiaohui Shen, Xin Lu, and Thomas S Huang. Free-form image inpainting with gated convolution. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 4471–4480, 2019.
- [Zhang and Aeron, 2017] Zemin Zhang and Shuchin Aeron. Exact tensor completion using t-SVD. *IEEE Transactions on Signal Processing*, 65(6):1511–1526, 2017.
- [Zhang *et al.*, 2021] Kai Zhang, Yawei Li, Wangmeng Zuo, Lei Zhang, Luc Van Gool, and Radu Timofte. Plug-and-play image restoration with deep denoiser prior. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2021.
- [Zhao *et al.*, 2020] Xi-Le Zhao, Wen-Hao Xu, Tai-Xiang Jiang, Yao Wang, and Michael K Ng. Deep plug-and-play prior for low-rank tensor completion. *Neurocomputing*, 400:137–149, 2020.