
Tensor denoising and completion based on ordinal observations

Chanwoo Lee¹ Miaoyan Wang¹

Abstract

Higher-order tensors arise frequently in applications such as neuroimaging, recommendation system, and social network analysis. We consider the problem of low-rank tensor estimation from possibly incomplete, ordinal-valued observations. Two related problems are studied, one on tensor denoising and another on tensor completion. We propose a multi-linear cumulative link model, develop a rank-constrained M-estimator, and obtain theoretical accuracy guarantees. Our mean squared error bound enjoys a faster convergence rate than previous results, and we show that the proposed estimator is minimax optimal under the class of low-rank models. Furthermore, the procedure developed serves as an efficient completion method which guarantees consistent recovery of an order- K $(d, \dots, d)$ -dimensional low-rank tensor using only $\tilde{O}(Kd)$ noisy, quantized observations. We demonstrate the outperformance of our approach over previous methods on the tasks of clustering and collaborative filtering.

1. Introduction

Multidimensional arrays, a.k.a. tensors, arise in a variety of applications including recommendation systems (Baltrunas et al., 2011), social networks (Nickel et al., 2011), genomics (Wang et al., 2019), and neuroimaging (Zhou et al., 2013). There is a growing need to develop general methods that can address two main problems for analyzing these noisy, high-dimensional datasets. The first problem is tensor denoising which aims to recover a signal tensor from its noisy entries (Hong et al., 2020; Wang & Zeng, 2019). The second problem is tensor completion which examines the minimum number of entries needed for a consistent recovery (Ghadermarzy et al., 2018; Montanari & Sun, 2018). Low-rankness is often imposed to the signal

tensor, thereby efficiently reducing the intrinsic dimension in both problems.

A number of low-rank tensor estimation methods have been proposed (Hong et al., 2020; Wang & Song, 2017), revitalizing classical methods such as CANDECOMP/PARAFAC (CP) decomposition (Hitchcock, 1927) and Tucker decomposition (Tucker, 1966). These tensor methods treat the entries as continuous-valued. In many cases, however, we encounter datasets of which the entries are qualitative. For example, the Netflix problem records the ratings of users on movies over time. Each data entry is a rating on a nominal scale $\{\text{very like}, \text{like}, \text{neutral}, \text{dislike}, \text{very dislike}\}$. Another example is in the signal processing, where the digits are frequently rounded or truncated so that only integer values are available. The qualitative observations take values in a limited set of categories, making the learning problem harder compared to continuous observations.

Ordinal entries are categorical variables with an ordering among the categories; for example, $\text{very like} \prec \text{like} \prec \text{neutral} \prec \dots$. The analyses of tensors with the ordinal entries are mainly complicated by two key properties needed for a reasonable model. First, the model should be invariant under a reversal of categories, say, from the Netflix example, $\text{very like} \succ \text{like} \succ \text{neutral} \succ \dots$, but not under arbitrary label permutations. Second, the parameter interpretations should be consistent under merging or splitting of contiguous categories. The classical continuous tensor model (Kolda & Bader, 2009; Ghadermarzy et al., 2019) fails in the first aspect, whereas the binary tensor model (Ghadermarzy et al., 2018) lacks the second property. An appropriate model for ordinal tensors has yet to be studied.

Our contributions. We are among the first to establish the recovery theory for signal tensors and quantization operators simultaneously from a limited number of highly discrete entries. Our main contributions are summarized in Table 1. We propose a cumulative link model for higher-order tensors, develop a rank-constrained M-estimator, and obtain theoretical accuracy guarantees. The mean squared error bound is established, and we show that the obtained bound has minimax optimal rate in high dimensions under the low-rank model. Furthermore, our proposal guarantees consistent recovery of an order- K $(d, \dots, d)$ -dimensional low-rank tensor using only $\tilde{O}(Kd)$ noisy, quantized observations.

¹Department of Statistics, University of Wisconsin-Madison, Wisconsin, USA. Correspondence to: Miaoyan Wang <miaoyan.wang@wisc.edu>.

	Bhaskar [2016]	Ghadermarzy et al. [2018]	This paper
Higher-order tensors ($K \geq 3$)	$\times$	$\checkmark$	$\checkmark$
Multi-level categories ($L \geq 3$)	$\checkmark$	$\times$	$\checkmark$
Error rate for tensor denoising	d^{-1} for $K = 2$	$d^{-(K-1)/2}$	$d^{-(K-1)}$
Optimality guarantee under low-rank models	unknown	$\times$	$\checkmark$
Sample complexity for tensor completion	d^K	Kd	Kd

Table 1. Comparison with previous work. For ease of presentation, we summarize the error rate and sample complexity assuming equal tensor dimension in all modes. K : tensor order; L : number of ordinal levels; d : dimension at each mode.

Related work. Our work is connected to non-Gaussian tensor decomposition. Existing work focuses exclusively on univariate observations such as binary- or continuous-valued entries (Wang & Li; Ghadermarzy et al., 2018). The problem of ordinal/quantized tensor is fundamentally more challenging than previously well-studied tensors for two reasons: (a) the entries do not belong to exponential family distribution, and (b) the observation contains much less information, because neither the underlying signal nor the quantization operator is known. The limited information makes the statistical recovery notably hard.

We address the challenge by proposing a cumulative link model that enjoys palindromic invariance (McCullagh, 1980). A distinctive non-monotonic, phase-transition pattern is demonstrated, as we show in Section 6. We prove that the recovery from quantized tensors achieves equally good information-theoretical convergence as the continuous tensors. These results fills the gap between classical and non-classical (ordinal) observations, thereby greatly enriching the tensor model literature.

From algorithm perspective, we address the challenge using the (non-convex) alternating algorithm. Earlier work has proposed an approximate (convex) algorithm for binary tensor completion (Ghadermarzy et al., 2018). Unlike matrix problems, convex-relaxation for low-rank tensors suffers from both computational intractability (Hillar & Lim, 2013) and statistical suboptimality. We improve the error bound from $\mathcal{O}(d^{-(K-1)/2})$ in Ghadermarzy et al. (2018) to $\mathcal{O}(d^{-(K-1)})$ and numerically compare the two approaches.

We also highlight the challenge associated with higher-order tensors. Matrix completion has been proposed for binary observations (Cai & Zhou, 2013; Davenport et al., 2014; Bhaskar & Javanmard, 2015) and for ordinal observations (Bhaskar, 2016). We show that, applying existing matrix methods to higher-order tensors results in suboptimal estimates. A full exploitation of the higher-order structure is needed; this is another challenge we address in this paper.

2. Preliminaries

Let $\mathcal{Y} \in \mathbb{R}^{d_1 \times \dots \times d_K}$ denote an order- K ($d_1, \dots, d_K$)-dimensional tensor. We use y_ω to denote the tensor entry

indexed by ω , where $\omega \in [d_1] \times \dots \times [d_K]$. The Frobenius norm of $\mathcal{Y}$ is defined as $\|\mathcal{Y}\|_F = \sum_\omega y_\omega^2$ and the infinity norm is defined as $\|\mathcal{Y}\|_\infty = \max_\omega |y_\omega|$. We use $\mathcal{Y}_{(k)}$ to denote the unfolded matrix of size d_k -by- $\prod_{i \neq k} d_i$, obtained by reshaping the tensor along the mode $k \in [K]$. The Tucker rank of $\mathcal{Y}$ is defined as a length- K vector $\mathbf{r} = (r_1, \dots, r_K)$, where r_k is the rank of matrix $\mathcal{Y}_{(k)}$ for $k \in [K]$.

We use lower-case letters ($a, b, \dots$) for scalars/vectors, upper-case boldface letters ($\mathbf{A}, \mathbf{B}, \dots$) for matrices, and calligraphy letters ($\mathcal{A}, \mathcal{B}, \dots$) for tensors of order three or greater. An event A is said to occur “with very high probability” if $\mathbb{P}(A)$ tends to 1 faster than any polynomial of tensor dimension $d_{\min} = \min\{d_1, \dots, d_K\} \rightarrow \infty$. The indicator function of an event A is denoted as $\mathbb{1}\{A\}$. For ease of notation, we allow basic arithmetic operators (e.g., $\leq, +, -$) to be applied to pairs of tensors in an element-wise manner. We use the shorthand $[n]$ to denote $\{1, \dots, n\}$ for $n \in \mathbb{N}_+$.

3. Model formulation and motivation

3.1. Observation model

Let $\mathcal{Y}$ denote an order- K ($d_1, \dots, d_K$)-dimensional data tensor. Suppose the entries of $\mathcal{Y}$ are ordinal-valued, and the observation space consists of L ordered levels, denoted by $[L] = \{1, \dots, L\}$. We propose a cumulative link model for the ordinal tensor $\mathcal{Y} = \llbracket y_\omega \rrbracket \in [L]^{d_1 \times \dots \times d_K}$. Specifically, assume the entries y_ω are (conditionally) independently distributed with cumulative probabilities,

$$\mathbb{P}(y_\omega \leq \ell) = f(b_\ell - \theta_\omega), \text{ for all } \ell \in [L-1], \quad (1)$$

where $\mathbf{b} = (b_1, \dots, b_{L-1})$ is a set of unknown scalars satisfying $b_1 < \dots < b_{L-1}$, $\Theta = \llbracket \theta_\omega \rrbracket \in \mathbb{R}^{d_1 \times \dots \times d_K}$ is a continuous-valued parameter tensor satisfying certain low-dimensional structure (to be specified later), and $f(\cdot): \mathbb{R} \mapsto [0, 1]$ is a known, strictly increasing function. We refer to $\mathbf{b}$ as the cut-off points and f the link function.

The formulation (1) imposes an additive model to the transformed probability of cumulative categories. This modeling choice is to respect the ordering structure among the categories. For example, if we choose the inverse link

$f^{-1}(x) = \log \frac{x}{1-x}$ to be the log odds, then the model (1) implies a linear spacing between the proportional odds,

$$\log \frac{\mathbb{P}(y_\omega \leq \ell)}{\mathbb{P}(y_\omega > \ell)} - \log \frac{\mathbb{P}(y_\omega \leq \ell-1)}{\mathbb{P}(y_\omega > \ell-1)} = b_\ell - b_{\ell-1}, \quad (2)$$

for all tensor entries y_ω . When there are only two categories in the observation space (e.g., for binary tensors), the cumulative model (1) is equivalent to the usual binomial link model. In general, however, when the number of categories $L \geq 3$, the proportional odds assumption (2) is more parsimonious, in that, the ordered categories can be envisaged as contiguous intervals on the continuous scale, where the points of division are exactly $b_1 < \dots < b_{L-1}$. This interpretation will be made explicit in the next section.

3.2. Latent-variable interpretation

The ordinal tensor model (1) with certain types of link f has the equivalent representation as an L -level quantization model on $\mathcal{Y} = \llbracket y_\omega \rrbracket$:

$$y_\omega = \begin{cases} 1, & \text{if } y_\omega^* \in (-\infty, b_1], \\ 2, & \text{if } y_\omega^* \in (b_1, b_2], \\ \vdots & \vdots \\ L, & \text{if } y_\omega^* \in (b_{L-1}, \infty), \end{cases} \quad (3)$$

for all $\omega \in [d_1] \times \dots \times [d_K]$. Here, $\mathcal{Y}^* = \llbracket y_\omega^* \rrbracket$ is a latent continuous-valued tensor following an additive noise model,

$$\underbrace{\mathcal{Y}^*}_{\text{latent continuous-valued tensor}} = \underbrace{\Theta}_{\text{signal tensor}} + \underbrace{\mathcal{E}}_{\text{i.i.d. noise}}, \quad (4)$$

where $\mathcal{E} = \llbracket \varepsilon_\omega \rrbracket \in \mathbb{R}^{d_1 \times \dots \times d_K}$ is a noise tensor with independent and identically distributed (i.i.d.) entries according to distribution $\mathbb{P}(\varepsilon)$. From the viewpoint of (4), the parameter tensor Θ can be interpreted as the latent signal tensor prior to contamination and quantization.

The equivalence between the latent-variable model (3) and the cumulative link model (1) is established if the link f is chosen to be the cumulative distribution function of noise ε , i.e., $f(\theta) = \mathbb{P}(\varepsilon \leq \theta)$. We describe two common choices of link f , or equivalently, the distribution of ε .

Example 1 (Logistic model). The logistic model is characterized by (1) with $f(\theta) = (1 + e^{-\theta/\sigma})^{-1}$, where $\sigma > 0$ is the scale parameter. Equivalently, the noise ε_ω in (3) follows i.i.d. logistic distribution with scale parameter σ .

Example 2 (Probit model). The probit model is characterized by (1) with $f(\theta) = \mathbb{P}(z \leq \theta/\sigma)$, where $z \sim N(0, 1)$. Equivalently, the noise ε_ω in (3) follows i.i.d. $N(0, \sigma^2)$.

Other link functions are also possible, such as Laplace, Cauchy, etc (McCullagh, 1980). These latent variable models share the property that the ordered categories can be

thought of as contiguous intervals on some continuous scale. We should point out that, although the latent-variable interpretation is incisive, our estimation procedure does not refer to the existence of $\mathcal{Y}^*$. Therefore, our model (1) is general and still valid in the absence of quantization process. More generally, we make the following assumptions about the link f .

Assumption 1. The link function is assumed to satisfy:

1. $f(\theta)$ is strictly increasing and twice-differentiable in $\theta \in \mathbb{R}$.
2. $f'(\theta)$ is strictly log-concave and symmetric with respect to $\theta = 0$.

3.3. Problem 1: Tensor denoising

The first question we aim to address is tensor denoising:

(P1) Given the quantization process induced by f and the cut-off points $\mathbf{b}$, how accurately can we estimate the latent signal tensor Θ from the ordinal observation $\mathcal{Y}$?

Clearly, the problem (P1) cannot be solved uniformly for all possible Θ with no assumptions. We focus on a class of “low-rank” and “flat” signal tensors, which is a plausible assumption in practical applications (Zhou et al., 2013; Bhaskar & Javanmard, 2015). Specifically, we consider the parameter space,

$$\mathcal{P} = \{\Theta \in \mathbb{R}^{d_1 \times \dots \times d_K} : \text{rank}(\Theta) \leq \mathbf{r}, \|\Theta\|_\infty \leq \alpha\}, \quad (5)$$

where $\mathbf{r} = (r_1, \dots, r_K)$ denotes the Tucker rank of Θ .

The parameter tensor of our interest satisfies two constraints. The first is that Θ is a low-rank tensor, with $r_k = \mathcal{O}(1)$ as $d_{\min} \rightarrow \infty$ for all $k \in [K]$. Equivalently, Θ admits the Tucker decomposition:

$$\Theta = \mathcal{C} \times_1 \mathbf{M}_1 \times_1 \dots \times_K \mathbf{M}_K, \quad (6)$$

where $\mathcal{C} \in \mathbb{R}^{r_1 \times \dots \times r_K}$ is a core tensor, $\mathbf{M}_k \in \mathbb{R}^{d_k \times r_k}$ are factor matrices with orthogonal columns, and $\times_k$ denotes the tensor-by-matrix multiplication (Kolda & Bader, 2009). The Tucker low-rankness is popularly imposed in tensor analysis, and the rank determines the tradeoff between model complexity and model flexibility. Note that, unlike matrices, there are various notions of tensor low-rankness, such as CP rank (Hitchcock, 1927) and train rank (Oseledets, 2011). Some notions of low-rankness may lead to mathematically ill-posed optimization; for example, the best low CP-rank tensor approximation may not exist (De Silva & Lim, 2008). We choose Tucker representation for well-posedness of optimization and easy interpretation.

The second constraint is that the entries of Θ are uniformly bounded in magnitude by a constant $\alpha \in \mathbb{R}_+$. In view of (4),

we refer to α as the signal level. The boundedness assumption is a technical condition that avoids the degeneracy in probability estimation with ordinal observations.

3.4. Problem 2: Tensor completion

Motivated by applications in collaborative filtering, we also consider a more general setup when only a subset of tensor entries y_ω are observed. Let $\Omega \subset [d_1] \times \cdots \times [d_K]$ denote the set of observed indices. The second question we aim to address is stated as follows:

(P2) Given an incomplete set of ordinal observations $\{y_\omega\}_{\omega \in \Omega}$, how many sampled entries do we need to consistently recover Θ based on the model (1)?

The answer to (P2) depends on the choice of Ω . We consider a general model on Ω that allows both uniform and non-uniform sampling. Specifically, let $\Pi = \{\pi_{i_1, \dots, i_K}\}$ denote a predefine probability distribution over the index set such that $\sum_{\omega \in [d_1] \times \cdots \times [d_K]} \pi_\omega = 1$. We assume that each index in Ω is drawn with replacement using distribution Π . This sampling model relaxes the uniform sampling in literature and is arguably a better fit in applications.

We consider the same parameter space (5) for the completion problem. In addition to the reasons mentioned in Section 3.3, the entrywise bound assumption also serves as the incoherence requirement for completion. In classical matrix completion, the incoherence is often imposed on the singular vectors. This assumption is recently relaxed for “flat” matrices with bounded magnitude (Negahban et al., 2011; Cai & Zhou, 2013; Bhaskar & Javanmard, 2015). We adopt the same assumption for higher-order tensors.

4. Rank-constrained M-estimator

We present a general treatment to both problems mentioned above. With a little abuse of notation, we use Ω to denote either the full index set $\Omega = [d_1] \times \cdots \times [d_K]$ (for the tensor denoising) or a random subset induced from the sampling distribution Π (for the tensor completion). Define $b_0 = -\infty$, $b_L = \infty$, $f(-\infty) = 0$ and $f(\infty) = 1$. The log-likelihood associated with the observed entries is

$$\mathcal{L}_{\mathcal{Y}, \Omega}(\Theta, \mathbf{b}) = \sum_{\omega \in \Omega} \sum_{\ell \in [L]} \left\{ \mathbb{1}\{y_\omega = \ell\} \log [f(b_\ell - \theta_\omega) - f(b_{\ell-1} - \theta_\omega)] \right\}. \quad (7)$$

We propose a rank-constrained maximum likelihood estimator (a.k.a. M-estimator) for Θ ,

$$\hat{\Theta} = \arg \max_{\Theta \in \mathcal{P}} \mathcal{L}_{\mathcal{Y}, \Omega}(\Theta, \mathbf{b}), \text{ where}$$

$$\mathcal{P} = \{\Theta \in \mathbb{R}^{d_1 \times \cdots \times d_K} : \text{rank}(\Theta) \leq r, \|\Theta\|_\infty \leq \alpha\} \quad (8)$$

In practice, the cut-off points $\mathbf{b}$ are unknown and should be jointly estimated with Θ . We will present the theory and

algorithm with known $\mathbf{b}$ in the main paper. The adaptation for unknown $\mathbf{b}$ is addressed in Section 5 and the Supplement.

We define a few key quantities that will be used in our theory. Let $g_\ell = f(\theta + b_\ell) - f(\theta + b_{\ell-1})$ for all $\ell \in [L]$, and

$$A_\alpha = \min_{\ell \in [L], |\theta| \leq \alpha} g_\ell(\theta), \quad U_\alpha = \max_{\ell \in [L], |\theta| \leq \alpha} \frac{|\dot{g}_\ell(\theta)|}{g_\ell(\theta)},$$

$$L_\alpha = \min_{\ell \in [L], |\theta| \leq \alpha} \left[\frac{\dot{g}_\ell^2(\theta)}{g_\ell^2(\theta)} - \frac{\ddot{g}_\ell(\theta)}{g_\ell(\theta)} \right],$$

where $\dot{g}(\theta) = dg(\theta)/d\theta$, and α is the entrywise bound of Θ . In view of equation (4), these quantities characterize the geometry including flatness and convexity of the latent noise distribution. Under Assumption 1, all these quantities are strictly positive and independent of tensor dimension.

4.1. Estimation error for tensor denoising

For the tensor denoising problem, we assume that the full set of tensor entries are observed. We assess the estimation accuracy using the mean squared error (MSE):

$$\text{MSE}(\hat{\Theta}, \Theta^{\text{true}}) = \frac{1}{\prod_k d_k} \|\Theta - \Theta^{\text{true}}\|_F^2.$$

The next theorem establishes the upper bound for the MSE of the proposed $\hat{\Theta}$ in (8).

Theorem 4.1 (Statistical convergence). *Consider an ordinal tensor $\mathcal{Y} \in [L]^{d_1 \times \cdots \times d_K}$ generated from model (1), with the link function f and the true coefficient tensor $\Theta^{\text{true}} \in \mathcal{P}$. Define $r_{\max} = \max_k r_k$. Then, with very high probability, the estimator in (8) satisfies*

$$\text{MSE}(\hat{\Theta}, \Theta^{\text{true}}) \leq \min \left(4\alpha^2, \frac{c_1 U_\alpha^2 r_{\max}^{K-1} \sum_k d_k}{L_\alpha^2 \prod_k d_k} \right), \quad (9)$$

where $c_1 > 0$ is a constant that depends only on K .

Theorem 4.1 establishes the statistical convergence for the estimator (8). In fact, the proof of this theorem (see the Supplement) shows that the same statistical rate holds, not only for the global optimizer (8), but also for any local optimizer $\check{\Theta}$ in the level set $\{\check{\Theta} \in \mathcal{P} : \mathcal{L}_{\mathcal{Y}, \Omega}(\check{\Theta}) \geq \mathcal{L}_{\mathcal{Y}, \Omega}(\Theta^{\text{true}})\}$. This suggests that the local optimality itself is not necessarily a severe concern in our context, as long as the convergent objective is large enough. In Section 5, we perform empirical studies to assess the algorithmic stability.

To gain insight into the bound (9), we consider a special setting with equal dimension in all modes, i.e., $d_1 = \cdots = d_K = d$. In such a case, our bound (9) reduces to

$$\text{MSE}(\hat{\Theta}, \Theta^{\text{true}}) \asymp d^{-(K-1)}, \quad \text{as } d \rightarrow \infty.$$

Hence, our estimator achieves consistency with polynomial convergence rate. We compare the bound with existing

literature. In the special case $L = 2$, Ghadermarzy et al. (2018) proposed a max-norm constrained estimator $\hat{\Theta}$ with $\text{MSE}(\hat{\Theta}, \Theta^{\text{true}}) \asymp d^{-(K-1)/2}$. In contrast, our estimator converges at a rate of $d^{-(K-1)}$, which is substantially faster than theirs. This provides a positive answer to the open question posed in Ghadermarzy et al. (2018) whether the square root in the bound is removable. The improvement stems from the fact that we have used the exact low-rankness of Θ , whereas the surrogate rank measure employed in Ghadermarzy et al. (2018) is scale-sensitive.

Our bound also generalizes the previous results on ordinal matrices. The convergence rate for rank-constrained matrix estimation is $\mathcal{O}(1/\sqrt{d})$ (Bhaskar, 2016), which fits into our special case when $K = 2$. Furthermore, our result (9) reveals that the convergence becomes favorable as the order of data tensor increases. Intuitively, the sample size for analyzing a data tensor is the number of entries, $\prod_k d_k$, and the number of free parameters is roughly on the order of $\sum_k d_k$, assuming $r_{\max} = \mathcal{O}(1)$. A higher tensor order implies higher effective sample size per parameter, thus achieving a faster convergence rate in high dimensions.

A similar conclusion is obtained for the prediction error, measured in Kullback-Leibler (KL) divergence, between the categorical distributions in the observation space.

Corollary 1 (Prediction error). Assume the same set-up as in Theorem 4.1. Let $\mathbb{P}_{\mathcal{Y}}$ and $\hat{\mathbb{P}}_{\mathcal{Y}}$ denote the distributions generating the L -level ordinal tensor $\mathcal{Y}$, given the true parameter Θ and its estimator $\hat{\Theta}$, respectively. Assume $L \geq 2$. Then, with very high probability,

$$\text{KL}(\mathbb{P}_{\mathcal{Y}} \parallel \hat{\mathbb{P}}_{\mathcal{Y}}) \leq \frac{c_1 U_{\alpha}^2 r_{\max}^{K-1}}{L_{\alpha}^2} \frac{(4L-6) \dot{f}^2(0)}{A_{\alpha}} \frac{\sum_k d_k}{\prod_k d_k},$$

where $c_1 > 0$ is the same constant as in Theorem 4.1.

We next show the statistical optimality of our estimator $\hat{\Theta}$. The result is based on the information theory and applies to all estimators in $\mathcal{P}$, including but not limited to $\hat{\Theta}$ in (8).

Theorem 4.2 (Minimax lower bound). Assume the same set-up as in Theorem 4.1, and $d_{\max} = \max_k d_k \geq 8$. Let $\inf_{\hat{\Theta}}$ denote the infimum over all estimators $\hat{\Theta} \in \mathcal{P}$ based on the ordinal tensor observation $\mathcal{Y} \in [L]^{d_1 \times \dots \times d_K}$. Then, under the model (1),

$$\inf_{\hat{\Theta}} \sup_{\Theta^{\text{true}} \in \mathcal{P}} \mathbb{P} \left\{ \text{MSE}(\hat{\Theta}, \Theta^{\text{true}}) \geq c \min \left(\alpha^2, \frac{C r_{\max} d_{\max}}{\prod_k d_k} \right) \right\} \geq \frac{1}{8},$$

where $C = C(\alpha, L, f, \mathbf{b}) > 0$ and $c > 0$ are constants independent of tensor dimension and the rank.

We see that the lower bound matches the upper bound in (9) on the polynomial order of tensor dimension. Therefore, our estimator (8) is rate-optimal.

4.2. Sample complexity for tensor completion

We now consider the tensor completion problem, when only a subset of entries Ω are observed. We consider a general sampling procedure induced by Π . The recovery accuracy is assessed by the weighted squared error,

$$\begin{aligned} \|\Theta - \hat{\Theta}\|_{F, \Pi}^2 &\stackrel{\text{def}}{=} \frac{1}{|\Omega|} \mathbb{E}_{\Omega \sim \Pi} \|\Theta - \hat{\Theta}\|_F^2 \\ &= \sum_{\omega \in [d_1] \times \dots \times [d_K]} \pi_{\omega} (\Theta_{\omega} - \hat{\Theta}_{\omega})^2. \end{aligned} \quad (10)$$

Note that the recovery error depends on the distribution Π . In particular, tensor entries with higher sampling probabilities have more influence on the recovery accuracy, compared to the ones with lower sampling probabilities.

Remark 1. If we assume each entry is sampled with some strictly positive probability; i.e., there exists a constant $\mu > 0$ such that

$$\pi_{\omega} \geq \frac{1}{\mu \prod_k d_k}, \quad \text{for all } \omega \in [d_1] \times \dots \times [d_K],$$

then the error in (10) provides an upper bound for MSE:

$$\|\Theta - \hat{\Theta}\|_{F, \Pi}^2 \geq \frac{\|\Theta - \hat{\Theta}\|_F^2}{\mu \prod_k d_k} = \frac{1}{\mu} \text{MSE}(\hat{\Theta}, \Theta^{\text{true}}).$$

The equality is attained under uniform sampling with $\mu = 1$.

Theorem 4.3. Assume the same set-up as in Theorem 4.1. Suppose that we observe a subset of tensor entries $\{y_{\omega}\}_{\omega \in \Omega}$, where Ω is chosen at random with replacement according to a probability distribution Π . Let $\hat{\Theta}$ be the solution to (8), and assume $r_{\max} = \mathcal{O}(1)$. Then, with very high probability,

$$\|\Theta - \hat{\Theta}\|_{F, \Pi}^2 \rightarrow 0, \quad \text{as } \frac{|\Omega|}{\sum_k d_k} \rightarrow \infty.$$

Theorem 4.3 shows that our estimator achieves consistent recovery using as few as $\tilde{\mathcal{O}}(Kd)$ noisy, quantized observations from an order- K ($d, \dots, d$)-dimensional tensor. Note that $\tilde{\mathcal{O}}(Kd)$ roughly matches the degree of freedom for an order- K tensor of fixed rank r , suggesting the optimality of our sample requirement. This sample complexity substantially improves over earlier result $\mathcal{O}(d^{\lceil K/2 \rceil})$ based on square matrixization (Mu et al., 2014), or $\mathcal{O}(d^{K/2})$ based on tensor nuclear-norm regularization (Yuan & Zhang, 2016). Existing methods that achieve $\tilde{\mathcal{O}}(Kd)$ sample complexity require either a deterministic cross sampling design (Zhang, 2019) or univariate measurements (Ghadermarzy et al., 2018). Our method extends the conclusions to multi-level measurements under a broader class of sampling schemes.

5. Numerical Implementation

We describe the algorithm to seek the optimizer of (7). In practice, the cut-off points $\mathbf{b}$ are often unknown, so we

Algorithm 1 Ordinal tensor decomposition

Input: Ordinal data tensor $\mathcal{Y} \in [L]^{d_1 \times \dots \times d_K}$, rank $\mathbf{r} \in \mathbb{N}_+^K$, entry-wise bound $\alpha \in \mathbb{R}_+$.

Output: $(\hat{\Theta}, \hat{\mathbf{b}}) = \arg \max_{(\Theta, \mathbf{b}) \in \mathcal{P} \times \mathcal{B}} \mathcal{L}_{\mathcal{Y}, \Omega}(\Theta, \mathbf{b})$.

Random initialization of core tensor $\mathcal{C}^{(0)}$, factor matrices $\{\mathbf{M}_k^{(0)}\}$, and cut-off points $\mathbf{b}^{(0)}$.

for $t = 1, 2, \dots$, **do**

for $k = 1, 2, \dots, K$ **do**

 Update $\mathbf{M}_k^{(t+1)}$ while fixing other blocks:

$\mathbf{M}_k^{(t+1)} \leftarrow \arg \max_{\mathbf{M}_k \in \mathbb{R}^{d_k \times r_K}} \mathcal{L}_{\mathcal{Y}, \Omega}(\mathbf{M}_k)$,

 s.t. $\|\Theta^{(t+1)}\|_\infty \leq \alpha$, where $\Theta^{(t+1)}$ is the parameter tensor based on the current block estimates.

end for

 Update $\mathcal{C}^{(t+1)}$ while fixing other blocks:

$\mathcal{C}^{(t+1)} \leftarrow \arg \max_{\mathcal{C} \in \mathbb{R}^{r_1 \times \dots \times r_K}} \mathcal{L}_{\mathcal{Y}, \Omega}(\mathcal{C})$, s.t. $\|\Theta^{(t+1)}\|_\infty \leq \alpha$.

 Update $\Theta^{(t+1)}$ based on the current block estimates:

$\Theta^{(t+1)} \leftarrow \mathcal{C}^{(t+1)} \times_1 \mathbf{M}_1^{(t+1)} \dots \times_K \mathbf{M}_K^{(t+1)}$.

 Update $\mathbf{b}^{(t+1)}$ while fixing $\Theta^{(t+1)}$:

$\mathbf{b}^{(t+1)} \leftarrow \arg \max_{\mathbf{b} \in \mathcal{B}} \mathcal{L}_{\mathcal{Y}, \Omega}(\Theta^{(t+1)}, \mathbf{b})$.

end for

return $(\hat{\Theta}, \hat{\mathbf{b}})$

choose to maximize $\mathcal{L}_{\mathcal{Y}, \Omega}$ jointly over $(\Theta, \mathbf{b}) \in \mathcal{P} \times \mathcal{B}$ (see the Supplement for details). The objective $\mathcal{L}_{\mathcal{Y}, \Omega}$ is concave in $(\Theta, \mathbf{b})$ whenever f' is log-concave. However, the feasible set $\mathcal{P}$ is non-convex, which makes the optimization (7) a non-convex problem. We employ the alternating optimization approach by utilizing the Tucker representation of Θ . Specifically, based on (6) and (7), the objective function consists of $K + 2$ blocks of variables, one for the cut-off points $\mathbf{b}$, one for the core tensor $\mathcal{C}$, and K for the factor matrices $\mathbf{M}_k$'s. The optimization is a simple convex problem if any $K + 1$ out of the $K + 2$ blocks are fixed. We update one block at a time while holding others fixed, and we alternate the optimization throughout the iteration. The convergence is guaranteed whenever $\mathcal{L}_{\mathcal{Y}, \Omega}$ is bounded from above, since the alternating procedure monotonically increases the objective. The Algorithm 1 gives the full description.

We comment on two implementation details before concluding this section. First, the problem (8) is non-convex, so Algorithm 1 usually has no theoretical guarantee on global optimality. Nevertheless, as shown in Section 4.1, the desired rate holds not only for the global optimizer, but also for the local optimizer with $\mathcal{L}_{\mathcal{Y}, \Omega}(\hat{\Theta}) \geq \mathcal{L}_{\mathcal{Y}, \Omega}(\Theta^{\text{true}})$. In practice, we find the convergence point $\hat{\Theta}$ upon random initialization is often satisfactory, in that the corresponding objective $\mathcal{L}_{\mathcal{Y}, \Omega}(\hat{\Theta})$ is close to and actually slightly larger than the objective evaluated at the true parameter $\mathcal{L}_{\mathcal{Y}, \Omega}(\Theta^{\text{true}})$. Figure 1 shows the trajectory of the objective function that is output in the default setting of Algorithm 1,

with the input tensor generated from probit model (1) with $d_1 = d_2 = d_3 = d$ and $r_1 = r_2 = r_3 = r$. The dashed line is the objective value at the true parameter $\mathcal{L}_{\mathcal{Y}, \Omega}(\Theta^{\text{true}})$. We find that the algorithm generally converges quickly to a desirable value in reasonable number of steps. The actual running time per iteration is shown in the plot legend.

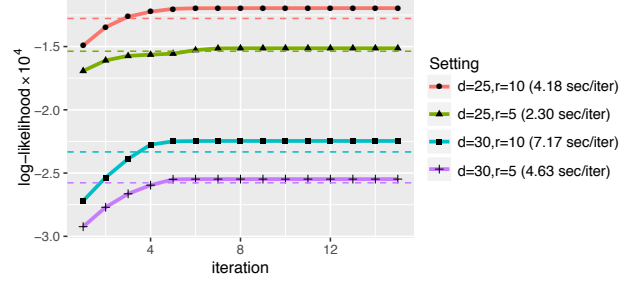


Figure 1. Trajectory of objective function with various d and r .

Second, the algorithm takes the rank $\mathbf{r}$ as an input. In practice, the rank $\mathbf{r}$ is hardly known and needs to be estimated from the data. We use Bayesian information criterion (BIC) and choose the rank that minimizes BIC; i.e.,

$$\begin{aligned} \hat{\mathbf{r}} &= \arg \min_{\mathbf{r} \in \mathbb{N}_+^K} \text{BIC}(\mathbf{r}) \\ &= \arg \min_{\mathbf{r} \in \mathbb{N}_+^K} \{-2\mathcal{L}_{\mathcal{Y}}(\hat{\Theta}(\mathbf{r}), \hat{\mathbf{b}}(\mathbf{r})) + p_e(\mathbf{r}) \log(\prod_k d_k)\}, \end{aligned}$$

where $\hat{\Theta}(\mathbf{r}), \hat{\mathbf{b}}(\mathbf{r})$ are the estimates given the rank $\mathbf{r}$, and $p_e(\mathbf{r}) \stackrel{\text{def}}{=} \sum_k (d_k - r_k)r_k + \prod_k r_k$ is the effective number of parameters in the model. We select $\hat{\mathbf{r}}$ that minimizes BIC through a grid search. The choice of BIC is intended to balance between the goodness-of-fit for the data and the degrees of freedom in the population model.

6. Experiments

In this section, we evaluate the empirical performance of our methods¹. We investigate both the complete and the incomplete settings, and we compare the recovery accuracy with other tensor-based methods. Unless otherwise stated, the ordinal data tensors are generated from model (1) using standard probit link f . We consider the setting with $K = 3$, $d_1 = d_2 = d_3 = d$, and $r_1 = r_2 = r_3 = r$. The parameter tensors are simulated based on (6), where the core tensor entries are i.i.d. drawn from $N(0, 1)$, and the factors $\mathbf{M}_k$ are uniformly sampled (with respect to Haar measure) from matrices with orthonormal columns. We set the cut-off points $b_\ell = f^{-1}(\ell/L)$ for $\ell \in [L]$, such that $f(b_\ell)$ are evenly spaced from 0 to 1. In each simulation study, we report the summary statistics across $n_{\text{sim}} = 30$ replications.

¹Software package: <https://CRAN.R-project.org/package=tensorordinal>

6.1. Finite-sample performance

The first experiment examines the performance under complete observations. We assess the empirical relationship between the MSE and various aspects of model complexity, such as dimension d , rank r , and signal level $\alpha = \|\Theta\|_\infty$. Figure 2a plots the estimation error versus the tensor dimension d for three different ranks $r \in \{3, 5, 8\}$. The decay in the error appears to behave on the order of d^{-2} , which is consistent with our theoretical results (9). We find that a higher rank leads to a larger error, as reflected by the upward shift of the curve as r increases. Indeed, a higher rank implies the higher number of parameters to estimate, thus increasing the difficulty of the estimation. Figure 2b shows the estimation error versus the signal level under $d = 20$. Interestingly, a larger estimation error is observed when the signal is either too small or too large. The non-monotonic behavior may seem surprising, but this is an intrinsic feature in the estimation with ordinal data. In view of the latent-variable interpretation (see Section 3.2), estimation from ordinal observation can be interpreted as an inverse problem of quantization. Therefore, the estimation error diverges in the absence of noise $\mathcal{E}$, because it is impossible to distinguish two different signal tensors, e.g., $\Theta_1 = \mathbf{a}_1 \otimes \mathbf{a}_2 \otimes \mathbf{a}_3$ and $\Theta_2 = \text{sign}(\mathbf{a}_1) \otimes \text{sign}(\mathbf{a}_2) \otimes \text{sign}(\mathbf{a}_3)$, from the quantized observations. This phenomenon (Davenport et al., 2014; Sur & Candès, 2019) is clearly contrary to the classical continuous-valued tensor problem.

The second experiment investigates the incomplete observations. We consider L -level tensors with $d = 20$, $\alpha = 10$ and choose a subset of tensor entries via uniform sampling. Figure 2c shows the estimation error of $\hat{\Theta}$ versus the fraction of observation $\rho = |\Omega|/d^K$. As expected, the error reduces with increased ρ or decreased r . Figure 2d evaluates the impact of ordinal levels L to estimation accuracy, under the setting $\rho = 0.5$. An improved performance is observed as L grows, especially from binary observations ($L = 2$) to multi-level ordinal observations ($L \geq 3$). The result showcases the benefit of multi-level observations compared to binary observations.

6.2. Comparison with alternative methods

Next, we compare our ordinal tensor method (**Ordinal-T**) with three popular low-rank methods.

- Continuous tensor decomposition (**Continuous-T**) (Acar et al., 2010) is a low-rank approximation method based on classical Tucker model.
- One-bit tensor completion (**1bit-T**) (Ghadermarzy et al., 2018) is a max-norm penalized tensor learning method based on partial binary observations.
- Ordinal matrix completion (**Ordinal-M**) (Bhaskar, 2016) is a rank-constrained matrix estimation method based on noisy, quantized observations.

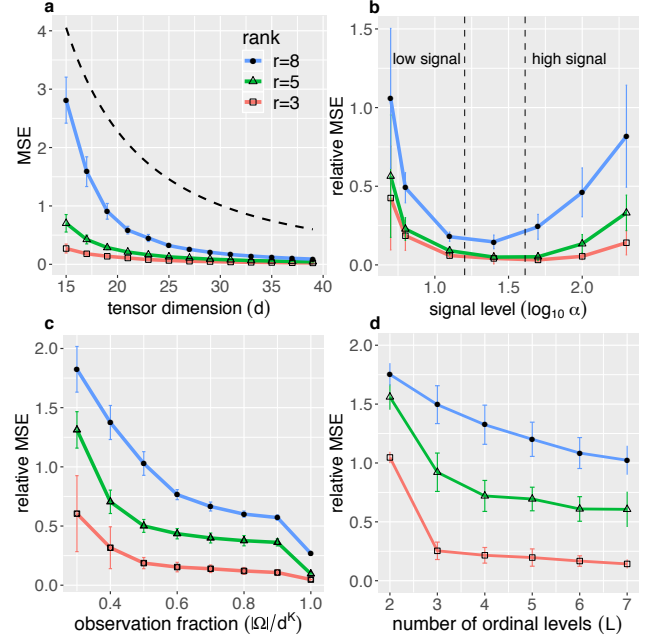


Figure 2. Empirical relationship between (relative) MSE versus (a) dimension d , (b) signal level α , (c) observation fraction ρ , and (d) number of ordinal levels L . In panels (b)-(d), we plot the relative MSE $= \|\hat{\Theta} - \Theta^{\text{true}}\|_F^2 / \|\Theta^{\text{true}}\|_F^2$ for better visualization.

We apply each of the above methods to L -level ordinal tensors $\mathcal{Y}$ generated from model (1). The **Continuous-T** is applied directly to $\mathcal{Y}$ by treating the L levels as continuous observations. The **Ordinal-M** is applied to the 1-mode unfolding matrix $\mathcal{Y}_{(1)}$. The **1bit-T** is applied to $\mathcal{Y}$ in two ways. The first approach, denoted **1bit-sign-T**, follows from Ghadermarzy et al. (2018) that transforms $\mathcal{Y}$ to a binary tensor, by taking the entrywise sign of the mean-adjusted tensor, $\mathcal{Y} - |\Omega|^{-1} \sum_{\omega} y_{\omega}$. The second approach, denoted **1bit-category-T**, transforms the order-3 ordinal tensor $\mathcal{Y}$ to an order-4 binary tensor, $\mathcal{Y}^{\#} = \llbracket y_{ijkl}^{\#} \rrbracket$, via dummy variable encoding; i.e., $y_{ijkl}^{\#} = \mathbb{1}\{y_{ijk} = \ell\}$ for $\ell \in [L-1]$. We evaluate the methods by their capabilities in predicting the most likely labels, $y_{\omega}^{\text{mode}} = \arg \max_{\ell} \mathbb{P}(y_{\omega} = \ell)$. Two performance metrics are considered: mean absolute deviation (MAD) $= d^{-K} \sum_{\omega} |y_{\omega}^{\text{mode}} - \hat{y}_{\omega}^{\text{mode}}|$, and misclassification rate (MCR) $= d^{-K} \sum_{\omega} \mathbb{1}\{y_{\omega}^{\text{mode}} \neq \text{round}(\hat{y}_{\omega}^{\text{mode}})\}$, where $\text{round}(\cdot)$ denotes the nearest integer. Note that MAD penalizes the large deviation more heavily than MCR.

Figure 3 compares the prediction accuracy under the setting $\alpha = 10$, $d = 20$, and $r = 5$. We find that our method outperforms the others in both MAD and MCR. In particular, methods built on multi-level observations (**Ordinal-T**, **Ordinal-M**, **1bit-category-T**) exhibit stable MCR over ρ and L , whereas the others two methods (**Continuous-T**, **1bit-sign-T**) generally fail except for $L = 2$ (Figures 3a-b). This observation highlights the benefits of multi-level modeling in the classification task. Although **1bit-category-**

T and our method **Ordinal-T** behave similarly for binary tensors ($L = 2$), the improvement of our method is substantial as L increases (Figures 3a and 3c). One possible reason is that our method incorporates the intrinsic ordering among the L levels via proportional odds assumption (2), whereas **1bit-category-T** ignores the ordinal structure and dependence among the induced binary entries.

Figures 3c-d assess the prediction accuracy with sample size. We see a clear advantage of our method (**Ordinal-T**) over the matricization (**Ordinal-M**) in both complete and non-complete observations. When the observation fraction is small, e.g., $|\Omega|/d^K = 0.4$, the tensor-based completion shows $\sim 30\%$ reduction in error compared to the matricization.

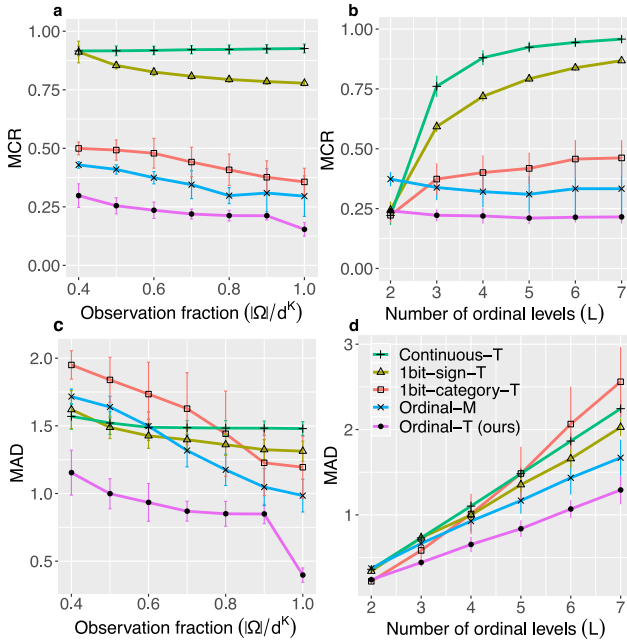


Figure 3. Performance comparison for predicting most likely labels. (a, c) Prediction errors versus sample complexity $\rho = |\Omega|/d^K$ when $L = 5$. (b, d) Prediction errors versus the number of ordinal levels L when $\rho = 0.8$.

We also compare the methods by their performance in predicting the median labels, $y_\omega^{\text{median}} = \min\{\ell: \mathbb{P}(y_\omega = \ell) \geq 0.5\}$. Under the latent variable model (4) and Assumption 1, the median label is the quantized θ_ω without noise; i.e., $y_\omega^{\text{median}} = \sum_\ell \mathbb{1}\{\theta_\omega \in (b_{\ell-1}, b_\ell]\}$. We utilize the same simulation setting as in the earlier experiment. Figure 4 shows that our method outperforms the others in both MCR and MAD. The improved accuracy comes from the incorporation of multilinear low-rank structure, multi-level observations, and the ordinal structure. Interestingly, the median estimator tends to yield smaller MAD than the mode estimator, $\text{MAD}(\mathcal{Y}^{\text{median}}, \hat{\mathcal{Y}}^{\text{median}}) \leq \text{MAD}(\mathcal{Y}^{\text{mode}}, \hat{\mathcal{Y}}^{\text{mode}})$ (Figures 3a-b vs. Figures 4a-b), for the three multilevel methods (**1bit-sign-T**, **Ordinal-M**, and **Ordinal-T**). The

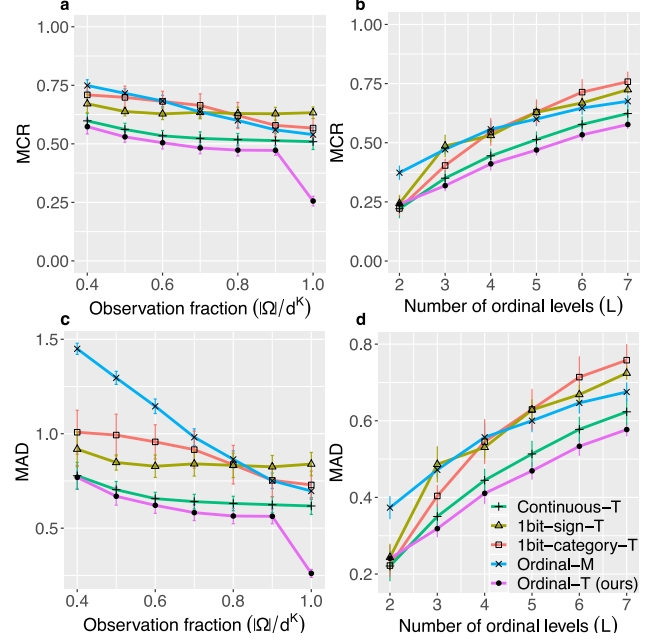


Figure 4. Performance comparison for predicting median labels. (a, c) Prediction error versus sample complexity $\rho = |\Omega|/d^K$ when $L = 5$. (b, d) Prediction error versus the number of ordinal levels L , when $\rho = 0.8$.

mode estimator, on the other hand, tends to yield smaller MCR than the median estimator, $\text{MCR}(\mathcal{Y}^{\text{mode}}, \hat{\mathcal{Y}}^{\text{mode}}) \leq \text{MCR}(\mathcal{Y}^{\text{median}}, \hat{\mathcal{Y}}^{\text{median}})$ (Figures 3c-d vs. Figures 4c-d). This tendency is from the property that the median estimator $\hat{y}_\omega^{(\text{median})}$ minimizes $R_1(z) = \mathbb{E}_{y_\omega} |y_\omega - z|$, whereas the mode estimator $\hat{y}_\omega^{(\text{mode})}$ minimizes $R_2(z) = \mathbb{E}_{y_\omega} \mathbb{1}\{y_\omega = z\}$. Here the expectation is taken over the testing data y_ω , conditional on the training data and thus parameters $\hat{\Theta}$ and $\hat{b}$.

7. Data Applications

We apply our ordinal tensor method to two real-world datasets. In the first application, we use our model to analyze an ordinal tensor consisting of structural connectivities among 68 brain regions for 136 individuals from Human Connectome Project (HCP) (Geddes, 2016). In the second application, we perform tensor completion to an ordinal dataset with missing values. The data tensor records the ratings of 139 songs on a scale of 1 to 5 from 42 users on 26 contexts (Baltrunas et al., 2011).

7.1. Human Connectome Project (HCP)

Each entry in the HCP dataset takes value on a nominal scale, $\{\text{high}, \text{moderate}, \text{low}\}$, indicating the strength level of fiber connection. We convert the dataset to a 3-level ordinal tensor $\mathcal{Y} \in [3]^{68 \times 68 \times 136}$ and apply the ordinal tensor method with a logistic link function. The BIC suggests $r = (23, 23, 8)$ with $\mathcal{L}_{\mathcal{Y}, \Omega}(\hat{\Theta}, \hat{b}) = -216, 646$. Based on

CLUSTER	I		
BRAIN NODES	L.FRONTALPOLE, L.TEMPORALPOLE, L.MEDIALORBITOFRONTAL, L.CUNEUS, L.PARAHIPPOCAMPAL, L.LINGUAL, R.FRONTALPOLE, R.TEMPORALPOLE, R.MEDIALORBITOFRONTAL, R.CUNEUS, R.PARAHIPPOCAMPAL		
CLUSTER	II		
BRAIN NODES	L.CAUDALMIDDLEFRONTAL, L.INFERIORPARIETAL, L.INSULA, L.ISTHMUSCINGULATE, L.LATERALOCIPITAL(2), L.PARSOPERCULARIS, L.PARSTRIANGULARIS, L.POSTCENTRAL, L.PRECUNEUS, L.SUPERIORFRONTAL, L.SUPERIORETEMPORAL(3)		
CLUSTER	III		
BRAIN NODES	R.CAUDALMIDDLEFRONTAL, R.INFERIORPARIETAL, R.INSULA, R.ISTHMUSCINGULATE, R.LATERALOCIPITAL(2), R.LINGUAL, R.PARSOPERCULARIS, R.PARSTRIANGULARIS, R.POSTCENTRAL, R.PRECENTRAL, R.PRECUNEUS, R.SUPERIORFRONTAL(3), R.SUPERIORPARIETAL, R.SUPERIORETEMPORAL(3)		
CLUSTER	IV	V	VI
BRAIN NODES	L.SUPRAMARGINAL(4)	L.INFERIORETEMPORAL(3)	L.MIDDELETEMPORAL(3)
CLUSTER	VII	VIII	VIII
BRAIN NODES	R.SUPRAMARGINAL(4)	R.INFERIORETEMPORAL(3)	R.MIDDELETEMPORAL(3)
CLUSTER	X	XI	
BRAIN NODES	L.SUPERIORFRONTAL(2)	L.PRECENTRAL, L.SUPERIORPARIETAL	

Table 2. Node clusters in the HCP analysis. The first alphabet in the node name indicates the left (L) or right (R) hemisphere. The number in the parentheses indicates the node count in each cluster.

the estimated Tucker factors $\{\tilde{M}_k\}$, we perform a clustering analysis via K-mean on the brain nodes (see detailed procedure in the Supplement). We find that the clustering successfully captures the spatial separation between brain regions (Table 2). In particular, cluster I represents the connection between the left and right hemispheres, whereas clusters II-III represent the connection within each of the half brains (Figure 5). Other smaller clusters represent local regions driving by similar nodes (Table 2). For example, the cluster IV/VII consists of nodes in the supramarginal gyrus region in the left/right hemisphere. This region is known to be involved in visual word recognition and reading (Stoeckel et al., 2009). The identified similarities among nodes without external annotations illustrate the applicability of our method to clustering analysis.

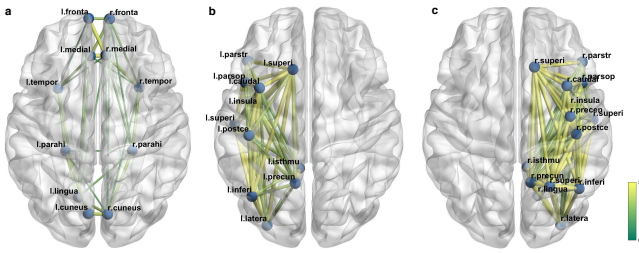


Figure 5. Top three clusters in the HCP analysis. (a) Cluster I reflects the connections between two brain hemispheres. (b)-(c) Cluster II/III consists of nodes within left/right hemisphere only. Node names are shown in abbreviation. Edges are colored based on estimated connection averaged across individuals.

We compare the goodness-of-fit of various tensor methods on the HPC data. Table 3 summarizes the prediction error via 5-fold stratified cross-validation averaged over 10 runs. Our method outperforms the others, especially in MAD.

7.2. InCarMusic recommendation system

We apply ordinal tensor completion to a recommendation system *InCarMusic*. *InCarMusic* is a mobile application

that offers music recommendation to passengers of cars based on contexts (Baltrunas et al., 2011). We conduct tensor completion on the 42-by-139-by-26 tensor with 2,844 observed entries only. Table 3 shows the averaged prediction error via 5-fold cross validation. The high missing rate makes the accurate classification challenging. Nevertheless, our method achieves the best performance among the three.

Human Connectome Project (HCP) dataset		
Method	MAD	MCR
Ordinal-T (ours)	0.1607 (0.0005)	0.1606 (0.0005)
Continuous-T	0.2530 (0.0002)	0.1599 (0.0002)
1bit-sign-T	0.3566 (0.0010)	0.1563 (0.0010)

InCarMusic dataset		
Method	MAD	MCR
Ordinal-T (ours)	1.37 (0.039)	0.59 (0.009)
Continuous-T	2.39 (0.152)	0.94 (0.027)
1bit-sign-T	1.39 (0.003)	0.81 (0.005)

Table 3. Comparison of prediction error in the HPC and *InCarMusic* analyses. Standard errors are reported in parentheses.

8. Conclusions

We have developed a low-rank tensor estimation method based on possibly incomplete, ordinal-valued observations. A sharp error bound is established, and we demonstrate the outperformance of our approach compared to other methods. The work unlocks several directions of future research. One interesting question would be the inference problem, i.e., to assess the uncertainty of the obtained estimates and the imputation. Other directions include the trade-off between (non)convex optimization and statistical efficiency. While we have provided numerical evidence for the success of nonconvex approach, the full landscape of the optimization remains open. The interplay between computational efficiency and statistical accuracy in general tensor problems warrants future research.

Acknowledgements

This research is supported in part by NSF grant DMS-1915978 and Wisconsin Alumni Research Foundation.

References

- Acar, E., Dunlavy, D. M., Kolda, T. G., and Mørup, M. Scalable tensor factorizations with missing data. In *Proceedings of the 2010 SIAM international conference on data mining*, pp. 701–712. SIAM, 2010.
- Baltrunas, L., Kaminskas, M., Ludwig, B., Moling, O., Ricci, F., Aydin, A., Lüke, K.-H., and Schwaiger, R. Incarmusic: Context-aware music recommendations in a car. In *International Conference on Electronic Commerce and Web Technologies*, pp. 89–100. Springer, 2011.
- Bhaskar, S. A. Probabilistic low-rank matrix completion from quantized measurements. *The Journal of Machine Learning Research*, 17(1):2131–2164, 2016.
- Bhaskar, S. A. and Javanmard, A. 1-bit matrix completion under exact low-rank constraint. In *2015 49th Annual Conference on Information Sciences and Systems (CISS)*, pp. 1–6. IEEE, 2015.
- Cai, T. and Zhou, W.-X. A max-norm constrained minimization approach to 1-bit matrix completion. *The Journal of Machine Learning Research*, 14(1):3619–3647, 2013.
- Davenport, M. A., Plan, Y., Van Den Berg, E., and Wootters, M. 1-bit matrix completion. *Information and Inference: A Journal of the IMA*, 3(3):189–223, 2014.
- De Silva, V. and Lim, L.-H. Tensor rank and the ill-posedness of the best low-rank approximation problem. *SIAM Journal on Matrix Analysis and Applications*, 30(3):1084–1127, 2008.
- Geddes, L. Human brain mapped in unprecedented detail. *Nature*, 2016. doi: 10.1038/nature.2016.20285.
- Ghadermarzy, N., Plan, Y., and Yilmaz, O. Learning tensors from partial binary measurements. *IEEE Transactions on Signal Processing*, 67(1):29–40, 2018.
- Ghadermarzy, N., Plan, Y., and Yilmaz, Ö. Near-optimal sample complexity for convex tensor completion. *Information and Inference: A Journal of the IMA*, 8(3):577–619, 2019.
- Hillar, C. J. and Lim, L.-H. Most tensor problems are NP-hard. *Journal of the ACM (JACM)*, 60(6):45, 2013.
- Hitchcock, F. L. The expression of a tensor or a polyadic as a sum of products. *Journal of Mathematics and Physics*, 6(1-4):164–189, 1927.
- Hong, D., Kolda, T. G., and Duersch, J. A. Generalized canonical polyadic tensor decomposition. *SIAM Review*, 62(1):133–163, 2020.
- Kolda, T. G. and Bader, B. W. Tensor decompositions and applications. *SIAM Review*, 51(3):455–500, 2009.
- McCullagh, P. Regression models for ordinal data. *Journal of the Royal Statistical Society: Series B (Methodological)*, 42(2):109–127, 1980.
- Montanari, A. and Sun, N. Spectral algorithms for tensor completion. *Communications on Pure and Applied Mathematics*, 71(11):2381–2425, 2018.
- Mu, C., Huang, B., Wright, J., and Goldfarb, D. Square deal: Lower bounds and improved relaxations for tensor recovery. In *International Conference on Machine Learning*, pp. 73–81, 2014.
- Negahban, S., Wainwright, M. J., et al. Estimation of (near) low-rank matrices with noise and high-dimensional scaling. *The Annals of Statistics*, 39(2):1069–1097, 2011.
- Nickel, M., Tresp, V., and Kriegel, H.-P. A three-way model for collective learning on multi-relational data. In *International Conference on Machine Learning*, volume 11, pp. 809–816, 2011.
- Oseledets, I. V. Tensor-train decomposition. *SIAM Journal on Scientific Computing*, 33(5):2295–2317, 2011.
- Stoeckel, C., Gough, P. M., Watkins, K. E., and Devlin, J. T. Supramarginal gyrus involvement in visual word recognition. *Cortex*, 45(9):1091–1096, 2009.
- Sur, P. and Candès, E. J. A modern maximum-likelihood theory for high-dimensional logistic regression. *Proceedings of the National Academy of Sciences*, 116(29):14516–14525, 2019.
- Tucker, L. R. Some mathematical notes on three-mode factor analysis. *Psychometrika*, 31(3):279–311, 1966.
- Wang, M. and Li, L. Learning from binary multiway data: Probabilistic tensor decomposition and its statistical optimality. *The Journal of Machine Learning Research*, to appear, *arXiv preprint arXiv:1811.05076*.
- Wang, M. and Song, Y. Tensor decompositions via two-mode higher-order SVD (HOSVD). In *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, pp. 614–622, 2017.
- Wang, M. and Zeng, Y. Multiway clustering via tensor block models. In *Advances in Neural Information Processing Systems*, pp. 713–723, 2019.

- Wang, M., Fischer, J., and Song, Y. S. Three-way clustering of multi-tissue multi-individual gene expression data using semi-nonnegative tensor decomposition. *The Annals of Applied Statistics*, 13(2):1103–1127, 2019.
- Yuan, M. and Zhang, C.-H. On tensor completion via nuclear norm minimization. *Foundations of Computational Mathematics*, 16(4):1031–1068, 2016.
- Zhang, A. Cross: Efficient low-rank tensor completion. *The Annals of Statistics*, 47(2):936–964, 2019.
- Zhou, H., Li, L., and Zhu, H. Tensor regression with applications in neuroimaging data analysis. *Journal of the American Statistical Association*, 108(502):540–552, 2013.