

Tensor Deflation for CANDECOMP/PARAFAC— Part II: Initialization and Error Analysis

Anh-Huy Phan, *Member, IEEE*, Petr Tichavský, *Senior Member, IEEE*, and Andrzej Cichocki, *Fellow, IEEE*

Abstract—In Part I of the study of the tensor deflation for CANDECOMP/PARAFAC, we have shown that the rank-1 tensor deflation is applicable under some conditions. Part II of the study presents several initialization algorithms suitable for the algorithm proposed in Part I. In addition, Part II contains an algorithm for the case when one or more factor matrices in the estimated model is constrained to be orthogonal. Finally, Part II provides an error analysis of the tensor deflation algorithm, which shows that there is a marginal loss of accuracy of the deflation algorithm compared to the ordinary CP decomposition.

Index Terms—CANDECOMP/PARAFAC, canonical polyadic decomposition (CPD), Cramér-Rao lower bound, tensor deflation.

I. INTRODUCTION

TENSOR deflation in this paper and the papers [1]–[3] is developed for tensors of rank- R or approximate rank- R where R does not exceed the tensor dimensions, or tensors comprising rank-1 tensors and multilinear rank tensors in their block forms. The major aim of tensor deflation is to extract one rank-1 tensor from a high rank tensor instead of estimating all components in the CANDECOMP/PARAFAC (CP) decomposition. The tensor deflation can be used in a sequential process to extract several rank-1 tensors. This kind of tensor decomposition has applications in extraction of one or a few rank-1 tensors from tensors of high rank, or tracking markers (components) of interest in an online learning system.

In the Part I of this work, we have presented the Alternating Subspace Update (ASU) algorithm for rank-1 tensor deflation, and a sequential process to extract a few rank-1 tensors. The tensor deflation is related to the rank-1 plus multilinear rank- $(R - 1, \dots, R - 1)$ block term decomposition (BTD) [4]. Practical results show that this kind of block term decomposition and the generalized rank- (L_r, M_r, N_r) BTD often get stuck in false local minima [4], [5]. In order to improve reliability of

algorithms, one can perform the decomposition with different initial points, then chooses the result achieving the lowest approximation error. So far, multiple random initialization is most commonly used although this method is not much efficient. In order to achieve good performance with fast convergence, besides good optimization criteria, good initialization is always prerequisite. This second part will introduce new initialization methods for tensor deflation, which can find solutions for the noise-free case. In addition, the paper simplifies the model to have one or several orthogonal factor matrices. The orthogonality constraint in factor matrices improves uniqueness of the CANDECOMP/PARAFAC decomposition (CPD), and avoids degeneracy which often occurs in CPD [6]. Last but not least, we provide Cramér-Rao Lower Bound (CRB) for the tensor deflation. The CRB will serve as an accuracy indicator of the estimation of rank-1 tensor. In comparison with the CRB for CPD [7], we found that there is no significant loss in accuracy of tensor deflation using our approach.

Notation used in this paper is similar to that in Part I [1], [8]. For example, we shall denote tensors by bold calligraphic letters, e.g., $\mathcal{A} \in \mathbb{R}^{I_1 \times I_2 \times \dots \times I_N}$, matrices by bold capital letters, e.g., $\mathbf{A} = [\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_R] \in \mathbb{R}^{I \times R}$, and vectors by bold italic letters, e.g., $\mathbf{a}_j$. The Kronecker product is denoted by $\otimes$. The mode- n matricization of tensor $\mathcal{Y}$ is denoted by $\mathbf{Y}_{(n)}$. The mode- n multiplication of a tensor $\mathcal{Y} \in \mathbb{R}^{I_1 \times I_2 \times \dots \times I_N}$ by a matrix $\mathbf{U} \in \mathbb{R}^{I_n \times R}$ is denoted by

$$\mathcal{Z} = \mathcal{Y} \times_n \mathbf{U} \in \mathbb{R}^{I_1 \times \dots \times I_{n-1} \times R \times I_{n+1} \times \dots \times I_N}.$$

The rest of the paper is organized as follows. Some properties of the exact rank-1 tensor extraction are presented in Section II. Section III presents initialization methods based on the singular value decomposition (SVD), the joint eigenvalue decomposition. In Section IV, a variant of the ASU algorithm is developed for the case when one factor matrix is orthogonal. Cramér-Rao lower bound for the tensor deflation is derived in Section V. Simulations in Section VI verify initialization algorithms and analyze the loss of accuracy of the tensor deflation compared to the ordinary CP decomposition. Section VII concludes the paper.

II. EXACT RANK-1 TENSOR EXTRACTION

As in Part I [1], for the rank-1 tensor deflation, we consider an order- N tensor $\mathcal{Y}$ of size $R \times R \times \dots \times R$, and a tensor decomposition of this tensor $\mathcal{Y}$ into two tensor blocks of rank-1 and multilinear rank- $(R - 1, \dots, R - 1)$

$$\mathcal{Y} \approx \llbracket \lambda; \mathbf{a}^{(1)}, \mathbf{a}^{(2)}, \dots, \mathbf{a}^{(N)} \rrbracket + \llbracket \mathcal{G}; \mathbf{U}^{(1)}, \mathbf{U}^{(2)}, \dots, \mathbf{U}^{(N)} \rrbracket, \quad (1)$$

Manuscript received September 25, 2014; revised March 26, 2015 and June 16, 2015; accepted July 04, 2015. Date of publication July 20, 2015; date of current version October 06, 2015. The associate editor coordinating the review of this manuscript and approving it for publication was Prof. Adel Belouchrani. The work of P. Tichavský was supported by the Czech Science Foundation through project No. 14-13713S.

A.-H. Phan is with the Lab for Advanced Brain Signal Processing, Brain Science Institute, RIKEN, Wako 351-0198, Japan (e-mail: phan@brain.riken.jp).

P. Tichavský is with the Institute of Information Theory and Automation of the Czech Academy of Sciences, Prague, 182 08, Czech Republic (e-mail: tichavsk@utia.cas.cz).

A. Cichocki is with the is with the Lab for Advanced Brain Signal Processing, RIKEN BSI, Wako 351-0198, Japan, and also with the Systems Research Institute, PAS, Warsaw 00-447, Poland, and SKOLTECH, Moscow 143026, Russia (e-mail: a.cichocki@riken.jp).

Color versions of one or more of the figures in this paper are available online at <http://ieeexplore.ieee.org>.

Digital Object Identifier 10.1109/TSP.2015.2458789

where $\mathbf{U}^{(n)}$ are of size $R \times (R - 1)$, and the core tensor $\mathcal{G}$ is of size $(R - 1) \times (R - 1) \times \dots \times (R - 1)$. If the tensor has sizes $I_n > R$, it should be compressed to the above size using the Tucker decomposition [9]. The CP decomposition and the rank-1 tensor extraction procedure are illustrated in [1, Fig. 2 of Part I].

According to the orthogonal normalization and rotation (Lemma 1 in [1]), the components $\mathbf{a}^{(n)}$ and factor matrices $\mathbf{U}^{(n)}$ can be assumed to be orthogonal, i.e., $(\mathbf{U}^{(n)})^T \mathbf{U}^{(n)} = \mathbf{I}_{R-1}$, $\rho_n = (\mathbf{u}_1^{(n)})^T \mathbf{a}^{(n)} \geq 0$ and $(\mathbf{u}_r^{(n)})^T \mathbf{a}^{(n)} = 0$ for $r = 2, \dots, R - 1$.

In this section, we present relation between components and factor matrices for tensors which admit the model in (1). The properties will be exploited to derive efficient initialization algorithms in Section III.

As similar to deriving the ASU algorithm in Part I [1], we express the components $\mathbf{a}^{(n)}$ through vectors $\mathbf{w}_n$ and the first columns $\mathbf{u}_n = \mathbf{u}_1^{(n)}$ of $\mathbf{U}^{(n)}$, that is

$$\mathbf{a}^{(n)} = \sqrt{1 - \rho_n^2} \mathbf{w}_n + \rho_n \mathbf{u}_n, \quad (2)$$

where $\mathbf{V}_n = [\mathbf{w}_n, \mathbf{U}^{(n)}]$ are orthonormal matrices, i.e., $\mathbf{V}_n^T \mathbf{V}_n = \mathbf{I}_R$. Let

$$\mathbf{C}_n = \mathbf{Y}_{(n)} \mathbf{Y}_{(n)}^T \in \mathbb{R}^{R \times R}, \quad (3)$$

$$\mathbf{s}_n = \mathbf{Y}_{(n)} \left(\bigotimes_{k \neq n} \mathbf{a}^{(k)} \right), \quad n = 1, \dots, N. \quad (4)$$

We will show that parameters $\mathbf{w}_n$, $\mathbf{U}_n$ and ρ_n can be found in closed-form through the components $\mathbf{a}^{(n)}$ for the exact decomposition.

Lemma 1 (Scaling Coefficient of Rank-1 Tensor): The weight λ associated with rank-1 tensor $\llbracket \{\mathbf{a}^{(n)}\} \rrbracket$ is given by

$$\lambda = \frac{1}{\mathbf{s}_n^T \mathbf{C}_n^{-1} \mathbf{a}^{(n)}}, \quad n = 1, 2, \dots, N. \quad (5)$$

Lemma 2 (Closed-Form Expressions for $\mathbf{w}_n$, ρ_n and $\mathbf{u}_n$):

$$\mathbf{w}_n = \frac{\mathbf{C}_n^{-1} \mathbf{s}_n}{\sqrt{\mathbf{s}_n^T \mathbf{C}_n^{-2} \mathbf{s}_n}} \quad (6)$$

$$\rho_n = \sqrt{1 - \frac{1}{\lambda^2 \mathbf{s}_n^T \mathbf{C}_n^{-2} \mathbf{s}_n}} \quad (7)$$

$$\mathbf{u}_n = \frac{\mathbf{a}^{(n)} - \sqrt{1 - \rho_n^2} \mathbf{w}_n}{\rho_n}. \quad (8)$$

Proofs of Lemma 1 and Lemma 2 are in Appendices A and B, respectively. The Lemmas give recipes, for a given rank-1 tensor $\llbracket \{\mathbf{a}^{(n)}\} \rrbracket$, how to compute the other parameters λ , $\mathbf{w}_n$ and $\mathbf{u}_n$. The last $(R - 2)$ columns of $\mathbf{U}^{(n)}$ are arbitrarily orthonormal basis of orthogonal complement to $[\mathbf{w}_n, \mathbf{u}_n]$.

III. INITIALIZATION FOR RANK-1 TENSOR EXTRACTION

A. Initialization Using Leading Singular Values

1) *Pre-Selection of Components:* The simplest initialization method is that all parameters are initialized by random values. Unfortunately, this type of initialization method often makes algorithms getting stuck in false local minima, and thus requires a large number of trials. For tensors which admit the CP decomposition, the vectors $\mathbf{a}^{(n)}$ and factor matrices $\mathbf{U}^{(n)}$ can be initialized using the Direct Trilinear Decomposition (DTLD) [10]

Algorithm 1: Simple Initialization

Input: Data tensor $\mathcal{Y}$: $(R \times R \times \dots \times R)$, rank R

Output: A rank-1 tensor $\llbracket \mathbf{a}^{(1)}, \dots, \mathbf{a}^{(N)} \rrbracket$
and $\{\mathbf{U}^{(1)}, \dots, \mathbf{U}^{(N)}\}$

begin

- 1 Initialize $\mathbf{V}_n = [\mathbf{v}_{n1}, \mathbf{v}_{n2}, \dots, \mathbf{v}_{nR}]$ by leading left singular vectors of $\mathbf{Y}_{(n)}$
for $r = 1, 2, \dots, R$ **do**
 for $n = 1, 2, \dots, N$ **do**
 - 2 $\mathbf{a}^{(n)} = \mathbf{v}_{n,r}$,
 - 3 $\mathbf{U}^{(n)} = [\mathbf{v}_{n,1}, \dots, \mathbf{v}_{n,r-1}, \mathbf{v}_{n,r+1}, \dots, \mathbf{v}_{n,R}]$
 - % Run ASU with a small number of iterations
 - 4 $\llbracket \{\mathbf{a}^{(n)}\}, \{\mathbf{U}^{(n)}\}, \varepsilon_r \rrbracket = \text{asu}(\mathcal{Y}, \{\mathbf{a}^{(n)}\}, \{\mathbf{U}^{(n)}\})$
 - 5 Choose the initial point yielding the smallest approximation error $\varepsilon_{r^*} = \min\{\varepsilon_1, \varepsilon_2, \dots, \varepsilon_R\}$
-

or the generalized rank annihilation method for (higher order) CPD [11], or simply but less efficiently using leading singular vectors of the tensor along modes, i.e., using the Higher Order SVD (HOSVD) or Higher Order Orthogonal Iteration (HOOI) algorithm [9].

For deflation of tensors which admit the block component model, we initialize components $\mathbf{a}^{(n)}$ of the rank-1 tensor, and factor matrices $\mathbf{U}^{(n)}$ of the multilinear rank- $(R - 1)$ tensor, but do not need to initialize the weight coefficient λ of the rank-1 tensor and the core tensor $\mathcal{G}$ of the second block.

If one block is considered to be fixed, then the other block can be obtained as the best rank-1 or multilinear rank- $(R - 1)$ tensor of a certain residue tensor. One can also initialize components using leading singular eigenvectors of mode- n matricizations of $\mathcal{Y}$ as in CPD or other matrix/tensor decompositions. However, simply choosing the first $(R - 1)$ leading singular vectors for the factor matrices $\mathbf{U}^{(n)}$ and one for $\mathbf{a}^{(n)}$ may not give a good initial point. Instead, one can choose the components $\mathbf{a}^{(n)}$ as one of R leading singular eigenvectors of $\mathbf{Y}_{(n)}$, and the rest $(R - 1)$ leading singular eigenvectors for $\mathbf{U}^{(n)}$ so that the ASU algorithm obtains the smallest approximation error after a small number of iterations (say 10). There are all R^N possible combinations of leading singular eigenvectors in all modes, but we only consider R combinations of singular eigenvectors in the same order. The procedure is summarized in Algorithm 1. Basically, complexity of the preselection is the same as that of the ASU algorithm, i.e., $\mathcal{O}(R^N)$ [1]. Together with the selection one among R rank-1 tensors, the computational cost is at most $\mathcal{O}(R^{N+1})$. We note that since ASU in each run in Step 3 executes only a few iterations, the preselection, in general, is still fast. Experiment results show that the procedure improves the initial points, and thereby improves the obtained solution.

2) *SVD-Based Initialization With Estimation of the Parameters ρ_n :* According to the orthogonal normalization, we can initialize $\mathbf{a}^{(n)}$ and $\mathbf{U}^{(n)}$ by orthogonal vectors which are naturally the leading left singular vectors of mode- n matricizations $\mathbf{Y}_{(n)}$. For tensors which admit the CP decomposition, the method works well. However, when the tensors do not fully admit CPD, the method may be less efficient. Instead, we initialize $\mathbf{w}_n$ with a singular vector and $\mathbf{U}^{(n)}$ by the rest $(R - 1)$

Algorithm 2: SVD Initialization with Estimation of Parameters

 ρ_n
Input: Data tensor $\mathcal{Y}$: $(R \times R \times \dots \times R)$, rank R
Output: A rank-1 tensor $\llbracket \mathbf{a}^{(1)}, \dots, \mathbf{a}^{(N)} \rrbracket$ and $\{\mathbf{U}^{(1)}, \dots, \mathbf{U}^{(N)}\}$
begin

- 1 Initialize $\llbracket \mathbf{w}_n, \mathbf{U}^{(n)} \rrbracket$ by singular eigenvectors of $\mathbf{Y}_{(n)}$
- 2 Initialize $0 \leq \rho_n < 1$ /* $\xi_n = \sqrt{1 - \rho_n^2}$ */
- 3 $\mathcal{X} = \mathcal{Y} \times \{\llbracket \mathbf{w}_n, \mathbf{u}_n \rrbracket^T\}$ /* $\mathbf{u}_n = \mathbf{U}^{(n)}(:, 1)$ */
- repeat**
- for** $n = 1, 2, \dots, N$ **do**
- 4 $\llbracket x_{n1}, x_{n2} \rrbracket^T = \mathcal{X} \times_{k \neq n} \{\llbracket \xi_n, \rho_n \rrbracket^T\}$
 $\quad - \llbracket 0, \rho_{-n} \mathcal{X}(2, \dots, 2) \rrbracket^T$
- 5 $\rho_n = \frac{|x_{n2}|}{\sqrt{x_{n1}^2(1 - \rho_{-n}^2) + x_{n2}^2}}$ /* $\rho_{-n} = \prod_{k \neq n} \rho_k$ */
- until** a stopping criterion is met
- 6 **for** $n = 1, \dots, N$ **do** $\mathbf{a}^{(n)} = \sqrt{1 - \rho_n^2} \mathbf{w}_n + \rho_n \mathbf{u}_n$

ones where $\mathbf{a}^{(n)} = \xi_n \mathbf{w}_n + \rho_n \mathbf{u}_n$, $\xi_n = \sqrt{1 - \rho_n^2}$, $\mathbf{u}_n$ are the first column vectors of $\mathbf{U}^{(n)}$ for $n = 1, \dots, N$. Then we suggest to seek parameters $\rho_1, \rho_2, \dots, \rho_N$ which can minimise the approximation error

$$\begin{aligned} D &= \frac{1}{2} \|\mathcal{Y} - \llbracket \lambda; \{\mathbf{a}^{(n)}\} \rrbracket - \llbracket \mathcal{G}; \{\mathbf{U}^{(n)}\} \rrbracket\|_F^2 \\ &= \frac{1}{2} \left\| \mathcal{Y} - \llbracket \mathcal{Y}; \{\mathbf{U}^{(n)} \mathbf{U}^{(n)T}\} \rrbracket - \lambda \left(\llbracket \{\mathbf{a}^{(n)}\} \rrbracket - \llbracket \rho; \{\mathbf{u}_n\} \rrbracket \right) \right\|_F^2, \end{aligned}$$

where $\mathcal{G}$ is substituted by its optimum counterpart $\mathcal{G}_{opt} = \mathcal{T} - \lambda \rho t_1 \mathcal{E}_1$ [1], $\rho = \prod_n \rho_n$, $\mathcal{T} = \mathcal{Y} \times \{\mathbf{U}^{(n)T}\}$, and $\mathcal{E}_1$ is tensor of the same size as $\mathcal{G}$ but has only one entry $\mathcal{E}_1(1) = 1$, and zeros elsewhere. As similarly to deriving in the ASU algorithm [1], the optimum weight λ_{opt} is given in closed-form as

$$\lambda_{opt} = \frac{\mathcal{Y} \times \{\mathbf{a}^{(n)T}\} - \rho t_1}{1 - \rho^2}, \quad (9)$$

where $t_1 = \mathcal{Y} \times_1 \mathbf{u}_1^T \cdots \times_N \mathbf{u}_N^T$ is the first entry of $\mathcal{T}$. Parameters ρ_n for $n = 1, \dots, N$ are sequentially updated as

$$\rho_n = \frac{|x_{n2}|}{\sqrt{x_{n1}^2(1 - \rho_{-n}^2) + x_{n2}^2}}, \quad (10)$$

where $\rho_{-n} = \prod_{k \neq n} \rho_k$, and

$$\begin{bmatrix} x_{n1} \\ x_{n2} \end{bmatrix} = \left(\mathcal{X} \times_{k \neq n} \left\{ \begin{bmatrix} \xi_k \\ \rho_k \end{bmatrix} \right\} \right) - \begin{bmatrix} 0 \\ \rho_{-n} t_1 \end{bmatrix} \quad (11)$$

with $\mathcal{X} = \mathcal{Y} \times \{\llbracket \mathbf{w}_n, \mathbf{u}_n \rrbracket^T\}$. Note that $t_1 = \mathcal{X}(2, \dots, 2)$.

Pseudo-code of the SVD-based initialization algorithm is summarized in Algorithm 2. The initialization in Step 1 computes the left singular vectors of $\mathbf{Y}_{(n)}$ with complexity of $\mathcal{O}(R^3)$. The tensor $\mathcal{X}$ computed only once in Step 3 is of size $2 \times 2 \times \dots \times 2$ with a cost of $\mathcal{O}(R^N)$. Small loop (Steps 4 and 5) runs a few iterations, usually less than 4, with a low cost per iteration of $\mathcal{O}(2^N)$. Therefore, Algorithm 2 has complexity of $\mathcal{O}(R^N)$ which is dominated by computation of the tensor $\mathcal{X}$.

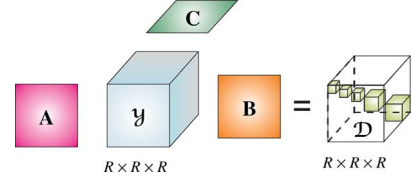


Fig. 1. Illustration of the tensor diagonalization as tool for initialization of the tensor deflation.

B. Initialization Using the Tensor Diagonalization (TEDIA)

An alternative initialization for the ASU algorithm consists in utilizing the tensor diagonalization TEDIA [5]. TEDIA was developed so far only for order-3 tensors. The method seeks factor matrices $\tilde{\mathbf{A}}$, $\tilde{\mathbf{B}}$ and $\tilde{\mathbf{C}}$ as shown in Fig. 1 which transform the tensor $\mathcal{Y}$ into a block diagonal tensor by minimizing the sum of squared off-diagonal elements of the tensor $\mathcal{T} = \mathcal{Y} \times_1 \tilde{\mathbf{A}} \times_2 \tilde{\mathbf{B}} \times_3 \tilde{\mathbf{C}}$. By inspecting the diagonal $\mathcal{T}(r, r, r)$ and corresponding off-diagonal coefficients, we can find suitable components $\mathbf{A}(:, r)$, $\mathbf{B}(:, r)$ and $\mathbf{C}(:, r)$ for the rank-1 tensor where $\mathbf{A} = \tilde{\mathbf{A}}^{-1}$, and the rest components of $\mathbf{A}$, $\mathbf{B}$ and $\mathbf{C}$ for the second block. Each iteration of the TEDIA algorithm costs $\mathcal{O}(R^4)$ as the same fast ALS algorithm for CPD. In simulations (see Section IV), the initialization by TEDIA alone looks costly, but since it typically provides a better initial solution than the other algorithms, TEDIA + ASU has the complexity comparable to ASU with other initialization methods.

C. Initialization Based on Joint Eigenvalue Decomposition

In this section, we first introduce rank-1 deflation of two matrices based on the generalized eigenvalue decomposition. The method is then applied to tensors comprising only two slices, compressed from the data tensor. For simplicity, the methods will be derived for order-3 tensors, then extended to higher order tensors. In addition, the components $\mathbf{a}$, $\mathbf{b}$ and $\mathbf{c}$ are used in place of $\mathbf{a}^{(1)}$, $\mathbf{a}^{(2)}$, $\mathbf{a}^{(3)}$.

1) Rank-1 Deflation of Two Matrices:

Theorem 1: Two real-valued full-rank matrices $\mathbf{X}_1$ and $\mathbf{X}_2$ of size $R \times R$ can be simultaneously decomposed into

$$\mathbf{X}_1 = \alpha_1 \mathbf{a} \mathbf{b}^T + \mathbf{U} \mathbf{G}_1 \mathbf{V}^T, \quad (12)$$

$$\mathbf{X}_2 = \alpha_2 \mathbf{a} \mathbf{b}^T + \mathbf{U} \mathbf{G}_2 \mathbf{V}^T, \quad (13)$$

where $\mathbf{a}$ and $\mathbf{b}$ are eigenvectors associated with the same eigenvalue λ for matrix pencils $(\mathbf{X}_2^{-1}, \mathbf{X}_1^{-1})$ and $(\mathbf{X}_2^{-1T}, \mathbf{X}_1^{-1T})$, $\mathbf{G}_1$ and $\mathbf{G}_2$ are matrices of size $(R-1) \times (R-1)$ and have rank $(R-1)$.

Proof of Theorem 1 is in Appendix C.

2) Multiple Rank-1 Deflations of Two Matrices: For deflation of a tensor $\mathcal{Y}$, one can use the Tucker decomposition to compress the tensor $\mathcal{Y}$ to order-3 tensors $\mathcal{X}$ consisting of only two slices, e.g., $R \times R \times 2$, or $R \times 2 \times R$, and apply Theorem 1 to the two slices to obtain estimates of the components $\mathbf{a}^{(n)}$ and matrices $\mathbf{U}^{(n)}$. The tensor $\mathcal{X}$ can be of lower order than $\mathcal{Y}$. For simplicity, we consider an order-3 tensor $\mathcal{Y}$, and compress this tensor along mode-3 to yield two frontal slices $\mathbf{X}_1^{(1,2)}$ and $\mathbf{X}_2^{(1,2)}$ of size $R \times R$ which share a common rank-1 matrix

$$\mathbf{X}_1^{(1,2)} \approx \alpha_1^{(1,2)} \mathbf{a} \mathbf{b}^T + \mathbf{U}_1 \check{\mathbf{G}}_1 \mathbf{U}_2^T,$$

$$\mathbf{X}_2^{(1,2)} \approx \alpha_2^{(1,2)} \mathbf{a} \mathbf{b}^T + \mathbf{U}_1 \check{\mathbf{G}}_2 \mathbf{U}_2^T,$$

where $\check{\mathbf{G}}_1$ and $\check{\mathbf{G}}_2$ are of size $(R-1) \times (R-1)$, and $\mathbf{a}$, $\mathbf{b}$ can be approximated by generalized eigenvectors $\mathbf{a}_{1,2}$, $\mathbf{b}_{1,2}$ of matrix pencils $(\mathbf{X}_1^{(1,2)}, \mathbf{X}_2^{(1,2)})$ and $(\mathbf{X}_1^{(1,2)T}, \mathbf{X}_2^{(1,2)T})$. For real-valued tensors, there may not be any real-valued eigenvectors. One can convert the complex eigenvalues to real 2-by-2 blocks on the diagonal using Matlab's routine "cdf2rdf". Even when real eigenvalues exist, there may have multiple pairs of eigenvectors $(\mathbf{a}_r^{(1,2)}, \mathbf{b}_r^{(1,2)})$ which are associated with the same real eigenvalues. In order to select a good estimate to $\mathbf{a}$ and $\mathbf{b}$, we perform a further compression of $\mathcal{Y}$ along mode-2 (or mode-3) and obtain estimates to $\mathbf{a}$ and $\mathbf{c}$ as eigenvectors $(\mathbf{a}_s^{(1,3)}, \mathbf{c}_s^{(1,3)})$. Loading components $(\mathbf{a}_r^{(1,2)}, \mathbf{b}_r^{(1,2)}, \mathbf{c}_s^{(1,3)})$ are chosen such that $\mathbf{a}_r^{(1,2)}$ is most correlated with $\mathbf{a}_s^{(1,3)}$, that is $\max (\mathbf{a}_r^{(1,2)})^T \mathbf{a}_s^{(1,3)}$.

Alternatively, one can perform three compressions to obtain pairs of eigenvectors $(\mathbf{a}_r^{(1,2)}, \mathbf{b}_r^{(1,2)})$, $(\mathbf{a}_s^{(1,3)}, \mathbf{c}_s^{(1,3)})$ and $(\mathbf{b}_t^{(2,3)}, \mathbf{c}_t^{(2,3)})$, and select a good estimate to $(\mathbf{a}, \mathbf{b}, \mathbf{c})$ such that

$$\max_{r,s,t} ((\mathbf{a}_r^{(1,2)})^T \mathbf{a}_s^{(1,3)}) ((\mathbf{b}_r^{(1,2)})^T \mathbf{b}_t^{(2,3)}) ((\mathbf{c}_s^{(1,3)})^T \mathbf{c}_t^{(2,3)}). \quad (14)$$

Matrices $\mathbf{U}_1$, $\mathbf{U}_2$ and $\mathbf{U}_3$ can be computed as in proof of Theorem 1, or using the method in Section II. Computational cost of this initialization method is dominated by the compression of the tensor of size $R \times R \times R$ to ones of two slices. This requires a cost of order $\mathcal{O}(R^3)$.

3) *Joint Eigenvalue Decomposition-Based Initialization (JEVD)*: The following initialization works with three compressions of the original tensor $\mathcal{Y}$, including $\mathcal{X}^{(1,2)}$ of size $R \times R \times 2$, $\mathcal{X}^{(1,3)}$ of size $R \times 2 \times R$ and $\mathcal{X}^{(2,3)}$ of size $2 \times R \times R$. We define matrices $\mathbf{F}_{i,j}$, for $1 \leq i \neq j \leq 3$ as

$$\begin{aligned} \mathbf{F}_{12} &= \mathbf{X}_1^{(1,2)} (\mathbf{X}_2^{(1,2)})^{-1}, & \mathbf{F}_{21} &= (\mathbf{X}_1^{(1,2)})^T (\mathbf{X}_2^{(1,2)T})^{-1}, \\ \mathbf{F}_{13} &= \mathbf{X}_1^{(1,3)} (\mathbf{X}_2^{(1,3)})^{-1}, & \mathbf{F}_{31} &= (\mathbf{X}_1^{(1,3)})^T (\mathbf{X}_2^{(1,3)T})^{-1}, \\ \mathbf{F}_{23} &= \mathbf{X}_1^{(2,3)} (\mathbf{X}_2^{(2,3)})^{-1}, & \mathbf{F}_{32} &= (\mathbf{X}_1^{(2,3)})^T (\mathbf{X}_2^{(2,3)T})^{-1}, \end{aligned}$$

where $\mathbf{X}_1^{(n,m)}$ and $\mathbf{X}_2^{(n,m)}$ are two slices of the tensor $\mathcal{X}^{(n,m)}$. As stated in Theorem 1, $\mathbf{a}$, $\mathbf{b}$ and $\mathbf{c}$ can be considered joint eigenvectors of matrices $\mathbf{F}_{i,j}$ and $\mathbf{F}_{j,i}$, that is,

$$\begin{aligned} \mathbf{F}_{12} \mathbf{a} &= \mathbf{a} \lambda_{12}, & \mathbf{F}_{13} \mathbf{a} &= \mathbf{a} \lambda_{13} \\ \mathbf{F}_{21} \mathbf{b} &= \mathbf{b} \lambda_{21}, & \mathbf{F}_{23} \mathbf{b} &= \mathbf{b} \lambda_{23} \\ \mathbf{F}_{31} \mathbf{c} &= \mathbf{c} \lambda_{31}, & \mathbf{F}_{32} \mathbf{c} &= \mathbf{c} \lambda_{32} \end{aligned}$$

with $\lambda_{12} = \lambda_{21}$, $\lambda_{13} = \lambda_{31}$ and $\lambda_{23} = \lambda_{32}$, $\mathbf{a}^T \mathbf{a} = 1$, $\mathbf{b}^T \mathbf{b} = 1$ and $\mathbf{c}^T \mathbf{c} = 1$. It follows that we can find unit-length vectors $\mathbf{a}$, $\mathbf{b}$ and $\mathbf{c}$ as solution to the following minimisation problem

$$\begin{aligned} \min_{\mathbf{a}, \mathbf{b}, \mathbf{c}} & \|\mathbf{F}_{12} \mathbf{a} - (\mathbf{a}^T \mathbf{F}_{12} \mathbf{a}) \mathbf{a}\|^2 + \|\mathbf{F}_{13} \mathbf{a} - (\mathbf{a}^T \mathbf{F}_{13} \mathbf{a}) \mathbf{a}\|^2 \\ & + \|\mathbf{F}_{21} \mathbf{b} - (\mathbf{b}^T \mathbf{F}_{21} \mathbf{b}) \mathbf{b}\|^2 + \|\mathbf{F}_{23} \mathbf{b} - (\mathbf{b}^T \mathbf{F}_{23} \mathbf{b}) \mathbf{b}\|^2 \\ & + \|\mathbf{F}_{31} \mathbf{c} - (\mathbf{c}^T \mathbf{F}_{31} \mathbf{c}) \mathbf{c}\|^2 + \|\mathbf{F}_{32} \mathbf{c} - (\mathbf{c}^T \mathbf{F}_{32} \mathbf{c}) \mathbf{c}\|^2 \\ & + \|\mathbf{a}^T \mathbf{F}_{12} \mathbf{a} - \mathbf{b}^T \mathbf{F}_{21} \mathbf{b}\|^2 + \|\mathbf{a}^T \mathbf{F}_{13} \mathbf{a} - \mathbf{c}^T \mathbf{F}_{31} \mathbf{c}\|^2 \\ & + \|\mathbf{b}^T \mathbf{F}_{23} \mathbf{b} - \mathbf{c}^T \mathbf{F}_{32} \mathbf{c}\|^2, \end{aligned}$$

Algorithm 3: Joint EVD-based Initialization

Input: Data tensor $\mathcal{Y}$: $(R \times R \times \dots \times R)$
Output: $\{\mathbf{a}^{(1)}, \dots, \mathbf{a}^{(N)}\}$ and $\{\mathbf{u}_1, \dots, \mathbf{u}_N\}$
begin
 % **Stage 1:** Estimate components $\mathbf{a}^{(n)}$ -----
 1 Initialize $\llbracket \{\mathbf{a}^{(n)}\} \rrbracket$
 2 Construct matrices $\mathbf{F}_{n,m}$ for $1 \leq n \neq m \leq N$
 repeat
 for $n = 1, 2, \dots, N$ **do**
 3 $[\mathbf{a}^{(n)}, \lambda_n] = \text{eigs}(\mathbf{Q}_n, 1, 'SA')$ /* $\mathbf{Q}_n$ in (16) */
 until a stopping criterion is met
 % **Stage 2:** Estimate λ of rank-1 tensor $\llbracket \{\mathbf{a}^{(n)}\} \rrbracket$ -----
 4 $\lambda = \frac{1}{N} \sum_n \frac{1}{s_n^T \mathbf{C}_n^{-1} \mathbf{a}^{(n)}}$
 % **Stage 3:** Estimate $\mathbf{w}_n$, ρ_n and $\mathbf{u}_n$ -----
 5 $\mathbf{w}_n = \frac{\mathbf{C}_n^{-1} \mathbf{s}_n}{\sqrt{\mathbf{s}_n^T \mathbf{C}_n^{-2} \mathbf{s}_n}}$, $\rho_n = \sqrt{1 - \frac{1}{\lambda^2 \mathbf{s}_n^T \mathbf{C}_n^{-2} \mathbf{s}_n}}$ and $\mathbf{u}_n = \frac{\mathbf{a}^{(n)} - \sqrt{1 - \rho_n^2} \mathbf{w}_n}{\rho_n}$
 eigs(C, 1, 'SA'): compute eigenvector associated with the smallest eigenvalue of C.
end

which can be simplified into

$$\begin{aligned} \min_{\mathbf{a}, \mathbf{b}, \mathbf{c}} & \mathbf{a}^T \mathbf{C}_a \mathbf{a} + \mathbf{b}^T \mathbf{C}_b \mathbf{b} + \mathbf{c}^T \mathbf{C}_c \mathbf{c} - 2(\mathbf{a}^T \mathbf{F}_{12} \mathbf{a})(\mathbf{b}^T \mathbf{F}_{21} \mathbf{b}) \\ & - 2(\mathbf{a}^T \mathbf{F}_{13} \mathbf{a})(\mathbf{c}^T \mathbf{F}_{31} \mathbf{c}) - 2(\mathbf{b}^T \mathbf{F}_{23} \mathbf{b})(\mathbf{c}^T \mathbf{F}_{32} \mathbf{c}) \end{aligned}$$

where $\mathbf{C}_a = \mathbf{F}_{12}^T \mathbf{F}_{12} + \mathbf{F}_{13}^T \mathbf{F}_{13}$, $\mathbf{C}_b = \mathbf{F}_{21}^T \mathbf{F}_{21} + \mathbf{F}_{23}^T \mathbf{F}_{23}$ and $\mathbf{C}_c = \mathbf{F}_{31}^T \mathbf{F}_{31} + \mathbf{F}_{32}^T \mathbf{F}_{32}$.

While fixing $\mathbf{b}$ and $\mathbf{c}$, $\mathbf{a}$ is solution to the eigenvalue problem

$$\arg \min_{\mathbf{a}} \mathbf{a}^T \mathbf{Q}_a \mathbf{a}, \quad (15)$$

where $\mathbf{Q}_a = \mathbf{C}_a - (\mathbf{b}^T \mathbf{F}_{21} \mathbf{b})(\mathbf{F}_{12} + \mathbf{F}_{12}^T) - (\mathbf{c}^T \mathbf{F}_{31} \mathbf{c})(\mathbf{F}_{13} + \mathbf{F}_{13}^T)$. That is $\mathbf{a}$ is the eigenvector associated with the smallest eigenvalue of the symmetric matrix $\mathbf{Q}_a$. Similarly, we can iteratively estimate $\mathbf{b}$ and $\mathbf{c}$ as eigenvectors of the matrices $\mathbf{Q}_b$ and $\mathbf{Q}_c$

$$\begin{aligned} \mathbf{Q}_b &= \mathbf{C}_b - (\mathbf{a}^T \mathbf{F}_{12} \mathbf{a})(\mathbf{F}_{21} + \mathbf{F}_{21}^T) - (\mathbf{c}^T \mathbf{F}_{32} \mathbf{c})(\mathbf{F}_{23} + \mathbf{F}_{23}^T), \\ \mathbf{Q}_c &= \mathbf{C}_c - (\mathbf{a}^T \mathbf{F}_{13} \mathbf{a})(\mathbf{F}_{31} + \mathbf{F}_{31}^T) - (\mathbf{b}^T \mathbf{F}_{23} \mathbf{b})(\mathbf{F}_{32} + \mathbf{F}_{32}^T). \end{aligned}$$

For higher order tensors, $\mathbf{a}^{(n)}$ are solved similarly as eigenvectors of the matrices

$$\mathbf{Q}_n = \sum_{m \neq n} \mathbf{F}_{n,m}^T \mathbf{F}_{n,m} - (\mathbf{a}^{(m)T} \mathbf{F}_{m,n} \mathbf{a}^{(m)})(\mathbf{F}_{n,m} + \mathbf{F}_{n,m}^T), \quad (16)$$

where $\mathbf{F}_{n,m} = \mathbf{X}_1^{(n,m)} (\mathbf{X}_2^{(n,m)})^{-1}$ and $\mathbf{F}_{m,n} = (\mathbf{X}_1^{(n,m)})^T (\mathbf{X}_2^{(n,m)T})^{-1}$ for $1 \leq n < m \leq N$.

The algorithm is summarized as Algorithm 3. Components $\mathbf{a}^{(n)}$ and matrices $\mathbf{U}^{(n)}$ can be initialized by eigenvectors in Section III-C2. The cost to construct the matrices $\mathbf{F}_{n,m}$ of size $R \times R$ is substantially due to compressions of the tensor $\mathcal{Y}$ to tensors of two slices $\mathcal{X}^{(n,m)}$ for $1 \leq n < m \leq N$, which are of order $\mathcal{O}(N(N-1)/2 R^N) = \mathcal{O}(N^2 R^N)$. For order-3 tensors, we only need to construct three tensors $\mathcal{X}^{(1,2)}$, $\mathcal{X}^{(1,3)}$ and $\mathcal{X}^{(2,3)}$. Moreover, it is worth noting that the matrices $\mathbf{F}_{n,m}$ and

the terms $\sum_{m \neq n} \mathbf{F}_{n,m}^T \mathbf{F}_{n,m}$ are computed once before the iterations in Step 2. Since matrices $\mathbf{Q}_n$ are of size $R \times R$, each iteration to estimate $\mathbf{a}^{(n)}$ in Step 3 requires an operation of $\mathcal{O}(R^3)$. Therefore, the total cost of JEVD is at most $\mathcal{O}(N^2 R^N)$.

IV. DEFLATION OF TENSOR WITH ORTHOGONAL FACTOR MATRICES

In this section, we present a special case of the rank-1 deflation for tensor decompositions with orthogonality constraints in one or several factor matrices. In the CP tensor decompositions, the orthogonality constraints can be imposed onto components in order to improve the stability of the decomposition, and avoid diverging components [6], [12]. In decomposition of received signals in a direct sequence code division multiple access (DS-CDMA) system [13], spreading codes can be orthogonal. The constraint has found applications for analyzing EEG signals [14], [15].

For simplicity, assuming that the decomposition seeks one orthonormal factor matrix, which is $\mathbf{A}^{(1)}$, i.e., $\mathbf{A}^{(1)T} \mathbf{A}^{(1)} = \mathbf{I}_R$. From the orthogonal normalization, we have that $[\mathbf{a}^{(1)}, \mathbf{U}^{(1)}]$ is an orthonormal matrix. That means $\rho_1 = \mathbf{u}_1^T \mathbf{a}^{(1)} = 0$, and thus, $\rho = 0$. The weight of the rank-1 tensor λ and the core tensor $\mathcal{G}$ are then given by

$$\lambda = \mathcal{Y} \times \{\mathbf{a}^{(n)T}\}, \quad \mathcal{G} = \mathcal{Y} \times \{\mathbf{U}^{(n)T}\}. \quad (17)$$

A. Estimation of the Orthogonal Components $\mathbf{a}^{(1)}$ and $\mathbf{U}^{(1)}$

In order to estimate $\mathbf{a}^{(1)}$ and $\mathbf{U}^{(1)}$, the cost function which minimizes the Frobenius norm between the tensor $\mathcal{Y}$ and its approximate

$$D = \frac{1}{2} \|\mathcal{Y} - [\lambda; \mathbf{a}^{(1)}, \mathbf{a}^{(2)}, \dots, \mathbf{a}^{(N)}]\|_F^2 - \|\mathcal{G}; \mathbf{U}^{(1)}, \mathbf{U}^{(2)}, \dots, \mathbf{U}^{(N)}\|_F^2$$

can be rewritten as

$$\begin{aligned} \min D &= \frac{1}{2} (\|\mathcal{Y}\|_F^2 - \|\mathcal{G}\|_F^2 - \lambda^2) \\ &= \frac{1}{2} (\|\mathcal{Y}\|_F^2 - \|\mathbf{U}^{(1)} \mathbf{Z}_1\|_F^2 - (\mathbf{w}_1^T \mathbf{s}_1)^2) \\ &= \frac{1}{2} (\|\mathcal{Y}\|_F^2 - \|\mathbf{Z}_1\|_F^2 + \mathbf{w}_1^T (\Phi_1 - \mathbf{s}_1 \mathbf{s}_1^T) \mathbf{w}_1), \end{aligned} \quad (18)$$

where $\Phi_n = \mathbf{Z}_n \mathbf{Z}_n^T$, where $\mathbf{Z}_n = \mathbf{Y}_{(n)} \left(\bigotimes_{k \neq n} \mathbf{U}^{(k)} \right)$. It follows that $\mathbf{a}^{(1)} = \mathbf{w}_1$ is the eigenvector associated with the smallest eigenvalue of $(\Phi_1 - \mathbf{s}_1 \mathbf{s}_1^T)$, and $\mathbf{U}^{(1)}$ is orthogonal complement to $\mathbf{w}_1$.

B. Estimation of Nonorthogonal Components $\mathbf{a}^{(n)}$ and $\mathbf{U}^{(n)}$

For other factors $\mathbf{a}^{(n)}$ and $\mathbf{U}^{(n)}$ with $n \neq 1$, since $\lambda = \mathbf{a}^{(n)T} \mathbf{s}_n$ the cost function (18) is simplified into

$$\begin{aligned} \min D &= \frac{1}{2} (\|\mathcal{Y}\|_F^2 - \|\mathbf{Z}_n\|_F^2 \\ &\quad + \mathbf{w}_n^T \Phi_n \mathbf{w}_n - (\xi_n \mathbf{w}_n^T \mathbf{s}_n + \rho_n \mathbf{u}_n^T \mathbf{s}_n)^2), \end{aligned}$$

where $\xi_n = \sqrt{1 - \rho_n^2}$. This yields the optimal parameter

$$\rho_n^{\text{opt}} = \frac{\mathbf{w}_n^T \mathbf{s}_n}{\sqrt{(\mathbf{w}_n^T \mathbf{s}_n)^2 + (\mathbf{u}_n^T \mathbf{s}_n)^2}}. \quad (19)$$

Algorithm 4: ASU with Orthogonality in Mode-1

Input: Data tensor $\mathcal{Y}$: $(R \times R \times \dots \times R)$, rank R

Output: A rank-1 tensor $[\lambda; \{\mathbf{a}^{(n)}\}]$ and multilinear rank- $(R-1)$ tensor $[\mathcal{G}; \{\mathbf{U}^{(n)}\}]$

begin

- 1 Initialize components $\mathbf{a}^{(n)}$ and $\mathbf{U}^{(n)}$
- repeat**
- 2 $[\mathbf{a}_1, \sigma] = \text{eigs}(\Phi_1 - \mathbf{s}_1 \mathbf{s}_1^T, 1, \text{'SA'})$
- for** $n = 2, 3, \dots, N$ **do**
- 3 $[\mathbf{w}_n, \sigma] = \text{eigs}(\Phi_n, 1, \text{'SA'})$
- 4 $\mathbf{u}_n = \frac{\mathbf{s}_n - \mathbf{w}_n (\mathbf{w}_n^T \mathbf{s}_n)}{\sqrt{\mathbf{s}_n^T \mathbf{s}_n - (\mathbf{w}_n^T \mathbf{s}_n)^2}}, \rho_n = \frac{\mathbf{w}_n^T \mathbf{s}_n}{\sqrt{(\mathbf{w}_n^T \mathbf{s}_n)^2 + (\mathbf{u}_n^T \mathbf{s}_n)^2}}$
- 5 $\mathbf{a}^{(n)} \leftarrow \mathbf{w}_n \sqrt{1 - \rho_n^2} + \mathbf{u}_n \rho_n$
- until** a stopping criterion is met
- 6 Select $\mathbf{U}^{(1)}$ as an orthogonal complement to $\mathbf{a}_1$
- for** $n = 2, \dots, N$ **do**
- 7 Select $\mathbf{U}_{2:R-1}^{(n)}$ as an orthogonal complement to $[\mathbf{w}_n, \mathbf{u}_n]$
- 8 $\lambda = \mathcal{Y} \times \{\mathbf{a}^{(n)T}\}, \mathcal{G} = \mathcal{Y} \times \{\mathbf{U}^{(n)T}\}$

The results can also be seen from the ASU algorithm after some manipulations. Once again, we come to that $\mathbf{w}_n$ and $\mathbf{u}_n$ are solution to the following optimization

$$\begin{aligned} \min \quad & \frac{1}{2} (\mathbf{w}_n^T (\Phi_n - \mathbf{s}_n \mathbf{s}_n^T) \mathbf{w}_n - (\mathbf{u}_n^T \mathbf{s}_n)^2) \\ \text{subject to} \quad & [\mathbf{w}_n, \mathbf{u}_n]^T [\mathbf{w}_n, \mathbf{u}_n] = \mathbf{I}_2. \end{aligned} \quad (20)$$

Closed-forms for $\mathbf{u}_n$ can be simplified

$$\mathbf{u}_n = \frac{\mathbf{s}_n - \mathbf{w}_n (\mathbf{w}_n^T \mathbf{s}_n)}{\sqrt{\mathbf{s}_n^T \mathbf{s}_n - (\mathbf{w}_n^T \mathbf{s}_n)^2}} \quad (21)$$

and $\mathbf{w}_n$ is the eigenvector of the smallest eigenvalue of the matrix $\Phi_n = \mathbf{Z}_n \mathbf{Z}_n^T$.

The whole ASU algorithm for decomposition with one orthogonal factor matrix is summarized in Algorithm 4. Similarly to the ASU algorithm [1], the matrices $\mathbf{U}^{(n)}$ need not be constructed except the first components $\mathbf{u}_n$ for $n = 2, \dots, N$.

V. CRAMÉR-RAO LOWER BOUND FOR TENSOR DEFLATION

In this section, we derive the Cramér-Rao Lower Bound (CRB) for the rank-1 plus multilinear rank- $(R-1, \dots, R-1)$ block tensor decomposition. We will assess the loss in accuracy of the tensor deflation by comparing its CRB with the similar bound obtained for the ordinary CP decomposition.

We consider the data model

$$\mathcal{Y} = \hat{\mathcal{Y}} + \mathcal{E} \quad (22)$$

where the deterministic part is

$$\hat{\mathcal{Y}} = \lambda [\{\mathbf{a}^{(n)}\}] + [\mathcal{G}; \mathbf{U}^{(1)}, \mathbf{U}^{(2)}, \dots, \mathbf{U}^{(N)}] \quad (23)$$

and $\mathcal{E}$ represents an error tensor which has i.i.d. normally distributed elements with zero mean and variance σ^2 . In the CP decomposition, $\mathcal{G}$ is spatially diagonal and $\mathbf{U}^{(n)}$ are arbitrary; in the tensor deflation $\mathcal{G}$ does not have any structure but $\mathbf{U}^{(n)}$ are orthogonal.

For now, the number of the estimated parameters is too high and the decomposition is not unique. Therefore, as in deriving the ASU algorithm, we perform reparameterization $\mathbf{a}^{(n)} = \mathbf{w}_n \xi_n + \mathbf{u}_n \rho_n$, where $\rho_n = [\xi_n, \rho_n]^T$ are unit-length vectors, and $[\mathbf{w}_n, \mathbf{U}_n]$ are orthonormal matrices for $n = 1, \dots, N$. The tensor $\mathcal{G}$ will be replaced by its maximum likelihood estimate

$$\hat{\mathcal{G}} = \mathcal{Y} \times \{\mathbf{U}^{(n)T}\} - \lambda \rho \mathcal{E}_1. \quad (24)$$

We follow the approach of [16] to derive the CRB from the so-called concentrated log-likelihood function. Inserting (24) into (23) we get

$$\hat{\mathcal{Y}} = \lambda [\{\mathbf{a}^{(n)}\}] - \lambda \rho [\{\mathbf{u}_n\}] + \mathcal{Y} \times \{\mathbf{I}_R - \mathbf{w}_n \mathbf{w}_n^T\}. \quad (25)$$

Considering that $\mathbf{a}^{(n)} = \mathbf{w}_n \xi_n + \mathbf{u}_n \rho_n$ and $\rho = \rho_1 \rho_2 \dots \rho_N$, the total number of parameters is reduced to $2N(R+1)+1$. The parameters are listed as elements of a vector parameter

$$\alpha = [\mathbf{w}_1^T, \mathbf{u}_1^T, \dots, \mathbf{w}_N^T, \mathbf{u}_N^T, \rho_1^T, \dots, \rho_N^T, \lambda]^T. \quad (26)$$

The parameter α obeys the constraint

$$\mathbf{f}(\alpha) = \begin{bmatrix} (\mathbf{w}_1^T \mathbf{w}_1 - 1)/2 \\ (\mathbf{u}_1^T \mathbf{u}_1 - 1)/2 \\ (\mathbf{w}_1^T \mathbf{u}_1)/\sqrt{2} \\ \vdots \\ (\mathbf{w}_N^T \mathbf{w}_N - 1)/2 \\ (\mathbf{u}_N^T \mathbf{u}_N - 1)/2 \\ (\mathbf{w}_N^T \mathbf{u}_N)/\sqrt{2} \\ (\xi_1^2 + \rho_1^2 - 1)/2 \\ \vdots \\ (\xi_N^2 + \rho_N^2 - 1)/2 \end{bmatrix} = \mathbf{0}. \quad (27)$$

The constrained Cramér-Rao bound on α [17], [18] is given by

$$\text{CRB}(\tilde{\alpha}) = \mathbf{F}_f(\alpha) (\mathbf{F}_f(\alpha)^T \mathbf{I}(\alpha) \mathbf{F}_f(\alpha))^{-1} \mathbf{F}_f(\tilde{\alpha})^T \quad (28)$$

where $\mathbf{F}_f(\alpha)$ is an orthonormal complement matrix whose vectors span the null space of the Jacobian matrix of the constraint functions with respect to α , and $\mathbf{I}(\alpha)$ is the Fisher information matrix

$$\mathbf{I}(\alpha) = \frac{1}{\sigma^2} \mathbf{H} \quad (29)$$

where $\mathbf{H}$ is the Hessian matrix of the error function $\frac{1}{2} \|\hat{\mathcal{Y}}(\alpha) - \mathcal{Y}\|_F^2$ with respect to α . Expressions of the matrix elements of $\mathbf{H}$ are given in Appendix D, whereas derivation of the orthogonal matrix $\mathbf{F}_f(\alpha)$ appearing in (28) is summarized in Appendix E.

Now, CRB on $\mathbf{a}^{(n)} = \mathbf{a}^{(n)}(\alpha)$ and the Cramér-Rao Induced Bound (CRIB) on squared angular errors (SAE)¹ of $\mathbf{a}^{(n)}$ can be computed as [7]

$$\text{CRB}(\{\mathbf{a}^{(n)}\}) = \mathbf{J}_a(\alpha) \text{CRB}(\alpha) \mathbf{J}_a(\alpha)^T, \quad (30)$$

$$\text{CRIB}(\mathbf{a}^{(n)}) = \text{tr}(\text{CRB}(\mathbf{a}^{(n)})) - \mathbf{a}^{(n)T} \text{CRB}(\mathbf{a}^{(n)}) \mathbf{a}^{(n)}, \quad (31)$$

¹Squared Angular Error between two vectors $\mathbf{a}$ and $\hat{\mathbf{a}}$ is defined as $\text{SAE}(\mathbf{a}, \hat{\mathbf{a}}) \triangleq \arccos\left(\frac{\mathbf{a}^T \hat{\mathbf{a}}}{\|\mathbf{a}\|_2 \|\hat{\mathbf{a}}\|_2}\right)^2$, and $-10 \log_{10}(\text{SAE})$ in dB.

where $\mathbf{J}_a(\alpha)$ represents the Jacobian matrix of the function $\mathbf{a}^{(n)} = \mathbf{a}^{(n)}(\alpha) = [\mathbf{w}_n, \mathbf{u}_n] \rho_n$

$$\mathbf{J}_a(\alpha) = \begin{bmatrix} \xi_1 \mathbf{I}_R & \rho_1 \mathbf{I}_R & \mathbf{w}_1 & \mathbf{u}_1 & & \\ & \ddots & \ddots & \ddots & \ddots & \\ & & \xi_N \mathbf{I}_R & \rho_N \mathbf{I}_R & \mathbf{w}_N & \mathbf{u}_N & 0 \end{bmatrix}. \quad (32)$$

The CRBs in (28), (30) and (31) derived in this paper not only serve for assessing performance of the block tensor decomposition, but also provide error bound for extraction of components of a rank-1 tensor in the CP decomposition. For the case when components $\mathbf{a}^{(n)}$ are orthogonal to $\mathbf{U}^{(n)}$, i.e., $\rho_n = 0$, the CRB should be derived with simplified constraint functions. Comparison between the CRBs derived for tensor deflation and CPD using the method developed in [7] shows that there is no significant loss in accuracy using our tensor deflation approach.

VI. SIMULATIONS

Example 1 Loss of Accuracy in Tensor Deflation Compared With CPD: This example compares CRB on extracting a rank-1 tensor in the tensor deflation with the CRB of this rank-1 tensor in CPD. The loss of accuracy in estimation of components will be shown as the difference between two similar CRBs for tensor deflation and the ordinary CPD.

The tensors in this analysis are of size $R \times R \times R$ and of rank R where $R = 10, 20$. The weight coefficients λ_r are set to 1, whereas collinearity degrees between components $\mathbf{a}_r^{(n)}$ and $\mathbf{a}_s^{(n)}$ for all $r \neq s$ are identical to a specific value c , which is varied in the range $[0, 1)$, $\mathbf{a}_r^{(n)T} \mathbf{a}_s^{(n)} = c$ and $\mathbf{a}_r^{(n)T} \mathbf{a}_r^{(n)} = 1$ for all n . Appendix F presents a method to generate the factor matrices $\mathbf{A}^{(n)}$.

The parameter ρ_n between a component $\mathbf{a}_r^{(n)}$ and the rest $(R-1)$ columns of $\mathbf{A}^{(n)}$ is a function of c and R defined as (see proof in Appendix G)

$$\rho_n = c \sqrt{\frac{R-1}{(R-2)c+1}}. \quad (33)$$

In Fig. 2(a), we compare the Cramér-Rao Induced bound (CRIB) [7] on the squared angular error (SAE) in dB, i.e., $-10 \log_{10} \text{SAE}$, of components $\mathbf{a}^{(n)}$, and the CRIB in (31). The CRBs are shown when the Gaussian noise is at SNR = 30 dB.

The two CRB curves are almost overlapped for all c , and both imply that decomposition becomes difficult, when ρ_n is close to 1, i.e., $\mathbf{a}^{(n)}$ is highly correlated to the other components, and easier when ρ_n approaches 0, i.e., components are mutually orthogonal.

The loss of accuracy in estimation of components $\mathbf{a}_r^{(n)}$ using the tensor deflation is shown in Fig. 2(b) for various ranks R . The loss is small when $c \leq 0.5$, but it tends to be high when c is closed to 1, and the rank R is small. Nevertheless, the curves reveal that the differences between two similar CRIBs were less than 0.1 dB and insignificant. Note that the loss is independent of the noise level.

Similar curves which compare the squared angular error of the ASU algorithm and Cramér-Rao bound on this error are shown in Part I [1].

Example 2 Loss of Accuracy in Tensor Deflation When Components are of Different Magnitudes: This example shows another comparison between CRIBs for tensor deflation and CPD

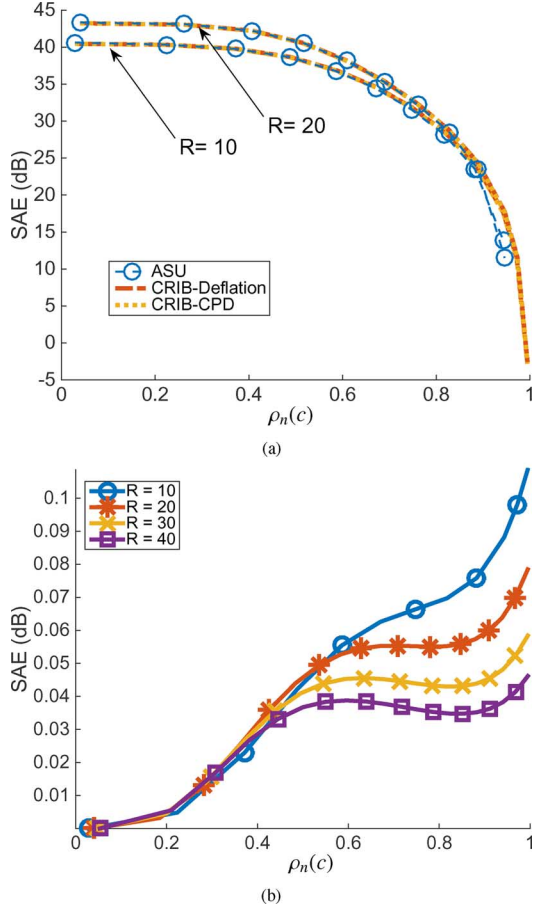


Fig. 2. Comparison of CRBs for CPD and tensor deflation on squared angular errors of components $\mathbf{a}^{(n)}$. (a) CRBs are illustrated as function of collinearity degree $c = \mathbf{a}_r^{(n)T} \mathbf{a}_s^{(n)}$ varying in the range $[0, 1)$. Tensor size $R = 10$ and 20 . (b) Loss of the squared angular error in (dB) as the difference between CRIB for tensor deflation in (31) and the similar bound for CPD in [7]. (a) SAE vs c . (b) Loss of SAE = $-10 \log_{10}(\text{CRIB-CPD [7]}/\text{CRIB in (31)})$.

[7]. Similarly to the previous example, components $\mathbf{a}_r^{(n)}$ have identical collinearity degree $c = 0.01$ and 0.6 . However, their magnitudes λ_r are different and take values as $\lambda_r = r$ for $r = 1, \dots, R$.

The two CRBs for CPD and tensor deflation are shown in Fig. 3 when $\text{SNR} = 30$ dB. Again the bounds are almost identical. Components of rank-1 tensors with large λ_r are easier to extract than ones with smaller weights. The loss of squared angular error using the tensor deflation does not exceed 0.07 dB.

Example 3 Comparison of the Initialization Algorithms: This example aims to verify efficiency of various initialization methods, including the random initialization, the SVD-based method with seeking parameters ρ_n and preselection in Algorithm 2, the tensor diagonalization tool (TEDIA) [5], and the JEVD method in Algorithm 3. For the random initialization, for each decomposition, we randomly generate 10 sets of initial points, and select the one achieving the lowest approximation error after passing through the ASU algorithm with a small number of iterations (say 10). Simulations were run on a computer consisted of Intel Xeon 2 processors clocked at 3.33 GHz, 64 GB of main memory.

The simulations are similar to those in Example 1 in Part I [1]. The tensors $\mathcal{Y}$ of size $R \times R \times R$ were randomly generated,

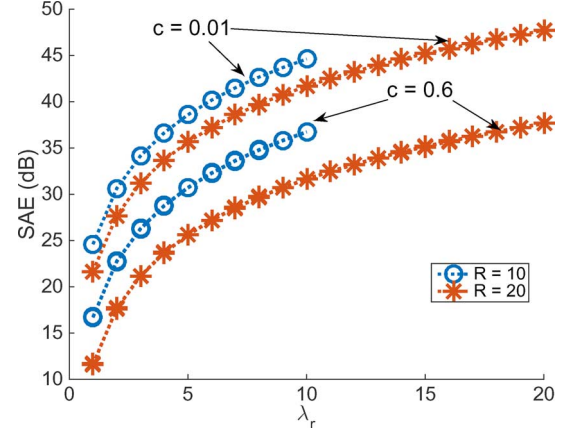


Fig. 3. Comparison of CRBs for CPD and tensor deflation on squared angular errors of components $\mathbf{a}^{(n)}$ whose associated weights are $\lambda_r = 1, 2, \dots, R$.

and corrupted by additive Gaussian noise $\mathcal{E}$ at specific signal-to-noise ratios SNR (dB)

$$\mathcal{Y} = [\mathbf{a}^{(1)}, \mathbf{a}^{(2)}, \mathbf{a}^{(3)}] + [\mathcal{G}; \mathbf{U}^{(1)}, \mathbf{U}^{(2)}, \mathbf{U}^{(3)}] + \sigma \mathcal{E}, \quad (34)$$

where $\mathbf{U}^{(n)}$ are matrices of size $R \times (R-1)$ with orthonormal columns, $\mathcal{E}$ is a zero-mean, unit-variance normal measurement error, and σ^2 denotes the noise variance. The tensors were normalized to have $\|\mathcal{G}\|_F = 1$.

For this example, we used the ASU algorithm [1]. Comparison of performance of algorithms is shown in Part I [1]. The ASU algorithm ran until differences between consecutive approximation errors were small enough, $|\varepsilon_k - \varepsilon_{k+1}| < 10^{-6} \varepsilon_k$ where $\varepsilon = \|\mathcal{Y} - \hat{\mathcal{Y}}\|_F^2$, or when the number of iterations exceeded 1000. Performances were assessed through SAE in estimation of components $\mathbf{a}^{(n)}$. There were 100 independent runs for each setting of rank and SNR, $R = 10, 20, 30$ and $\text{SNR} = 10, 20$ and 30 dB.

Fig. 4 illustrates box-plots of the squared angular errors in dB of the estimated components $\mathbf{a}^{(n)}$. Each box represents the median (central mark), 25th and 75th percentiles, and outliers. The average Cramér-Rao induced bounds on the angular errors over independent runs are shown as horizontal lines in Fig. 4.

The results indicate that when purely initialized by multiple random points, ASU failed in most test cases even for low noise levels, e.g., $\text{SNR} = 30$ dB. Similar observation was also reported for the block component decompositions using random initialization in [4]. Performances of algorithms were much improved when the components were initialized using SVD-based algorithm, JEVD and TEDIA. Except for the case using the random method, the SAEs achieved the average Cramér-Rao bound on the angular error.

Regarding running times, computing leading singular vectors was very fast, which took less than 0.02 seconds. In companion with the estimation of parameters ρ_n and pre-selection using ASU, the SVD-based initialization method in Algorithm 2 on average took 0.22, 0.33 and 0.67 seconds when $R = 10, 20$ and 30 , respectively. The JEVD method was slightly faster. For example, this initialization took an average 0.53 seconds when $R = 30$. A full graphical comparison of running times between initialization methods can be seen in Fig. 5. The random initialization was always simple and fast, but did not generate good

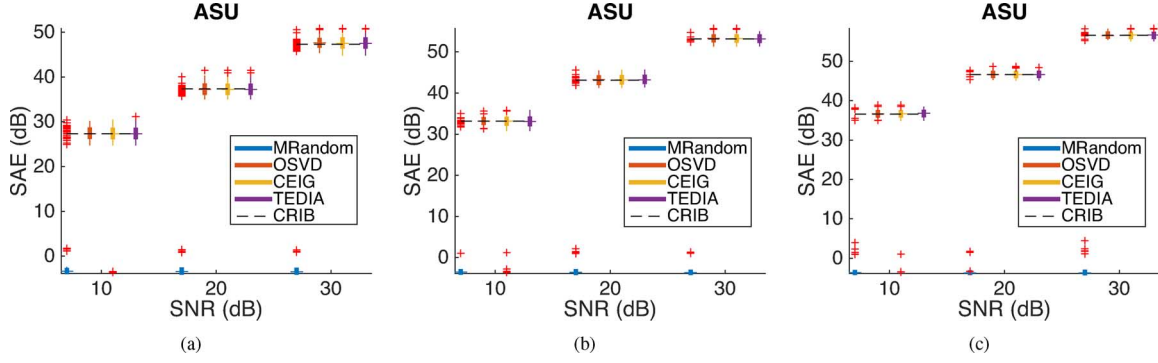


Fig. 4. Box-plots of the squared angular error achieved by algorithms using four different initialization methods: 1 – multiple random initial points, 2- leading singular vectors for $[\mathbf{w}_n, \mathbf{u}_n]$ plus estimation of parameters ρ_n , 3-JEVD, 4- TEDIA. Tensors are of size $R \times R \times R$ where $R = 10, 20$ and 30 . CRIB represents the Cramér-Rao induced bounds on the squared angular errors (SAE), which were averaged over 100 independent runs. Each box represents the median (central mark), 25th and 75th percentiles, whereas the outliers are marked with red crosses “+”. (a) $R = 10$. (b) $R = 20$. (c) $R = 30$.

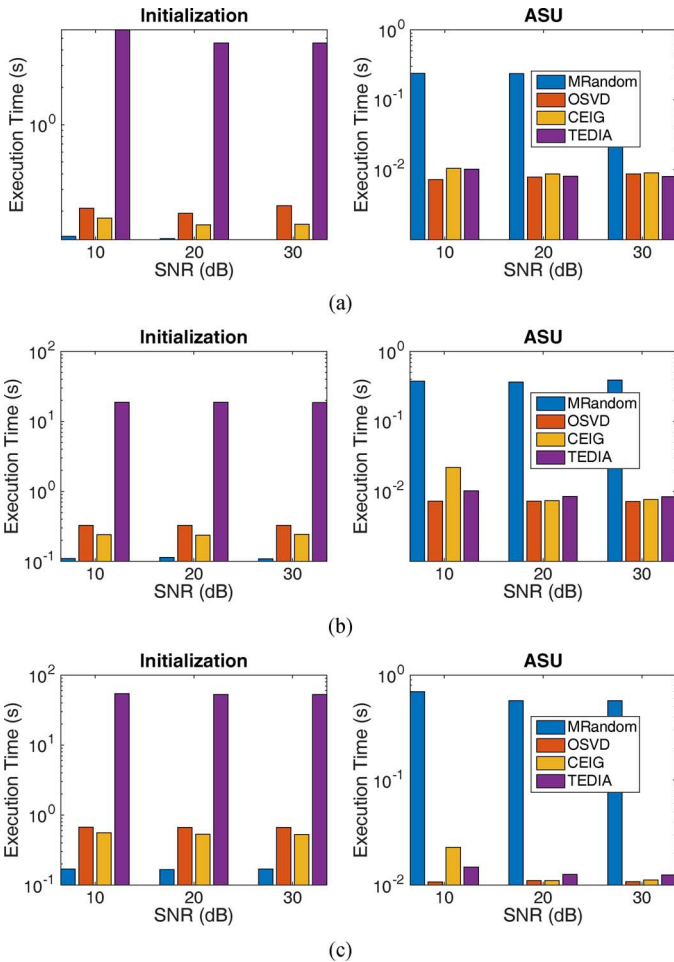


Fig. 5. Comparison of execution times of the ASU algorithm using different initialization methods for Example 3. Initialization methods are numbered as: 1- multiple random initial points, 2- leading singular vectors for $[\mathbf{w}_n, \mathbf{u}_n]$ plus estimation of parameters ρ_n , 3-JEVD, 4- TEDIA. (a) $R = 10$. (b) $R = 20$. (c) $R = 30$.

initial points. Using this initialization, ASU as well as other algorithms often did not converge to the desired solution within the prescribed maximum number of iterations (1000). However, this algorithm converged quickly after approximately 0.02 seconds with other initialization methods.

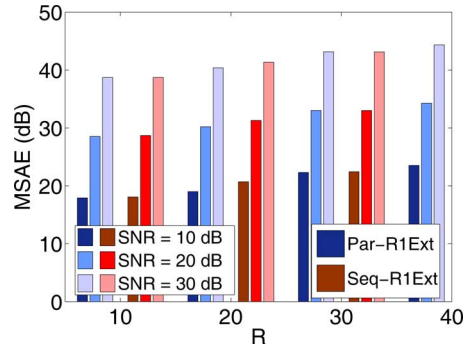


Fig. 6. Performance of the parallel rank-1 extraction (Par-R1Ext) and sequential rank-1 extraction (Seq-R1Ext) with orthogonality constraints in one mode in Example 4. Results for Par-R1Ext are shown for $R = 10, 20, 30$, and 40 , respectively, whereas results for Seq-R1Ext for $R = 10, 20$, and 30 , respectively. The color intensity indicates the SNR level.

1) *Example 4 Rank-One Extraction With Orthogonality Constraint in One Mode:* In this example, we verify the algorithm in Section IV for deflation of tensors whose one factor matrix is column-wise orthogonal. Three way tensors of size $R \times R \times R$ were generated from random factor matrices $\mathbf{A}^{(n)}$ of size $R \times R$, in which $\mathbf{A}^{(1)}$ was an orthonormal matrix. The tensors were then corrupted with additive Gaussian noise of signal-to-noise ratio $\text{SNR} = -10 \log_{10} \frac{\|\mathbf{Y}\|_F^2}{\sigma^2 \prod_{n=2}^N I_n}$, where σ^2 denotes the noise variance, and $\|\mathbf{Y}\|_F$ is the Frobenius norm of $\mathbf{Y}$. The decompositions were executed over 100 independent runs for each tensor rank $R = 10, 20, 30, 40$ and each noise level $\text{SNR} = 10, 20$ and 30 dB. Parameters of the deflation were initialized using the Direct Trilinear Decomposition (DTLD) [10] without preselection. That is, we did not choose the best rank-1 tensor among R rank-1 tensors generated by DTLD, which yields the smallest approximation error after a small number of iterations. Instead, each rank-1 tensor was assigned to the first block. Thereby, the deflation was proceeded with R different initial values. In Fig. 6, the method is labeled by “Par-R1Ext”, while Seq-R1Ext denotes the sequential rank-1 deflation. The mean SAEs (in dB) $= (-10 \log_{10}(\text{SAE}))$ on estimation of components were compared in Fig. 6. For most cases, the MSAEs (dB) achieved with Par-R1Ext were only slightly worse than those by Seq-R1Ext. The results indicate that we can accurately retrieve all components of the tensors $\mathbf{Y}$ using different initial points.

VII. CONCLUSIONS

We have proposed several initialization methods for the rank-1 tensor deflation, and verified their efficiency and performance. The SVD-based initialization method works efficiently when accompanied by the preselection. Although the computational cost of this procedure is $\mathcal{O}(R^{N+1})$, this initialization is still fast because the preselection is performed by running the ASU within a small number of iterations (e.g., 10). The JEVD initialization works well, and has relatively low complexity of order $\mathcal{O}(N^2 R^N)$ which is dominated by the tensor compressions, while the core implementation requires a cost only $\mathcal{O}(R^3)$. TEDIA is an alternative to the SVD-based and JEVD initializations.

Together with good initialization algorithms, we have shown that there is no significant loss of accuracy in estimating rank-tensors using our tensor deflation methods compared with the ordinary CP decomposition. The loss in squared angular error was at most 0.1 dB for difficult scenarios when components were highly collinear despite the noise ratio.

Finally, the initialization methods, CRLB as well as the ASU algorithm for tensor deflation are implemented in the Matlab package TENSORBOX which is available online at: <http://www.bsp.brain.riken.jp/phan/tensorbox.php>.

APPENDIX A PROOF OF LEMMA 3

In order to derive expression for λ in Lemma 1, we introduce the following results.

Lemma 3: Given a tensor $\mathcal{Y}$ which has an exact CPD or rank-1 plus multilinear rank- $(R-1, \dots, R-1)$ BTD in (1). If a rank-one tensor $\llbracket \mathbf{a}^{(1)}, \mathbf{a}^{(2)}, \dots, \mathbf{a}^{(N)} \rrbracket$ of $\mathcal{Y}$ is uniquely identified then for arbitrary orthogonal matrices $\mathbf{F}_n$ that span the row space of $\mathbf{Y}_{(n)}$

$$\left(\bigotimes_{k \neq N} \mathbf{a}^{(k)} \right)^T \mathbf{F}_n \mathbf{F}_n^T \left(\bigotimes_{k \neq N} \mathbf{a}^{(k)} \right) = 1, \quad (35)$$

where $\bigotimes_{k \neq N} \mathbf{a}^{(k)} = \mathbf{a}^{(N)} \otimes \dots \otimes \mathbf{a}^{(n+1)} \otimes \mathbf{a}^{(n-1)} \otimes \dots \otimes \mathbf{a}^{(1)}$.

Proof: For simplicity, the proof is given for order-3 tensors in block term decomposition. Mode-1 matricization for the exact decomposition for order-3 tensors in (1) shows that

$$\mathbf{Y}_{(1)} = [\mathbf{a}^{(1)}, \mathbf{U}^{(1)}] \begin{bmatrix} \lambda (\mathbf{a}^{(3)} \otimes \mathbf{a}^{(2)})^T \\ \mathbf{G}_{(1)} (\mathbf{U}^{(3)} \otimes \mathbf{U}^{(2)})^T \end{bmatrix}, \quad (36)$$

where $[\mathbf{a}^{(1)}, \mathbf{U}^{(1)}]$ is a full-rank matrix. Since we have assumed that $\mathbf{a}^{(1)}, \mathbf{a}^{(2)}$ and $\mathbf{a}^{(3)}$ are uniquely identified, and $\mathbf{F}_1$ spans the row space of $\mathbf{Y}_{(1)}$, there exists a vector $\mathbf{v}$ such that

$$\mathbf{F}_1 \mathbf{v} = \mathbf{a}^{(3)} \otimes \mathbf{a}^{(2)}. \quad (37)$$

Since $\mathbf{a}^{(2)T} \mathbf{a}^{(2)} = \mathbf{a}^{(3)T} \mathbf{a}^{(3)} = 1$, we have $\mathbf{v} = \mathbf{F}_1^T (\mathbf{a}^{(3)} \otimes \mathbf{a}^{(2)})$ and $\mathbf{v}^T \mathbf{v} = 1$. Thereby,

$$(\mathbf{a}^{(3)} \otimes \mathbf{a}^{(2)})^T \mathbf{F}_1 \mathbf{F}_1^T (\mathbf{a}^{(3)} \otimes \mathbf{a}^{(2)}) = 1. \quad (38)$$

Proof: Since $\mathbf{Y}_{(n)}^T (\mathbf{Y}_{(n)} \mathbf{Y}_{(n)}^T)^{-1} \mathbf{Y}_{(n)}$ is the projection matrix onto the subspace spanned by the rows of $\mathbf{Y}_{(n)}$, we can replace $\mathbf{F}_n \mathbf{F}_n^T$ in (35) by $\mathbf{Y}_{(n)}^T (\mathbf{Y}_{(n)} \mathbf{Y}_{(n)}^T)^{-1} \mathbf{Y}_{(n)}$, and rewrite the expression using the new notion $\mathbf{s}_n$ to obtain (39). ■

Proof of Lemma 1: When the rank-1 tensor $\llbracket \{\mathbf{a}^{(n)}\} \rrbracket$ is given, $\mathbf{U}^{(n)}$ are orthonormal basis of column spaces of mode- n matricizations of the tensor $\mathcal{Y} - \llbracket \lambda; \{\mathbf{a}^{(n)}\} \rrbracket$ or column spaces of matrices

$$\begin{aligned} \Psi_n &= \begin{pmatrix} \mathbf{Y}_{(n)} - \lambda \mathbf{a}^{(n)} \left(\bigotimes_{k \neq n} \mathbf{a}^{(k)} \right)^T \\ \left(\mathbf{Y}_{(n)} - \lambda \mathbf{a}^{(n)} \left(\bigotimes_{k \neq n} \mathbf{a}^{(k)} \right)^T \right)^T \end{pmatrix} \\ &= \mathbf{Y}_{(n)} \mathbf{Y}_{(n)}^T + \mathbf{T}_n \begin{bmatrix} \lambda^2 & -\lambda \\ -\lambda & 0 \end{bmatrix} \mathbf{T}_n^T, \end{aligned} \quad (40)$$

with $\mathbf{T}_n = [\mathbf{a}^{(n)}, \mathbf{s}_n]$. For the exact case, λ is selected such that ranks of the matrices $\left(\mathbf{Y}_{(n)} - \lambda \mathbf{a}^{(n)} \left(\bigotimes_{k \neq n} \mathbf{a}^{(k)} \right)^T \right)$ or Ψ_n are reduced by 1. This implies that λ is a solution to equations

$$\det(\Psi_n) = 0, \quad \text{for } n = 1, \dots, N,$$

where $\det(\Psi_n)$ are quadratic polynomials given by

$$\begin{aligned} \det(\Psi_n) &= \beta_n \det \left(\mathbf{I}_2 + \mathbf{K}_n \begin{bmatrix} \lambda^2 & -\lambda \\ -\lambda & 0 \end{bmatrix} \right) \\ &= \beta_n (\alpha_{n2} \lambda^2 + \alpha_{n1} \lambda + 1), \end{aligned} \quad (41)$$

where $\mathbf{K}_n = \mathbf{T}_n^T (\mathbf{Y}_{(n)} \mathbf{Y}_{(n)}^T)^{-1} \mathbf{T}_n$ are matrices of size 2×2 , $\beta_n = \det(\mathbf{Y}_{(n)} \mathbf{Y}_{(n)}^T)$, $\alpha_{n2} = \mathbf{K}_n(1, 1) - \det(\mathbf{K}_n)$ and $\alpha_{n1} = -2 \mathbf{K}_n(1, 2)$.

According to (39), we have

$$\mathbf{K}_n(2, 2) = \mathbf{s}_n^T \mathbf{C}_n^{-1} \mathbf{s}_n = 1.$$

Thereby, $\alpha_{n2} = \mathbf{K}_n(1, 2)^2$, or

$$\det(\Psi_n) = \beta_n (\mathbf{K}_n(1, 2) \lambda - 1)^2, \quad (42)$$

giving us closed-form expression in (5). ■

APPENDIX B PROOF OF LEMMA 2

Proof: Multiplying along mode-1 both sides of $\mathcal{Y} = \llbracket \lambda; \{\mathbf{a}^{(n)}\} \rrbracket + \llbracket \mathcal{G}; \{\mathbf{U}^{(n)}\} \rrbracket$ with

$$\tilde{\mathbf{w}}_1 = \frac{\mathbf{w}_1}{\sqrt{1 - \rho_1^2}} = \frac{\mathbf{a}^{(1)} - \rho_1 \mathbf{u}_1}{1 - \rho_1^2}, \quad (43)$$

we get

$$\mathcal{Y} \times_1 \tilde{\mathbf{w}}_1^T = \lambda \llbracket \mathbf{a}^{(2)}, \dots, \mathbf{a}^{(N)} \rrbracket, \quad (44)$$

since $\tilde{\mathbf{w}}_1^T \mathbf{a}^{(1)} = 1$. Consider the Frobenius norm of the difference between the two sides of (44) which achieves zero for the exact decomposition

$$\begin{aligned} D(\tilde{\mathbf{w}}_1) &= \|\mathcal{Y} \times_1 \tilde{\mathbf{w}}_1^T - \llbracket \mathbf{a}^{(2)}, \dots, \mathbf{a}^{(N)} \rrbracket\|_F^2 \\ &= \tilde{\mathbf{w}}_1^T \mathbf{C}_1 \tilde{\mathbf{w}}_1 + \lambda^2 - 2\lambda \tilde{\mathbf{w}}_1^T \mathbf{s}_1. \end{aligned}$$

Corollary 1: Assuming that $\mathbf{Y}_{(n)}^T$ has a full column rank, then

$$\mathbf{s}_n^T \mathbf{C}_n^{-1} \mathbf{s}_n = 1. \quad (39)$$

It is straightforward to see that $\lambda \mathbf{C}_1^{-1} \mathbf{s}_1$ is root of the above function since $\lambda \mathbf{a}^{(1)T} \mathbf{C}_1^{-1} \mathbf{s}_1 = 1$ and

$$D(\lambda \mathbf{C}_1^{-1} \mathbf{s}_1) = \lambda^2 \mathbf{s}_1^T \mathbf{C}_1^{-1} \mathbf{C}_1 \mathbf{C}_1^{-1} \mathbf{s}_1 + \lambda^2 - 2 \lambda^2 \mathbf{s}_1^T \mathbf{C}_1^{-1} \mathbf{s}_1 = \lambda^2 + \lambda^2 - 2\lambda^2 = 0.$$

From (43), we obtain $\mathbf{w}_1 = \sqrt{1 - \rho_1^2} \lambda \mathbf{C}_1^{-1} \mathbf{s}_1$. Since $\mathbf{w}_1$ is a unit-length vector, this vector takes the form

$$\mathbf{w}_1 = \frac{\mathbf{C}_1^{-1} \mathbf{s}_1}{\sqrt{\mathbf{s}_1^T \mathbf{C}_1^{-2} \mathbf{s}_1}} \quad (45)$$

and ρ_1 can be derived from

$$(1 - \rho_1^2) \lambda^2 = \frac{1}{\mathbf{s}_1^T \mathbf{C}_1^{-2} \mathbf{s}_1}, \quad (46)$$

which yields the optimal ρ_1 as in (7). Similarly, the expressions (6) and (7) hold for other $n = 2, \dots, N$. When given $\mathbf{w}_n$ and ρ_n , the vectors $\mathbf{u}_n$ are computed from (2). ■

APPENDIX C

PROOF OF THEOREM 1

Proof: Since $\mathbf{X}_1 \mathbf{X}_2^{-1} = \mathbf{X}_2 (\mathbf{X}_2^{-1} \mathbf{X}_1) \mathbf{X}_2^{-1}$, $\mathbf{X}_1 \mathbf{X}_2^{-1}$ and $\mathbf{X}_2^{-1} \mathbf{X}_1$ are similar, and have the same eigenvalues λ

$$\begin{aligned} \mathbf{X}_1 \mathbf{X}_2^{-1} \mathbf{a} &= \mathbf{a} \lambda \\ \mathbf{b}^T \mathbf{X}_2^{-1} \mathbf{X}_1 &= \mathbf{b}^T \lambda. \end{aligned}$$

Hence, $\mathbf{a}$ and $\tilde{\mathbf{b}} = \frac{(\mathbf{X}_2^{-1})^T \mathbf{b}}{\mathbf{b}^T \mathbf{X}_2^{-1} \mathbf{a}}$ with $\tilde{\mathbf{b}}^T \mathbf{a} = 1$, are left and right eigenvectors of the matrix $\mathbf{X}_1 \mathbf{X}_2^{-1}$. Let $\mathbf{B}$ be an orthogonal complement to $\tilde{\mathbf{b}}$, we define a matrix $\mathbf{Q}$ of size $R \times R$

$$\mathbf{Q} = [\mathbf{a}, \mathbf{B}],$$

whose inverse is given by

$$\mathbf{Q}^{-1} = \begin{bmatrix} \tilde{\mathbf{b}}^T \\ \mathbf{B}^T (\mathbf{I}_R - \mathbf{a} \tilde{\mathbf{b}}^T) \end{bmatrix}.$$

It is obvious that

$$\mathbf{Q}^{-1} \mathbf{X}_1 \mathbf{X}_2^{-1} \mathbf{Q} = \begin{bmatrix} \lambda & \\ & \mathbf{C} \end{bmatrix},$$

where $\mathbf{C} = \mathbf{B}^T (\mathbf{I}_R - \mathbf{a} \tilde{\mathbf{b}}^T) \mathbf{X}_1 \mathbf{X}_2^{-1} \mathbf{B}$ is of size $(R-1) \times (R-1)$, and has rank- $(R-1)$. It follows that

$$\begin{aligned} \mathbf{X}_1 &= \lambda \mathbf{a} \tilde{\mathbf{b}}^T \mathbf{X}_2 + \mathbf{B} \mathbf{C} \mathbf{B}^T (\mathbf{I}_R - \mathbf{a} \tilde{\mathbf{b}}^T) \mathbf{X}_2 \\ &= \frac{1}{\mathbf{b}^T \mathbf{X}_1^{-1} \mathbf{a}} \mathbf{a} \mathbf{b}^T + \mathbf{B} \mathbf{C} \mathbf{B}^T (\mathbf{I}_R - \mathbf{a} \tilde{\mathbf{b}}^T) \mathbf{X}_2, \end{aligned} \quad (47)$$

which indicates that $\mathbf{X}_1$ can be split into a rank-1 matrix composed by generalized eigenvectors $\frac{1}{\mathbf{b}^T \mathbf{X}_1^{-1} \mathbf{a}} \mathbf{a} \mathbf{b}^T$, and a rank- $(R-1)$ term.

Next, $\mathbf{U}$ and $\mathbf{V}$ can be chosen arbitrarily as orthonormal basis vectors of the column and row space of the second term in (47) or of the matrix $(\mathbf{X}_1 - \lambda \mathbf{X}_2)$. The matrices are then normalized to be orthogonal to $\mathbf{a}$ and $\mathbf{b}$, respectively, i.e., $\mathbf{U}^T \mathbf{a} = \rho_1 [1, 0, \dots, 0]^T$, $\mathbf{V}^T \mathbf{b} = \rho_2 [1, 0, \dots, 0]^T$ where $0 \leq \rho_1 < 1$ and $0 \leq \rho_2 < 1$. ■

APPENDIX D

THE FIM $\mathbf{I}(\hat{\alpha})$ IN (29)

Consider the error function

$$D(\mathbf{y}|\boldsymbol{\alpha}) = \frac{1}{2} \|\mathbf{y} - \lambda[\{\mathbf{a}^{(n)}\}] + \lambda \rho[\{\mathbf{u}_n\}] - \mathbf{y} \times \{\mathbf{I}_R - \mathbf{w}_n \mathbf{w}_n^T\}\|_F^2. \quad (48)$$

The Hessian of the error function with respect to $\boldsymbol{\alpha}$ is evaluated at $\mathbf{y} = \hat{\mathbf{y}}$, i.e., when the tensor-valued function $\mathcal{E}$ inside the Frobenius norm in (48) goes to zero. For such case, the Hessian is given by

$$\mathbf{H} = \mathbf{J}_D(\boldsymbol{\alpha})^T \mathbf{J}_D(\boldsymbol{\alpha}) \quad (49)$$

where $\mathbf{J}_D(\boldsymbol{\alpha})$ is the Jacobian of the tensor-valued function $\mathcal{E}$ with respect to $\boldsymbol{\alpha}$ and evaluated when $\mathbf{y} = \hat{\mathbf{y}}$.

The Hessian $\mathbf{H}$ of size $(2N(R+1)+1) \times (2N(R+1)+1)$ is given in block form as

$$\mathbf{H} = \begin{bmatrix} \mathbf{H}_{\mathbf{V}, \mathbf{V}} & \mathbf{H}_{\mathbf{V}, \rho} & \mathbf{h}_{\mathbf{V}, \lambda} \\ \mathbf{H}_{\mathbf{V}, \rho}^T & \mathbf{H}_{\rho, \rho} & \mathbf{h}_{\rho, \lambda} \\ \mathbf{h}_{\mathbf{V}, \lambda}^T & \mathbf{h}_{\rho, \lambda}^T & 1 - \rho^2 \end{bmatrix} \quad (50)$$

where $\mathbf{H}_{\mathbf{V}, \mathbf{V}} = [\mathbf{H}_{\mathbf{V}_n, \mathbf{V}_m}]_{n,m=1}^N$ is an $N \times N$ partitioned matrix of submatrices $\mathbf{H}_{\mathbf{V}_n, \mathbf{V}_m}$ of size $2R \times 2R$, $\mathbf{H}_{\mathbf{V}, \rho} = [\mathbf{H}_{\mathbf{V}_n, \rho_m}]_{n,m=1}^N$ is a partitioned matrix of $\mathbf{H}_{\mathbf{V}_n, \rho_m}$ of size $2R \times 2$, $\mathbf{H}_{\rho, \rho} = [\mathbf{H}_{\rho_n, \rho_m}]$ are of size 2×2 , $\mathbf{h}_{\mathbf{V}, \lambda} = [\mathbf{h}_{\mathbf{w}_1, \lambda}^T, \mathbf{h}_{\mathbf{u}_1, \lambda}^T, \dots, \mathbf{h}_{\mathbf{w}_N, \lambda}^T, \mathbf{h}_{\mathbf{u}_N, \lambda}^T]^T$ is of length $2NR$, and $\mathbf{h}_{\rho, \lambda} = [\mathbf{h}_{\rho_1, \lambda}^T, \dots, \mathbf{h}_{\rho_N, \lambda}^T]^T$

$$\mathbf{H}_{\mathbf{V}_n, \mathbf{V}_m} = \begin{bmatrix} \mathbf{H}_{\mathbf{w}_n, \mathbf{w}_m} & \mathbf{H}_{\mathbf{w}_n, \mathbf{u}_m} \\ \mathbf{H}_{\mathbf{u}_n, \mathbf{w}_m} & \mathbf{H}_{\mathbf{u}_n, \mathbf{u}_m} \end{bmatrix}, \quad \mathbf{H}_{\mathbf{V}_n, \rho_m} = \begin{bmatrix} \mathbf{H}_{\mathbf{w}_n, \rho_m} \\ \mathbf{H}_{\mathbf{u}_n, \rho_m} \end{bmatrix} \quad (51)$$

where $\mathbf{V}_n = [\mathbf{w}_n, \mathbf{u}_n]$.

The element sub-matrices to construct the Hessian $\mathbf{H}$, i.e., $\mathbf{H}_{\mathbf{w}_n, \mathbf{w}_m}$, $\mathbf{H}_{\mathbf{w}_n, \mathbf{u}_m}$, $\mathbf{H}_{\mathbf{u}_n, \mathbf{u}_m}$, $\dots$, are given by

$$\begin{aligned} \mathbf{H}_{\mathbf{w}_n, \mathbf{w}_n} &= \Phi_n + \lambda^2 \xi_n^2 (1 - \rho_{-n}^2) \mathbf{I}_R \\ \mathbf{H}_{\mathbf{u}_n, \mathbf{u}_n} &= \lambda^2 \rho_n^2 (1 - \rho_{-n}^2) \mathbf{I}_R \\ \mathbf{H}_{\mathbf{w}_n, \mathbf{u}_n} &= \lambda^2 \xi_n \rho_n (1 - \rho_{-n}^2) \mathbf{I}_R \\ \mathbf{H}_{\mathbf{u}_n, \mathbf{u}_m} &= \lambda^2 \rho_n \rho_m \mathbf{a}^{(n)} \mathbf{a}^{(m)T} + \lambda^2 \rho^2 \mathbf{u}_n \mathbf{u}_m^T \\ &\quad - \lambda^2 \rho \rho_{-n} \mathbf{a}^{(n)} \mathbf{u}_m^T - \lambda^2 \rho \rho_{-m} \mathbf{u}_n \mathbf{a}^{(m)T} \\ \mathbf{H}_{\mathbf{w}_n, \mathbf{w}_m} &= -2 \lambda \xi_n \xi_m \rho_{-[n,m]} \mathbf{Z}_{nm} + \lambda^2 \xi_n \xi_m \mathbf{a}^{(n)} \mathbf{a}^{(m)T} \\ &\quad + \lambda^2 \xi_n \xi_m \rho_{-[n,m]}^2 (2 \xi_n \xi_m \mathbf{w}_n \mathbf{w}_m^T - \rho_n \rho_m \mathbf{u}_n \mathbf{u}_m^T) \\ \mathbf{H}_{\mathbf{w}_n, \mathbf{u}_m} &= -\lambda \xi_n \rho_{-[n,m]} \mathbf{Z}_{n,m} \\ &\quad + \lambda^2 \xi_n \rho_m (1 - \rho_{-[n,m]}^2) \mathbf{a}^{(n)} \mathbf{a}^{(m)T} \\ &\quad + \lambda^2 \xi_n \xi_m \rho_m \rho_{-[n,m]}^2 (2 \mathbf{a}^{(n)} - \rho_n \mathbf{u}_n) \mathbf{w}_m^T \\ \mathbf{H}_{\mathbf{u}_n, \rho_n} &= \lambda^2 \rho_n (1 - \rho_{-n}^2) [\mathbf{w}_n, \mathbf{u}_n] \\ \mathbf{H}_{\mathbf{w}_n, \rho_n} &= \lambda^2 \xi_n (1 - \rho_{-n}^2) [\mathbf{w}_n, \mathbf{u}_n] - \lambda \rho_{-n} \mathbf{z}_n [1, 0] \\ \mathbf{H}_{\mathbf{u}_n, \rho_m} &= \lambda^2 \rho_n \mathbf{a}^{(n)} [\xi_m, \rho_m (1 - \rho_{-[n,m]}^2)] \\ \mathbf{H}_{\mathbf{w}_n, \rho_m} &= \lambda^2 \xi_n \mathbf{a}^{(n)} [\xi_m, \rho_m (1 - \rho_{-[n,m]}^2)] \\ &\quad - \lambda \xi_n \rho_{-[n,m]} \mathbf{z}_n [0, 1] \\ \mathbf{H}_{\rho_n, \rho_n} &= \lambda^2 \text{diag}(1, 1 - \rho_{-n}^2) \\ \mathbf{H}_{\rho_n, \rho_m} &= \lambda^2 (\rho_n \rho_m^T - \text{diag}(0, \rho_{-n} \rho_{-m})) \\ \mathbf{h}_{\mathbf{u}_n, \lambda} &= \lambda \rho_n (1 - \rho_{-n}^2) \mathbf{a}^{(n)} \\ \mathbf{h}_{\mathbf{w}_n, \lambda} &= \lambda \xi_n (1 - \rho_{-n}^2) \mathbf{a}^{(n)} - \xi_n \rho_{-n} \mathbf{z}_n \\ \mathbf{h}_{\rho_n, \lambda} &= \lambda [\xi_n, \rho_n (1 - \rho_{-n}^2)]^T \end{aligned}$$

where $\rho_{-n} = \frac{\rho}{\rho_n}$, $\rho_{-[n,m]} = \frac{\rho}{\rho_n \rho_m}$ and

$$\begin{aligned} \mathbf{z}_n &= \hat{\mathbf{Y}} \times_{k \neq n} \{\mathbf{u}_k^T\} = \lambda \rho_{-n} \mathbf{a}^{(n)} + \mathbf{U}^{(n)} \mathbf{g}_1^{(n)} \\ \mathbf{Z}_{n,m} &= \hat{\mathbf{Y}} \times_{\substack{k \neq n, \\ k \neq m}} \{\mathbf{u}_k^T\} \\ &= \lambda \rho_{-[n,m]} \mathbf{a}^{(n)} \mathbf{a}^{(m)T} + \mathbf{U}^{(n)} \mathbf{G}_1^{(n,m)} \mathbf{U}^{(m)T} \\ \Phi_n &= \hat{\mathbf{Y}}_{(n)} \left(\bigotimes_{k \neq n} (\mathbf{I} - \mathbf{w}_k \mathbf{w}_k^T) \right) \hat{\mathbf{Y}}_{(n)}^T \\ &= \lambda^2 \rho_{-n}^2 \mathbf{a}^{(n)} \mathbf{a}^{(n)T} + \mathbf{U}^{(n)} (\mathbf{G}_{(n)} \mathbf{G}_{(n)}^T) \mathbf{U}^{(n)T} \\ &\quad + \lambda \rho_{-n} \mathbf{U}^{(n)} \mathbf{g}_1^{(n)} \mathbf{a}^{(n)T} + \lambda \rho_{-n} \mathbf{a}^{(n)} \mathbf{g}_1^{(n)T} \mathbf{U}^{(n)T} \end{aligned}$$

where $\mathbf{g}_1^{(n)}$ is the first column of $\mathbf{G}_{(n)}$, i.e., the fiber of the tensor $\mathcal{G}$ along mode- n with all other indices $i_k = 1$, $k \neq n$, $\mathbf{G}_1^{(n,m)}$ is the slice of $\mathcal{G}$ along modes n and m with all other indices $i_k = 1$, $k \neq n, m$.

APPENDIX E

JACOBIAN OF THE FUNCTIONAL EQUALITY CONSTRAINTS AND ITS ORTHONORMAL COMPLEMENT

The Jacobian matrix of the constraint functions $\mathbf{f}(\boldsymbol{\alpha})$ in (27) is given in block form as

$$\mathbf{J}_f(\boldsymbol{\alpha}) = \text{bdiag}(\mathbf{J}_1, \dots, \mathbf{J}_N, \boldsymbol{\rho}_1^T, \dots, \boldsymbol{\rho}_N^T, 0) \quad (52)$$

where bdiag constructs a block diagonal matrix, and

$$\mathbf{J}_n = \begin{bmatrix} \mathbf{w}_n^T & \mathbf{u}_n^T \\ \frac{1}{\sqrt{2}} \mathbf{u}_n^T & \frac{1}{\sqrt{2}} \mathbf{w}_n^T \end{bmatrix} \in \mathbb{R}^{3 \times 2R}, \quad n = 1, \dots, N. \quad (53)$$

We define a matrix $\mathbf{F}_n$ of size $2R \times (2R - 3)$ as

$$\mathbf{F}_n = \begin{bmatrix} \tilde{\mathbf{U}}^{(n)} & \frac{1}{\sqrt{2}} \mathbf{u}_n \\ \tilde{\mathbf{U}}^{(n)} & \frac{-1}{\sqrt{2}} \mathbf{w}_n \end{bmatrix} \quad (54)$$

where $\tilde{\mathbf{U}}^{(n)} = \mathbf{U}_{2:R-1}^{(n)}$ takes the last $(R - 2)$ columns of $\mathbf{U}^{(n)}$. It is straightforward to see that $[\mathbf{J}_n^T, \mathbf{F}_n]$ forms an orthonormal matrix of size $2R \times 2R$

$$[\mathbf{J}_n^T, \mathbf{F}_n]^T [\mathbf{J}_n^T, \mathbf{F}_n] = \mathbf{I}_{2R}.$$

Therefore, we can chose the orthogonal basis matrix $\mathbf{F}_f(\boldsymbol{\alpha})$ of size $(2N(R + 1) + 1) \times (2N(R - 1) + 1)$ as follows

$$\mathbf{F}_f(\boldsymbol{\alpha}) = \text{bdiag}(\mathbf{F}_1, \dots, \mathbf{F}_N, \begin{bmatrix} -\rho_1 \\ \xi_1 \end{bmatrix}, \dots, \begin{bmatrix} -\rho_N \\ \xi_N \end{bmatrix}, 1). \quad (55)$$

APPENDIX F

GENERATION OF MATRIX WITH SPECIFIC COLLINEARITY DEGREE

This appendix presents a method to generate a random matrix $\mathbf{A}$ of size $I \times R$ with $I \geq R$ whose collinearity degree between columns are identical to a specific value c , that is

$$\mathbf{A}^T \mathbf{A} = (1 - c) \mathbf{I}_R + c \mathbf{1}_{R \times R}. \quad (56)$$

We define a matrix $\mathbf{Q}$ of size $R \times R$ as

$$\mathbf{Q} = \begin{bmatrix} 1 & c_1 & c_1 & \dots & c_1 \\ & b_1 & c_2 & \dots & c_2 \\ & & \ddots & \ddots & \vdots \\ & & & b_{R-2} & c_{R-1} \\ & & & & b_{R-1} \end{bmatrix} \quad (57)$$

where $c_1 = c$, $c_r = \frac{1}{b_{r-1}} \left(c - \sum_{i=1}^{r-1} c_i^2 \right)$, $b_r = \sqrt{1 - \sum_{i=1}^r c_i^2}$ for $r = 1, 2, \dots, R - 1$. It is obvious to see that $\boldsymbol{\Omega} = \mathbf{Q}^T \mathbf{Q} = (1 - c) \mathbf{I}_R + c \mathbf{1}_{R \times R}$, because

$$\begin{aligned} \omega_{r,r} &= \mathbf{q}_r^T \mathbf{q}_r = \sum_{i=1}^{r-1} c_i^2 + b_{r-1}^2 = 1 \\ \omega_{r,s} &= \mathbf{q}_r^T \mathbf{q}_s = \sum_{i=1}^{r-1} c_i^2 + b_{r-1} c_r = c, \quad 1 \leq r < s \leq R. \end{aligned}$$

Now with an arbitrary orthogonal matrix $\mathbf{U}$ of size $I \times R$, $\mathbf{U}^T \mathbf{U} = \mathbf{I}_R$, the matrix $\mathbf{A} = \mathbf{U} \mathbf{Q}$ satisfies the condition in (56).

APPENDIX G

RELATION OF COLLINEARITY DEGREE c AND PARAMETER ρ_n IN EXAMPLE 1

We consider a matrix $\mathbf{A} = [\mathbf{a}, \mathbf{B}]$ of size $R \times R$ with $\mathbf{a}_r^T \mathbf{a}_r = 1$ and $\mathbf{a}_r^T \mathbf{a}_s = c$ for $r \neq s$. $\mathbf{a}$ is the first column, and $\mathbf{B}$ comprises the last $(R - 1)$ columns of $\mathbf{A}$. Since

$$\begin{aligned} \mathbf{B}^T \mathbf{B} &= (1 - c) \mathbf{I}_{R-1} + c \mathbf{1}_{R-1} \mathbf{1}_{R-1}^T \\ &= \frac{(R - 2)c + 1}{R - 1} \mathbf{1} \mathbf{1}^T + (1 - c) \left(\mathbf{I}_{R-1} - \frac{1}{R - 1} \mathbf{1} \mathbf{1}^T \right), \end{aligned}$$

the first eigenpair $(\mathbf{v}_1, \sigma_1^2)$ of the matrix $\mathbf{B}^T \mathbf{B} = \mathbf{V} \boldsymbol{\Sigma}^2 \mathbf{V}^T$, $\boldsymbol{\Sigma} = \text{diag}(\sigma_1, \dots, \sigma_{R-1})$ is given by

$$\mathbf{v}_1 = \frac{1}{\sqrt{R - 1}} \mathbf{1}_{R-1}, \quad \sigma_1 = \sqrt{(R - 2)c + 1} \quad (58)$$

the rest $(R - 2)$ eigenvalues of $\mathbf{B}^T \mathbf{B}$ are identical

$$\sigma_r = \sqrt{1 - c}, \quad r = 2, \dots, R - 1, \quad (59)$$

and $\mathbf{V}_{2:R-1} = [\mathbf{v}_2, \dots, \mathbf{v}_{R-1}]$ can take arbitrary orthonormal basis of orthogonal complement to $\mathbf{v}_1$, i.e., $\mathbf{V}_{2:R-1}^T \mathbf{1}_{R-1} = \mathbf{0}$. This gives us singular value decomposition of the matrix $\mathbf{B}$

$$\mathbf{B} = \mathbf{U} \boldsymbol{\Sigma} \mathbf{V}^T \quad (60)$$

where $\mathbf{U}$ of size $R \times (R - 1)$. We denote by $\mathbf{w}$ a unit-length vector of orthogonal complement to $\mathbf{U}$, i.e., $[\mathbf{U}, \mathbf{w}]$ forms an orthonormal matrix of size $R \times R$. The vector $\mathbf{a}$ can be always expressed as linear combination of $[\mathbf{U}, \mathbf{w}]$

$$\mathbf{a} = \mathbf{U} \boldsymbol{\beta} + \xi \mathbf{w} \quad (61)$$

where $\boldsymbol{\beta}$ is a vector of length $R - 1$. Since

$$c \mathbf{1}_{R-1} = \mathbf{B}^T \mathbf{a} = \mathbf{V} \boldsymbol{\Sigma} \boldsymbol{\beta},$$

we obtain

$$\begin{aligned}\beta &= c \Sigma^{-1} \mathbf{V}^T \mathbf{1}_{R-1} = c \Sigma^{-1} \left[\sqrt{R-1}, \mathbf{0}_{R-2}^T \right]^T \\ &= c \sqrt{\frac{R-1}{(R-2)c+1}} \mathbf{e}_1,\end{aligned}\quad (62)$$

indicating that

$$\mathbf{a} = c \sqrt{\frac{R-1}{(R-2)c+1}} \mathbf{u}_1 + \xi \mathbf{w}.\quad (63)$$

Therefore, for Example 1, the parameter ρ_n defined for $\mathbf{a}_r^{(n)}$ and the rest $(R-1)$ columns of $\mathbf{A}^{(n)}$ is given by

$$\rho_n = c \sqrt{\frac{R-1}{(R-2)c+1}}.\quad (64)$$

ACKNOWLEDGMENT

The authors would like to thank the associate editor and the referees for their time and efforts to review this paper, and for very constructive comments which helped to improve the quality and presentation of the paper.

REFERENCES

- [1] A.-H. Phan, P. Tichavský, and A. Cichocki, "Tensor deflation for CANDECOMP/PARAFAC—Part I: Alternating subspace update algorithm," *IEEE Trans. Signal Process.*, vol. 63, no. 22, pp. 5924–5938, 2015.
- [2] A.-H. Phan, P. Tichavský, and A. Cichocki, "Deflation method for CANDECOMP/PARAFAC tensor decomposition," in *Proc. IEEE Int. Conf. Acoust., Speech, Signal Process. (ICASSP)*, May 2014, pp. 6736–6740.
- [3] A.-H. Phan, P. Tichavský, and A. Cichocki, "Tensor deflation for CANDECOMP/PARAFAC. Part 3: Rank splitting," *ArXiv e-prints* 2015 [Online]. Available: <http://arxiv.org/abs/1506.04971>
- [4] L. De Lathauwer and D. Nion, "Decompositions of a higher-order tensor in block terms—Part III: Alternating least squares algorithms," *SIAM J. Matrix Anal. Appl.*, vol. 30, no. 3, pp. 1067–1083, 2008.
- [5] P. Tichavský, A.-H. Phan, and A. Cichocki, "Tensor diagonalization—A new tool for PARAFAC and block-term decomposition," 2014, arXiv:1402.1673 [Online]. Available: <http://arxiv.org/abs/1402.1673v1>, to be published
- [6] W. P. Krijnen, T. K. Dijkstra, and A. Stegeman, "On the non-existence of optimal solutions and the occurrence of degeneracy in the CANDECOMP/PARAFAC model," *Psychometrika*, vol. 73, pp. 431–439, 2008.
- [7] P. Tichavský, A.-H. Phan, and Z. Koldovský, "Cramér-Rao-induced bounds for CANDECOMP/PARAFAC tensor decomposition," *IEEE Trans. Signal Process.*, vol. 61, no. 8, pp. 1986–1997, 2013.
- [8] A. Cichocki, R. Zdunek, A.-H. Phan, and S. Amari, *Nonnegative Matrix and Tensor Factorizations: Applications to Exploratory Multi-way Data Analysis and Blind Source Separation*. Chichester, U.K.: Wiley, 2009.
- [9] L. De Lathauwer, B. De Moor, and J. Vandewalle, "On the best rank-1 and rank-(R1, R2, ..., RN) approximation of higher-order tensors," *SIAM J. Matrix Anal. Appl.*, vol. 21, no. 4, pp. 1324–1342, 2000.
- [10] E. Sanchez and B. R. Kowalski, "Tensorial resolution: A direct trilinear decomposition," *J. Chemometr.*, vol. 4, pp. 29–45, 1990.
- [11] A.-H. Phan, P. Tichavský, and A. Cichocki, "CANDECOMP/PARAFAC decomposition of high-order tensors through tensor reshaping," *IEEE Trans. Signal Process.*, vol. 61, no. 19, pp. 4847–4860, 2013.
- [12] A. Stegeman, "CANDECOMP/PARAFAC: From diverging components to a decomposition in block terms," *SIAM J. Matrix Anal. Appl.*, vol. 33, no. 2, pp. 291–316, 2012.
- [13] N. Sidiropoulos, G. Giannakis, and R. Bro, "Blind PARAFAC receivers for DS-CDMA systems," *IEEE Trans. Signal Process.*, vol. 48, no. 3, pp. 810–823, 2000.

- [14] A. S. Field and D. Graupe, "Topographic component (parallel factor) analysis of multichannel evoked potentials: Practical issues in trilinear spatiotemporal decomposition," *Brain Topograph.*, vol. 3, no. 4, pp. 407–423, 1991.
- [15] K. Vanderperren, B. Mijović, N. Novitskiy, B. Vanrumste, P. Stiers, and B. R. H. Van den Bergh *et al.*, "Single trial ERP reading based on parallel factor analysis," *Psychophysiol.*, vol. 50, no. 1, pp. 97–110, 2013.
- [16] B. M. Hochwald and A. Nehorai, "Concentrated Cramér-Rao bound expressions," *IEEE Trans. Inf. Theory*, vol. 40, no. 2, pp. 363–371, 1994.
- [17] P. Stoica and B. C. Ng, "On the Cramér-Rao bound under parametric constraints," *IEEE Signal Process. Lett.*, vol. 5, no. 7, pp. 177–179, Jul. 1998.
- [18] T. J. Moore, "A Theory of Cramér-Rao bounds for constrained parametric models," Ph.D. dissertation, Univ. of Maryland, College Park, MD, USA, 2010.



Anh-Huy Phan (M'15) received the master's degree from Hochiminh University of Technology, Vietnam, in 2005, and the Ph.D. degree from the Kyushu Institute of Technology, Japan, in 2011. He is a research scientist at the Laboratory for Advanced Brain Signal Processing, RIKEN. He has served on the editorial board of the International Journal of Computational Mathematics. His research interests include multilinear algebra, tensor computation, blind source separation and brain computer interface.



Petr Tichavský (M'98–SM'04) received the Ph.D. degree in theoretical cybernetics from the Czechoslovak Academy of Sciences in 1992. Since that time he is with the Institute of Information Theory and Automation of the Czech Academy of Sciences in Prague. In 1994 he received the Fulbright grant for a 10 month fellowship at Yale University, Department of Electrical Engineering, in New Haven, U.S.A. He is author and co-author of research papers in the area of sinusoidal frequency/frequency-rate estimation, adaptive filtering and tracking of time varying signal parameters, algorithm-independent bounds on achievable performance, sensor array processing, independent component analysis and blind source separation, and tensor decompositions. Petr Tichavský served as associate editor of the IEEE SIGNAL PROCESSING LETTERS from 2002 to 2004, and as associate editor of the IEEE TRANSACTIONS ON SIGNAL PROCESSING from 2005 to 2009 and from 2011 to now. Petr Tichavský has also served as a general co-chair of the 36th IEEE International Conference on Acoustics, Speech and Signal Processing ICASSP 2011 in Prague, Czech Republic, and as general co-chair of the 12th International Conference on Latent Variable Analysis and Signal Separation LVA/ICA 2015 in Liberec, Czech Republic.



Andrzej Cichocki (M'96–SM'07–F'13) received the M.Sc. (with honors), Ph.D. and Dr.Sc. (Habilitation) degrees, all in electrical engineering from Warsaw University of Technology (Poland). He spent several years at University Erlangen (Germany) as an Alexander-von-Humboldt Research Fellow and Guest Professor. In 1995–1997 he was a team leader of the Laboratory for Artificial Brain Systems, at Frontier Research Program RIKEN (Japan), in the Brain Information Processing Group. He is currently a Senior Team Leader and Head of the laboratory for Advanced Brain Signal Processing, at RIKEN Brain Science Institute (Japan) and affiliated full professor in several universities. He has served as an Associated Editor of the IEEE TRANSACTIONS ON NEURAL NETWORKS, the IEEE TRANSACTIONS ON CYBERNETICS, the IEEE TRANSACTIONS ON SIGNALS PROCESSING, the *Journal of Neuroscience Methods* and as founding Editor in Chief for the *Journal Computational Intelligence and Neuroscience*. He is (co)author of more than 400 technical journal papers and 4 monographs in English (two of them translated to Chinese). His publications currently report over 25,000 citations according to Google Scholar, with an h-index of 70. His researches focus on tensor decompositions, brain machine interface, human robot interactions, and EEG hyper-scanning, and their practical applications.