

Supplementary Materials: Proofs and Technical Details for “Parsimonious Tensor Response Regression” Lexin Li and Xin Zhang

A Some preliminary results

We will apply the following two results repeatedly. For a positive definite matrix $\mathbf{B}$,

$$\mathbf{B} = \arg \min_{\mathbf{A} > 0} \{ \log |\mathbf{A}| + \text{tr}(\mathbf{A}^{-1} \mathbf{B}) \}.$$

In addition (Kolda, 2006; Proposition 3.7),

$$\mathbf{A} = \mathbf{B} \times_1 \mathbf{C}_1 \times_2 \cdots \times_N \mathbf{C}_N \Leftrightarrow \mathbf{A}_{(n)} = \mathbf{C}_n \mathbf{B}_{(n)} (\mathbf{C}_N \otimes \cdots \otimes \mathbf{C}_{n+1} \otimes \mathbf{C}_{n-1} \otimes \cdots \otimes \mathbf{C}_1)^\top.$$

For the envelope parameterization of $\mathbf{B} = \llbracket \boldsymbol{\Theta}; \boldsymbol{\Gamma}_1, \dots, \boldsymbol{\Gamma}_m, \mathbf{I}_p \rrbracket$, we have

$$\mathbf{B}_{(m+1)} = \boldsymbol{\Theta}_{(m+1)} (\boldsymbol{\Gamma}_m^\top \otimes \cdots \otimes \boldsymbol{\Gamma}_1^\top), \text{ and } \boldsymbol{\Theta}_{(m+1)} = \mathbf{B}_{(m+1)} (\boldsymbol{\Gamma}_m \otimes \cdots \otimes \boldsymbol{\Gamma}_1).$$

As for $\mathbf{Y} = \llbracket \mathbf{Z}; \boldsymbol{\Gamma}_1, \dots, \boldsymbol{\Gamma}_m \rrbracket \in \mathbb{R}^{r_1 \times \cdots \times r_m}$, we have the following results by treating $\mathbf{Y}$ as $(m+1)$ -th order tensor with $r_{m+1} = 1$, and treating $\text{vec}^\top(\mathbf{Y})$ as the mode- $(m+1)$ matricization.

$$\begin{aligned} \text{vec}^\top(\mathbf{Z}) &= \text{vec}^\top(\mathbf{Y} \times_1 \boldsymbol{\Gamma}_1^\top \times_2 \cdots \times_m \boldsymbol{\Gamma}_m^\top) = \text{vec}^\top(\mathbf{Y}) (\boldsymbol{\Gamma}_m \otimes \cdots \otimes \boldsymbol{\Gamma}_1), \\ \text{vec}(\mathbf{Z}) &= (\boldsymbol{\Gamma}_m^\top \otimes \cdots \otimes \boldsymbol{\Gamma}_1^\top) \text{vec}(\mathbf{Y}). \end{aligned}$$

B Proof for Proposition 1 and 2

We first prove Proposition 1. Then Proposition 2 follows directly from the results of Proposition 1 and the definitions of reducing subspace and Tucker decomposition.

Under the tensor linear model $\mathbf{Y} = \mathbf{B}^{\bar{\times}_{(m+1)}} \mathbf{X} + \boldsymbol{\epsilon}$, we can write

$$\mathbf{Y} \times_k \mathbf{Q}_k = (\mathbf{B}^{\bar{\times}_{(m+1)}} \mathbf{X}) \times_k \mathbf{Q}_k + \boldsymbol{\epsilon} \times_k \mathbf{Q}_k = (\mathbf{B} \times_k \mathbf{Q}_k)^{\bar{\times}_{(m+1)}} \mathbf{X} + \boldsymbol{\epsilon} \times_k \mathbf{Q}_k,$$

which implies that $\mathbf{Y} \times_k \mathbf{Q}_k | \mathbf{X} \sim \mathbf{Y} \times_k \mathbf{Q}_k$ is equivalent to $\mathbf{B} \times_k \mathbf{Q}_k = 0$.

Similarly, $\mathbf{Y} \times_k \mathbf{Q}_k \perp\!\!\!\perp \mathbf{Y} \times_k \mathbf{P}_k | \mathbf{X}$ is equivalent to $\boldsymbol{\epsilon} \times_k \mathbf{Q}_k \perp\!\!\!\perp \boldsymbol{\epsilon} \times_k \mathbf{P}_k$, where the independence of two tensor-valued random variable is defined as independence of their vectorized forms: $\text{vec}(\boldsymbol{\epsilon} \times_k \mathbf{Q}_k) \perp\!\!\!\perp \text{vec}(\boldsymbol{\epsilon} \times_k \mathbf{P}_k)$. Because of the tensor normal distribution, the independence is equivalent to $\text{cov}\{\text{vec}(\boldsymbol{\epsilon} \times_k \mathbf{Q}_k), \text{vec}(\boldsymbol{\epsilon} \times_k \mathbf{P}_k)\} = 0$, and hence to

$$\boldsymbol{\Sigma}_m \otimes \cdots \otimes \boldsymbol{\Sigma}_{k+1} \otimes \mathbf{Q}_k \boldsymbol{\Sigma}_k \mathbf{P}_k \otimes \boldsymbol{\Sigma}_{k-1} \otimes \cdots \otimes \boldsymbol{\Sigma}_1 = 0.$$

Therefore, we have shown that

$$\mathbf{Y} \times_k \mathbf{Q}_k \perp\!\!\!\perp \mathbf{Y} \times_k \mathbf{P}_k | \mathbf{X} \Leftrightarrow \mathbf{Q}_k \boldsymbol{\Sigma}_k \mathbf{P}_k = 0 \Leftrightarrow \boldsymbol{\Sigma}_k = \mathbf{P}_k \boldsymbol{\Sigma}_k \mathbf{P}_k + \mathbf{Q}_k \boldsymbol{\Sigma}_k \mathbf{Q}_k.$$

C Proof for Proposition 3

We first show that for subspaces $\mathcal{E} = \mathcal{E}_m \otimes \cdots \otimes \mathcal{E}_1$ and covariances $\boldsymbol{\Sigma} = \boldsymbol{\Sigma}_m \otimes \cdots \otimes \boldsymbol{\Sigma}_1$, where $\mathcal{E}_k \subseteq \mathbb{R}^{r_k}$ and $\boldsymbol{\Sigma}_k \in \mathbb{R}^{r_k \times r_k}$, $k = 1, \dots, m$, the two conditions are equivalent: (1) $\mathcal{E}$ reduces $\boldsymbol{\Sigma}$; (2) $\mathcal{E}_k$ reduces $\boldsymbol{\Sigma}_k$ for all $k = 1, \dots, m$.

The statement (2) implies (1) is straightforward. We need to show that (1), $\mathcal{E} = \mathcal{E}_m \otimes \cdots \otimes \mathcal{E}_1$ reduces $\boldsymbol{\Sigma} = \boldsymbol{\Sigma}_m \otimes \cdots \otimes \boldsymbol{\Sigma}_1$, implies (2), $\mathcal{E}_k$ reduces $\boldsymbol{\Sigma}_k$ for all $k = 1, \dots, m$. By definition, (1) implies that $\mathbf{P}_{\mathcal{E}_m \otimes \cdots \otimes \mathcal{E}_1} (\boldsymbol{\Sigma}_m \otimes \cdots \otimes \boldsymbol{\Sigma}_1) \mathbf{Q}_{\mathcal{E}_m \otimes \cdots \otimes \mathcal{E}_1} = 0$, where $\mathbf{Q}_{\mathcal{E}_m \otimes \cdots \otimes \mathcal{E}_1} = \mathbf{I} - \mathbf{P}_{\mathcal{E}_m \otimes \cdots \otimes \mathcal{E}_1} = \mathbf{I} - \mathbf{P}_{\mathcal{E}_m} \otimes \cdots \otimes \mathbf{P}_{\mathcal{E}_1}$. Expanding this, we see that

$$\mathbf{P}_{\mathcal{E}_m} \boldsymbol{\Sigma}_m \otimes \cdots \otimes \mathbf{P}_{\mathcal{E}_1} \boldsymbol{\Sigma}_1 - \mathbf{P}_{\mathcal{E}_m} \boldsymbol{\Sigma}_m \mathbf{P}_{\mathcal{E}_m} \otimes \cdots \otimes \mathbf{P}_{\mathcal{E}_1} \boldsymbol{\Sigma}_1 \mathbf{P}_{\mathcal{E}_1} = 0.$$

If we right-multiply $\mathbf{P}_{\mathcal{E}_m} \otimes \cdots \otimes \mathbf{P}_{\mathcal{E}_2} \otimes \mathbf{I}_{r_1}$ to both sides of the above equation, we obtain

$$\mathbf{P}_{\mathcal{E}_m} \boldsymbol{\Sigma}_m \mathbf{P}_{\mathcal{E}_m} \otimes \cdots \otimes \mathbf{P}_{\mathcal{E}_2} \boldsymbol{\Sigma}_2 \mathbf{P}_{\mathcal{E}_2} \otimes (\mathbf{P}_{\mathcal{E}_1} \boldsymbol{\Sigma}_1 - \mathbf{P}_{\mathcal{E}_1} \boldsymbol{\Sigma}_1 \mathbf{P}_{\mathcal{E}_1}) = 0,$$

which implies that $\mathbf{P}_{\mathcal{E}_1} \Sigma_1 - \mathbf{P}_{\mathcal{E}_1} \Sigma_1 \mathbf{P}_{\mathcal{E}_1} = 0$. Since $\mathbf{P}_{\mathcal{E}_1} \Sigma_1 = \mathbf{P}_{\mathcal{E}_1} \Sigma_1 (\mathbf{P}_{\mathcal{E}_1} + \mathbf{Q}_{\mathcal{E}_1})$, we conclude that $\mathbf{P}_{\mathcal{E}_1} \Sigma_1 \mathbf{Q}_{\mathcal{E}_1} = 0$ and hence $\mathcal{E}_1$ reduces Σ_1 . Using the similar argument, we can see that $\mathcal{E}_k$ reduces Σ_k , for all $k = 1, \dots, m$.

We next show that $\mathcal{E} = \mathcal{E}_m \otimes \dots \otimes \mathcal{E}_1$ contains $\text{span}(\mathbf{B}_{(m+1)}^\top)$ implies $\mathcal{E}_k$ contains $\text{span}(\mathbf{B}_{(k)})$. Let $\mathbf{G}_k \in \mathbb{R}^{r_k \times q_k}$ be semi-orthogonal matrix such that $\text{span}(\mathbf{G}_k) = \text{span}(\mathbf{B}_{(k)})$, then we have Tucker decomposition of $\mathbf{B}$ as $\mathbf{B} = [\![\boldsymbol{\eta}; \mathbf{G}_1, \dots, \mathbf{G}_m, \mathbf{I}_p]\!]$ for some $\boldsymbol{\eta} \in \mathbb{R}^{q_1 \times \dots \times q_m \times p}$. Therefore we can write $\mathbf{B}_{(m+1)}^\top = (\mathbf{G}_m \otimes \dots \otimes \mathbf{G}_1) \boldsymbol{\eta}_{(m+1)}^\top$. Hence $\mathcal{E} = \mathcal{E}_m \otimes \dots \otimes \mathcal{E}_1$ contains $\text{span}(\mathbf{B}_{(m+1)}^\top)$ implies that $\mathcal{E}_k$ contains $\text{span}(\mathbf{G}_k) = \text{span}(\mathbf{B}_{(k)})$.

D Proof for Lemma 1

Here we aim to show that $\Sigma_k^{(0)}$, the Kronecker covariance estimator from Manceur and Dutilleul (2013), is a $\sqrt{n}$ -consistent estimator for Σ_k , for all k , if $\boldsymbol{\varepsilon}_i$, $i = 1, \dots, n$, are i.i.d. from tensor normal distribution with mean 0 and covariances $\Sigma_1, \dots, \Sigma_m$.

First we note that the mode- k matricization $\boldsymbol{\varepsilon}_{i(k)} \in \mathbb{R}^{r_k \times (\prod_{j \neq k} r_j)}$, $i = 1, \dots, n$, are i.i.d. from a matrix normal distribution with mean 0 and covariances Σ_k (the row covariance matrix) and $\Delta_k \equiv \Sigma_m \otimes \dots \otimes \Sigma_{k+1} \otimes \Sigma_{k-1} \otimes \dots \otimes \Sigma_1$ (the column covariance matrix). Then following Gupta and Nagar (1999; Chapter 2, Theorem 2.3.5), we have the following second-order moments:

$$\mathbb{E}(\boldsymbol{\varepsilon}_{i(k)} \boldsymbol{\varepsilon}_{i(k)}^\top) = \Sigma_k \text{tr}(\Delta_k),$$

where $\text{tr}(\cdot)$ is the trace operator of matrices.

In the iterative algorithm of obtaining $\Sigma_k^{(0)}$, the starting value for $\Sigma_k^{(0)}$ is

$$\frac{1}{n \prod_{j \neq k} r_j} \sum_{i=1}^n \mathbf{e}_{i(k)} \mathbf{e}_{i(k)}^\top,$$

which is thus $\sqrt{n}$ -consistent for Σ_k up to a scalar difference due to $\text{tr}(\Delta_k)$. As we discussed in Section 4.1 regarding the identifiability of Σ_k 's, the scalar difference can be resolved after normalizing $\Sigma^{(0)} = \tau \Sigma_m^{(0)} \otimes \dots \otimes \Sigma_1^{(0)}$ according to the scalar $\tau = (n \prod_j r_j)^{-1} \sum_{i=1}^n \text{vec}(\mathbf{e}_i) \{(\Sigma_m^{(0)})^{-1} \otimes \dots \otimes (\Sigma_1^{(0)})^{-1}\} \text{vec}^\top(\mathbf{e}_i)$ at the end of each iteration.

Therefore, after the first iteration, we have obtained $\sqrt{n}$ -consistent estimators $\Sigma_k^{(0)}$ for Σ_k . In the iterations that follow, the updating equations of $\Sigma_k^{(0)}$ for $k = 1, \dots, m$, are obtained by maximizing the tensor normal likelihood function. Hence it is guaranteed that the final estimators of $\Sigma_k^{(0)}$ from the algorithm are also $\sqrt{n}$ -consistent for Σ_k .

E Proof for Theorem 1

E.1 Consistency of the one-step estimator

From Cook and Zhang (2016; Proposition 6) we know that if $\widehat{\mathbf{M}}$ and $\widehat{\mathbf{U}}$ are $\sqrt{n}$ -consistent estimators for $\mathbf{M} > 0$ and $\mathbf{U} \geq 0$, then the 1D algorithm for minimizing $J_n(\mathbf{G}) = \log |\mathbf{G}^\top \widehat{\mathbf{M}} \mathbf{G}| + \log |\mathbf{G}^\top (\widehat{\mathbf{M}} + \widehat{\mathbf{U}})^{-1} \mathbf{G}|$ produce $\sqrt{n}$ -consistent estimator for the projection onto the envelope $\mathcal{E}_{\mathbf{M}}(\mathbf{U})$. We use this result to prove the $\sqrt{n}$ -consistency of $\mathbf{P}_{\mathbf{F}_k}^{os}$, noting that $f_k^{*(0)}$ and Algorithm 2 are special instances of J_n and the 1D algorithm.

First $\Sigma_k^{(0)}$, the Kronecker covariance estimator from Manceur and Dutilleul (2013), is $\sqrt{n}$ -consistent estimator for Σ_k . Next, by comparing $f_k^{*(0)}(\mathbf{G}_k)$ to $J_n(\mathbf{G})$ we see that $\mathbf{N}_k^{(0)} = (n \prod_{j \neq k} r_j)^{-1} \sum_{i=1}^n \mathbf{Y}_{i(k)} \Sigma_{-k}^{(0)} \mathbf{Y}_{i(k)}^\top$ is analogous to $(\widehat{\mathbf{M}} + \widehat{\mathbf{U}})$ in $J_n(\mathbf{G})$, where we have defined

$$\Sigma_{-k}^{(0)} \equiv \left((\Sigma_m^{(0)})^{-1} \otimes \dots \otimes (\Sigma_{k+1}^{(0)})^{-1} \otimes (\Sigma_{k-1}^{(0)})^{-1} \otimes \dots \otimes (\Sigma_1^{(0)})^{-1} \right).$$

Therefore we focus on $\mathbf{N}_k^{(0)} - \Sigma_k^{(0)}$, which is analogous to $\widehat{\mathbf{U}}$ in $J_n(\mathbf{G})$. Recall that $\mathbf{Y}_i = \mathbf{B}^{(0)} \bar{\mathbf{x}}_{(m+1)} \mathbf{X}_i^\top + \mathbf{e}_i$ and that $\Sigma_k^{(0)}$ is obtained based on $\mathbf{e}_i$ as $\Sigma_k^{(0)} = (n \prod_{j \neq k} r_j)^{-1} \sum_{i=1}^n \mathbf{e}_{i(k)} \Sigma_{-k}^{(0)} \mathbf{e}_{i(k)}^\top$. Hence,

$$\mathbf{N}_k^{(0)} - \Sigma_k^{(0)} = (n \prod_{j \neq k} r_j)^{-1} \sum_{i=1}^n (\mathbf{Y}_i - \mathbf{e}_i)_{(k)} \Sigma_{-k}^{(0)} (\mathbf{Y}_i - \mathbf{e}_i)_{(k)}^\top,$$

where we recognize $\mathbf{Y}_i - \mathbf{e}_i = \mathbf{B}^{(0)} \bar{\mathbf{x}}_{(m+1)} \mathbf{X}_i^\top$. Define the following “scaled” regression coefficient tensor

$$\mathbf{B}^{\{k\}} = \llbracket \mathbf{B}; \Sigma_1^{-1/2}, \dots, \Sigma_{k-1}^{-1/2}, \mathbf{I}_{r_k}, \Sigma_{k+1}^{-1/2}, \dots, \Sigma_m^{-1/2}, \Sigma_{\mathbf{X}}^{1/2} \rrbracket.$$

Then we see that $\mathbf{N}_k^{(0)} - \Sigma_k^{(0)} = (\prod_{j \neq k} r_j)^{-1} \widehat{\mathbf{B}}^{\{k\}} (\widehat{\mathbf{B}}^{\{k\}})^\top$ for

$$\widehat{\mathbf{B}}^{\{k\}} = \llbracket \mathbf{B}; (\Sigma_1^{(0)})^{-1/2}, \dots, (\Sigma_{k-1}^{(0)})^{-1/2}, \mathbf{I}_{r_k}, (\Sigma_{k+1}^{(0)})^{-1/2}, \dots, (\Sigma_m^{(0)})^{-1/2}, \widehat{\Sigma}_{\mathbf{X}}^{1/2} \rrbracket,$$

where $\widehat{\Sigma}_{\mathbf{X}} = n^{-1} \sum_{i=1}^n \mathbf{X}_i \mathbf{X}_i^\top$. We claim that $\mathbf{N}_k^{(0)} - \Sigma_k^{(0)}$ is a $\sqrt{n}$ -consistent estimator for $\mathbf{B}_{(k)}^{\{k\}} (\mathbf{B}_{(k)}^{\{k\}})^\top$ up to an upfront scaling constant, $(\prod_{j \neq k} r_j)^{-1}$, based on the fact that the sample matrices $\Sigma_j^{(0)}$ and $\widehat{\Sigma}_{\mathbf{X}}$ are asymptotically independent of the OLS estimator $\mathbf{B}^{(0)}$. By Cook and Zhang (2016; Proposition 6) we now see that $\mathbf{P}_{\Gamma_k}^{os}$ is a $\sqrt{n}$ -consistent estimator for the projection onto the envelope $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)}^{\{k\}} (\mathbf{B}_{(k)}^{\{k\}})^\top) = \mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)}^{\{k\}})$. By definition of $\mathbf{B}^{\{k\}}$ and the property of Tucker operator, we have

$$\mathbf{B}_{(k)}^{\{k\}} = \mathbf{B}_{(k)} \left(\Sigma_1^{-1/2} \otimes \dots \otimes \Sigma_{k-1}^{-1/2} \otimes \Sigma_{k+1}^{-1/2} \otimes \dots \otimes \Sigma_m^{-1/2} \otimes \Sigma_{\mathbf{X}}^{1/2} \right),$$

which implies $\text{span}(\mathbf{B}_{(k)}^{\{k\}}) = \text{span}(\mathbf{B}_{(k)})$, and hence $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)}^{\{k\}}) = \mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$.

So far we have shown that $\mathbf{P}_{\Gamma_k}^{os}$ is a $\sqrt{n}$ -consistent estimator for the projection onto the envelope $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$. The second part of the proposition is based on the $\sqrt{n}$ -consistency of $\mathbf{B}^{(0)} = \mathbf{B}_{\text{OLS}}$ and $\mathbf{P}_{\Gamma_k}^{os}$, $k = 1, \dots, m$, and the definition of $\mathbf{B}^{os}$:

$$\mathbf{B}^{os} = \llbracket \mathbf{B}^{(0)}; \mathbf{P}_{\Gamma_1}^{os}, \dots, \mathbf{P}_{\Gamma_m}^{os}, \mathbf{I}_p \rrbracket.$$

E.2 Consistency of the likelihood-based estimator

The $\sqrt{n}$ -consistency of the likelihood-based estimator from minimizing $\ell(\mathbf{B}, \Sigma)$ relies on Shapiro's (1986) results on the asymptotics of over-parameterized structural models. The proof is parallel to the proof of Proposition 4 in Cook and Zhang (2015) and is thus omitted.

F Proof for Theorem 2 (including derivations for the updating equations in Section 4.1)

Following the discussion in Section 5.2 of the paper, the iterative estimator is not guaranteed to be necessarily the maximum likelihood estimator (MLE), due to the existence of multiple local minima. However, it is asymptotically equivalent to the MLE. This is because, under the tensor normal distribution, the initialization of Algorithm 1 is built upon $\sqrt{n}$ -consistent estimators, while each parameter in Algorithm 1 is iteratively obtained along the partial derivative of

the log-likelihood. From the classical theory of point estimation, we know that one Newton-Raphson step from the starting value provides an estimator that is asymptotically equivalent to the MLE even in the presence of multiple local minima (Lehmann and Casella, 1998, p. 454). Consequently, to prove Theorem 2, we only focus on the asymptotic properties of the theoretical MLE.

Specifically, we need to derive all the objective functions and updating equations in Section 4.1 from the negative normal log-likelihood function:

$$\ell(\mathbf{B}, \Sigma) = \log |\Sigma| + n^{-1} \sum_{i=1}^n \{ \text{vec}(\mathbf{Y}_i) - \mathbf{B}_{(m+1)}^\top \mathbf{X}_i \} \Sigma^{-1} \{ \text{vec}(\mathbf{Y}_i) - \mathbf{B}_{(m+1)}^\top \mathbf{X}_i \},$$

which is to be optimized over $\Sigma = \Sigma(\{\Gamma_k, \Omega_k, \Omega_{0k}\}_{k=1}^m)$ and $\mathbf{B} = \mathbf{B}(\{\Gamma_k\}_{k=1}^m, \Theta)$ for semi-orthogonal $\Gamma_k \in \mathbb{R}^{r_k \times u_k}$, positive definite and symmetric $\Omega_k \in \mathbb{S}^{u_k}$ and $\Omega_{0k} \in \mathbb{S}^{r_k - u_k}$, and $\Theta \in \mathbb{R}^{u_m \times \dots \times u_1 \times p}$. Since it is impossible to obtain explicit forms of MLEs, we show that the series of equations used in Algorithm 1 come from partially minimizing ℓ with specified parameters fixed. We summarize our findings in the following statements and give detailed derivations immediately after. Note that the updating equations in Section 4.1 is obtained from the following equations by superscripting (t) on the left hand sides and superscripting $(t+1)$ on the right hand sides of equations.

Under the tensor normal assumption, the MLEs satisfy the following equations,

$$\begin{aligned} \Theta &= \mathbb{Z} \times_{(m+1)} \{ (\mathbb{X}\mathbb{X}^\top)^{-1} \mathbb{X} \}, \\ \mathbf{B} &= \mathbb{Y} \times_1 \mathbf{P}_{\Gamma_1} \times_2 \dots \times_m \mathbf{P}_{\Gamma_m} \times_{(m+1)} \{ (\mathbb{X}\mathbb{X}^\top)^{-1} \mathbb{X} \}, \\ &= [\widehat{\mathbf{B}}_{\text{OLS}}; \mathbf{P}_{\Gamma_1}, \dots, \mathbf{P}_{\Gamma_m}, \mathbf{I}_p] \\ \Omega_k &= \frac{1}{n \prod_{j \neq k} r_j} \sum_{i=1}^n \mathbf{s}_{i(k)} \{ \Omega_m^{-1} \otimes \dots \otimes \Omega_{k+1}^{-1} \otimes \Omega_{k-1}^{-1} \otimes \dots \otimes \Omega_1^{-1} \} \mathbf{s}_{i(k)}^\top, \\ \Omega_{0k} &= \frac{1}{n \prod_{j \neq k} r_j} \sum_{i=1}^n \Gamma_{0k}^\top \mathbf{Y}_{i(k)} \{ \Sigma_m^{-1} \otimes \dots \otimes \Sigma_{k+1}^{-1} \otimes \Sigma_{k-1}^{-1} \otimes \dots \otimes \Sigma_1^{-1} \} \mathbf{Y}_{i(k)}^\top \Gamma_{0k}, \end{aligned}$$

where the data tensors $\mathbf{Z}_i$ and $\mathbb{Z}$ are defined according to $\mathbf{Z} = [\mathbb{Y}; \Gamma_1^\top, \dots, \Gamma_m^\top]$, and the residual tensor $\mathbf{s}_i = \mathbf{Z}_i - \Theta \bar{\times}_{(m+1)} \mathbf{X}_i$. Under the tensor normal assumption, the MLE of $\{\Gamma_k\}_{k=1}^m$ can be

obtained as minimizer of the following objective function

$$f_k(\Gamma_k) = \log |\Gamma_k^\top \mathbf{M}_k \Gamma_k| + \log |\Gamma_k^\top \mathbf{N}_k^{-1} \Gamma_k|,$$

where $\mathbf{M}_k, \mathbf{N}_k$ are defined as

$$\begin{aligned} \mathbf{M}_k &= (n \prod_{j \neq k} r_j)^{-1} \sum_{i=1}^n \delta_{i(k)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_{k+1}^{-1} \otimes \Sigma_{k-1}^{-1} \otimes \cdots \otimes \Sigma_1^{-1}) \delta_{i(k)}^\top, \\ \mathbf{N}_k &= (n \prod_{j \neq k} r_j)^{-1} \sum_{i=1}^n \mathbf{Y}_{i(k)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_{k+1}^{-1} \otimes \Sigma_{k-1}^{-1} \otimes \cdots \otimes \Sigma_1^{-1}) \mathbf{Y}_{i(k)}^\top, \end{aligned}$$

where $\delta_{i(k)}$ is the k -th matricization of the residual δ_i ,

$$\begin{aligned} \delta_i &= \mathbf{Y}_i - \mathbb{Y} \times_1 \mathbf{P}_{\Gamma_1} \times_2 \cdots \times_{(k-1)} \mathbf{P}_{\Gamma_{k-1}} \times_{(k+1)} \mathbf{P}_{\Gamma_{k+1}} \times_{k+2} \cdots \times_m \mathbf{P}_{\Gamma_m} \times_{(m+1)} \mathbf{X}_i^\top \{(\mathbb{X} \mathbb{X}^\top)^{-1} \mathbb{X}\} \\ &= \mathbf{Y}_i - [\widehat{\mathbf{B}}_{\text{OLS}}; \mathbf{P}_{\Gamma_1}, \dots, \mathbf{P}_{\Gamma_{k-1}}, \mathbf{I}_{r_k}, \mathbf{P}_{\Gamma_{k+1}}, \dots, \mathbf{P}_{\Gamma_m}, \mathbf{I}_p] \bar{\times}_{(m+1)} \mathbf{X}_i. \end{aligned}$$

F.1 Estimation of other parameters given $\{\Gamma_k\}_{k=1}^m$

We first decompose the $\log |\Sigma|$ term as

$$\log |\Sigma| = \log |\Sigma_m \otimes \cdots \otimes \Sigma_1| = \sum_{k=1}^m \{(\prod_{j \neq k} r_j) \log |\Sigma_k|\} = \sum_{k=1}^m \{(\prod_{j \neq k} r_j) (\log |\Omega_k| + \log |\Omega_{0k}|)\},$$

which is essentially $2m$ additive terms of $\log |\Omega_k|$ and $\log |\Omega_{0k}|$. The conditional log-likelihood can be separated into two independent parts regarding $\mathbb{P}(\mathbf{Y})|\mathbf{X}$ and $\mathbb{Q}(\mathbf{Y}) \sim \mathbb{Q}(\mathbf{Y})|\mathbf{X}$. Studying the regression of $\mathbb{P}(\mathbf{Y})$ on $\mathbf{X}$ is essentially studying that of $\mathbf{Z} = [\mathbf{Y}; \Gamma_1^\top, \dots, \Gamma_m^\top]$ on $\mathbf{X}$. Following the discussion in the paper, we have the MLEs for $\Theta, \mathbf{B}$ and $\{\Omega_k\}_{k=1}^m$ given $\{\Gamma_k\}_{k=1}^m$.

We next derive the MLE equations for $\{\Omega_{0k}\}_{k=1}^m$ with given $\{\Gamma_k\}_{k=1}^m$. Without loss of generality, we write down our derivations with respect to Ω_{01} . We decompose $\Sigma^{-1} = \Delta_1 + \Delta_{01}$, according to Ω_1 and Ω_{01} in the decomposition of $\Sigma_1^{-1} = \Gamma_1 \Omega_1^{-1} \Gamma_1^\top + \Gamma_{01} \Omega_{01}^{-1} \Gamma_{01}^\top$.

$$\begin{aligned} \Delta_1 &= \Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1} \otimes (\Gamma_1 \Omega_1^{-1} \Gamma_1^\top) \\ &= (\Gamma_m \Omega_m^{-1} \Gamma_m^\top + \Gamma_{0m} \Omega_{0m}^{-1} \Gamma_{0m}^\top) \otimes \cdots \otimes (\Gamma_2 \Omega_2^{-1} \Gamma_2^\top + \Gamma_{02} \Omega_{02}^{-1} \Gamma_{02}^\top) \otimes (\Gamma_1 \Omega_1^{-1} \Gamma_1^\top), \\ \Delta_{01} &= \Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1} \otimes (\Gamma_{01} \Omega_{01}^{-1} \Gamma_{01}^\top) \\ &= (\Gamma_m \Omega_m^{-1} \Gamma_m^\top + \Gamma_{0m} \Omega_{0m}^{-1} \Gamma_{0m}^\top) \otimes \cdots \otimes (\Gamma_2 \Omega_2^{-1} \Gamma_2^\top + \Gamma_{02} \Omega_{02}^{-1} \Gamma_{02}^\top) \otimes (\Gamma_{01} \Omega_{01}^{-1} \Gamma_{01}^\top). \end{aligned}$$

Hence, we can write the negative partial log-likelihood for solving Ω_{01} as

$$\ell(\Omega_{01}|\{\Gamma_k\}_{k=1}^m) \simeq \left(\prod_{j>1} r_j\right) \log |\Omega_{01}| + n^{-1} \sum_{i=1}^n \{\text{vec}(\mathbf{Y}_i) - \mathbf{B}_{(m+1)}^\top \mathbf{X}_i\}^\top \Delta_{01} \{\text{vec}(\mathbf{Y}_i) - \mathbf{B}_{(m+1)}^\top \mathbf{X}_i\}.$$

Recall that $\mathbf{B}_{(m+1)} = \boldsymbol{\eta}_{(m+1)}(\Gamma_m^\top \otimes \cdots \otimes \Gamma_1^\top)$, hence $\Delta_{01} \mathbf{B}_{(m+1)}^\top = 0$ as a result of $\Gamma_{01} \Omega_{01}^{-1} \Gamma_{01}^\top \cdot \Gamma_1 = 0$. Therefore,

$$\ell(\Omega_{01}|\{\Gamma_k\}_{k=1}^m) \simeq \left(\prod_{j>1} r_j\right) \log |\Omega_{01}| + n^{-1} \sum_{i=1}^n \text{vec}^\top(\mathbf{Y}_i) \Delta_{01} \text{vec}(\mathbf{Y}_i).$$

The quadratic form $\text{vec}^\top(\mathbf{Y}_i) \Delta_{01} \text{vec}(\mathbf{Y}_i)$ equals the squared norm of $\text{vec}(\mathbf{Y}_i \times_1 \Omega_{01}^{-1/2} \Gamma_{01}^\top \times_2 \Sigma_2^{-1/2} \times \cdots \times_m \Sigma_m^{-1/2}) \equiv \text{vec}(\mathbf{V}_i)$. By definition of tensor norm, $\|\mathbf{V}_i\|^2 = \|\text{vec}(\mathbf{V}_i)\|^2 = \|\mathbf{V}_{i(1)}\|_F^2 = \text{tr}(\mathbf{V}_{i(1)}^\top \mathbf{V}_{i(1)})$, where $\mathbf{V}_{i(1)}$ is the mode-1 matricization of $\mathbf{V}_i$:

$$\mathbf{V}_{i(1)} = (\Omega_{01}^{-1/2} \Gamma_{01}^\top) \mathbf{Y}_{i(1)} \left(\Sigma_m^{-1/2} \otimes \cdots \otimes \Sigma_2^{-1/2} \right).$$

Hence

$$\text{tr}(\mathbf{V}_{i(1)}^\top \mathbf{V}_{i(1)}) = \text{tr}(\mathbf{V}_{i(1)} \mathbf{V}_{i(1)}^\top) = \text{tr} \left\{ \Omega_{01}^{-1} \Gamma_{01}^\top \mathbf{Y}_{i(1)} \left(\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1} \right) \mathbf{Y}_{i(1)}^\top \Gamma_{01} \right\}.$$

The partial conditional log-likelihood becomes

$$\begin{aligned} \ell(\Omega_{01}|\{\Gamma_k\}_{k=1}^m) &\simeq \left(\prod_{j>1} r_j\right) \log |\Omega_{01}| + n^{-1} \sum_{i=1}^n \text{tr}(\mathbf{V}_{i(1)}^\top \mathbf{V}_{i(1)}) \\ &= \left(\prod_{j>1} r_j\right) \log |\Omega_{01}| + n^{-1} \text{tr} \left[\Omega_{01}^{-1} \sum_{i=1}^n \left\{ \Gamma_{01}^\top \mathbf{Y}_{i(1)} \left(\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1} \right) \mathbf{Y}_{i(1)}^\top \Gamma_{01} \right\} \right], \end{aligned}$$

which lead to the following equations for iteratively solving Ω_{01} .

$$\Omega_{01} = \left(n \prod_{j>1} r_j \right)^{-1} \sum_{i=1}^n \Gamma_{01}^\top \mathbf{Y}_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \mathbf{Y}_{i(1)}^\top \Gamma_{01},$$

where $\Sigma_k = \Gamma_k \Omega_k \Gamma_k^\top + \Gamma_{0k} \Omega_{0k} \Gamma_{0k}^\top$ for $k = 1, \dots, m$. It is then easy to obtain the following result, for any k ,

$$\Omega_{0k} = \left(n \prod_{j \neq k} r_j \right)^{-1} \sum_{i=1}^n \Gamma_{0k}^\top \mathbf{Y}_{i(k)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_{k+1}^{-1} \otimes \Sigma_{k-1}^{-1} \otimes \cdots \otimes \Sigma_1^{-1}) \mathbf{Y}_{i(k)}^\top \Gamma_{0k}.$$

F.2 Estimation of Γ_1 given $\{\Gamma_k, \Omega_k, \Omega_{0k}\}_{k=2}^m$

Treating the other parameters $\{\Gamma_k, \Omega_k, \Omega_{0k}\}_{k=2}^m$ as fixed constants, we write $\mathbf{B} = \mathbf{B}(\Gamma_1)$ and $\Sigma = \Sigma(\Gamma_1, \Omega_1(\Gamma_1), \Omega_{01}(\Gamma_1))$ as functions of Γ_1 . We then plug them into the likelihood $\ell(\mathbf{B}, \Sigma)$ to partially optimize over Ω_1 and Ω_{01} analytically, and then the objective function for optimizing over $\Gamma_1 \in \mathbb{R}^{r_1 \times u_1}$ as follows.

First, ignoring all the fixed constants, the log-determinant term in $\ell(\mathbf{B}, \Sigma)$ becomes

$$\log |\Sigma| \simeq \left(\prod_{j=2}^m r_j \right) \{ \log |\Omega_1(\Gamma_1)| + \log |\Omega_{01}(\Gamma_1)| \}.$$

Similar to the previous section, we can decompose $\Sigma^{-1} = \Sigma_m^{-1} \otimes \dots \otimes \Sigma_1^{-1}$ into two parts according to the decomposition $\Sigma_1^{-1} = \Gamma_1 \Omega_1^{-1} \Gamma_1^\top + \Gamma_{01} \Omega_{01}^{-1} \Gamma_{01}^\top$. Then we can write

$$\begin{aligned} & \sum_{i=1}^n \{ \text{vec}(\mathbf{Y}_i) - \mathbf{B}_{(m+1)}^\top \mathbf{X}_i \} \Sigma^{-1} \{ \text{vec}(\mathbf{Y}_i) - \mathbf{B}_{(m+1)}^\top \mathbf{X}_i \} \\ &= \sum_{i=1}^n \text{tr} \left\{ \Omega_{01}^{-1} \Gamma_{01}^\top \mathbf{Y}_{i(1)} (\Sigma_m^{-1} \otimes \dots \otimes \Sigma_2^{-1}) \mathbf{Y}_{i(1)}^\top \Gamma_{01} \right\} + \sum_{i=1}^n \text{tr} \left\{ \Omega_1^{-1} \Gamma_1^\top \mathbf{e}_{i(1)} (\Sigma_m^{-1} \otimes \dots \otimes \Sigma_2^{-1}) \mathbf{e}_{i(1)}^\top \Gamma_1 \right\}, \end{aligned}$$

where $\mathbf{e}_i \equiv \mathbf{e}_i(\Gamma_1) = \mathbf{Y}_i - \mathbf{B}(\Gamma_1) \times_{(m+1)} \mathbf{X}_i^\top$ and $\mathbf{B}(\Gamma_1) = \mathbb{Y} \times_1 \mathbf{P}_{\Gamma_1} \times_2 \dots \times_m \mathbf{P}_{\Gamma_m} \times_{(m+1)} \{(\mathbb{X}\mathbb{X}^\top)^{-1} \mathbb{X}\}$. Then,

$$\begin{aligned} \Gamma_1^\top \mathbf{e}_{i(1)} &= \left(\mathbf{Y}_i \times_1 \Gamma_1^\top - \mathbb{Y} \times_1 \Gamma_1^\top \mathbf{P}_{\Gamma_1} \times_2 \dots \times_m \mathbf{P}_{\Gamma_m} \times_{(m+1)} \mathbf{X}_i^\top \{(\mathbb{X}\mathbb{X}^\top)^{-1} \mathbb{X}\} \right)_{(1)} \\ &= \left(\mathbf{Y}_i \times_1 \Gamma_1^\top - \mathbb{Y} \times_1 \Gamma_1^\top \times_2 \dots \times_m \mathbf{P}_{\Gamma_m} \times_{(m+1)} \mathbf{X}_i^\top \{(\mathbb{X}\mathbb{X}^\top)^{-1} \mathbb{X}\} \right)_{(1)} \\ &= \left\{ \left(\mathbf{Y}_i - \mathbb{Y} \times_1 \mathbf{I}_{r_1} \times_2 \dots \times_m \mathbf{P}_{\Gamma_m} \times_{(m+1)} \mathbf{X}_i^\top \{(\mathbb{X}\mathbb{X}^\top)^{-1} \mathbb{X}\} \right) \times_1 \Gamma_1^\top \right\}_{(1)} \\ &= \Gamma_1^\top \boldsymbol{\delta}_{i(1)}, \end{aligned}$$

where $\boldsymbol{\delta}_i = \left(\mathbf{Y}_i - \mathbb{Y} \times_1 \mathbf{I}_{r_1} \times_2 \dots \times_m \mathbf{P}_{\Gamma_m} \times_{(m+1)} \mathbf{X}_i^\top \{(\mathbb{X}\mathbb{X}^\top)^{-1} \mathbb{X}\} \right)$ does not involve Γ_1 . The partially maximized negative log-likelihood now becomes

$$\begin{aligned}
\ell(\mathbf{B}, \Sigma) &= \log |\Sigma| + n^{-1} \sum_{i=1}^n \{ \text{vec}(\mathbf{Y}_i) - \mathbf{B}_{(m+1)}^\top \mathbf{X}_i \} \Sigma^{-1} \{ \text{vec}(\mathbf{Y}_i) - \mathbf{B}_{(m+1)}^\top \mathbf{X}_i \} \\
&\simeq \left(\prod_{j=2}^m r_j \right) \{ \log |\Omega_1(\Gamma_1)| + \log |\Omega_{01}(\Gamma_1)| \} \\
&+ n^{-1} \sum_{i=1}^n \text{tr} \left\{ \Omega_{01}^{-1} \Gamma_{01}^\top \mathbf{Y}_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \mathbf{Y}_{i(1)}^\top \Gamma_{01} \right\} \\
&+ n^{-1} \sum_{i=1}^n \text{tr} \left\{ \Omega_1^{-1} \Gamma_1^\top \delta_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \delta_{i(1)}^\top \Gamma_1 \right\},
\end{aligned}$$

which leads to partial MLE of $\Omega_1(\Gamma_1)$ and $\Omega_{01}(\Gamma_1)$ as

$$\begin{aligned}
\Omega_1(\Gamma_1) &= \left(n \prod_{j=2}^m r_j \right)^{-1} \sum_{i=1}^n \Gamma_1^\top \delta_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \delta_{i(1)}^\top \Gamma_1 \\
\Omega_{01}(\Gamma_1) &= \left(n \prod_{j=2}^m r_j \right)^{-1} \sum_{i=1}^n \Gamma_{01}^\top \mathbf{Y}_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \mathbf{Y}_{i(1)}^\top \Gamma_{01}.
\end{aligned}$$

Then, substitute these back to $\ell(\mathbf{B}, \Sigma)$ to get the likelihood-based objective function for Γ_1 :

$$\begin{aligned}
F_n(\Gamma_1) &= \log \left| \left(n \prod_{j=2}^m r_j \right)^{-1} \sum_{i=1}^n \Gamma_1^\top \delta_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \delta_{i(1)}^\top \Gamma_1 \right| \\
&+ \log \left| \left(n \prod_{j=2}^m r_j \right)^{-1} \sum_{i=1}^n \Gamma_{01}^\top \mathbf{Y}_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \mathbf{Y}_{i(1)}^\top \Gamma_{01} \right| \\
&\simeq \log |\Gamma_1^\top \mathbf{M}_1 \Gamma_1| + \log |\Gamma_1^\top \mathbf{N}_1^{-1} \Gamma_1|,
\end{aligned}$$

where $\mathbf{M}_1 = \left(n \prod_{j=2}^m r_j \right)^{-1} \sum_{i=1}^n \delta_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \delta_{i(1)}^\top$ and $\mathbf{N}_1 = \left(n \prod_{j=2}^m r_j \right)^{-1} \sum_{i=1}^n \mathbf{Y}_{i(1)} (\Sigma_m^{-1} \otimes \cdots \otimes \Sigma_2^{-1}) \mathbf{Y}_{i(1)}^\top$. The last step of the above equations come from the fact that $\log |\Gamma_{01}^\top \mathbf{A} \Gamma_{01}| \simeq \log |\Gamma_1 \mathbf{A}^{-1} \Gamma_1|$ for any positive definite symmetric matrix $\mathbf{A}$.

G Proof for Theorem 3

From the proof of Theorem 2, we have seen that $\widehat{\mathbf{B}}_\Gamma = \llbracket \widehat{\mathbf{B}}_{\text{OLS}}; \mathbf{P}_{\Gamma_1}, \dots, \mathbf{P}_{\Gamma_m}, \mathbf{I}_p \rrbracket$, where $\mathbf{P}_{\Gamma_k}$, $k = 1, \dots, m$, are all true projections onto the envelopes (i.e. population values). Then $\text{vec}(\widehat{\mathbf{B}}_\Gamma) = (\mathbf{I}_p \otimes \mathbf{P}_{\Gamma_m} \otimes \cdots \otimes \mathbf{P}_{\Gamma_1}) \text{vec}(\widehat{\mathbf{B}}_{\text{OLS}}) := \mathbf{P}_t \cdot \text{vec}(\widehat{\mathbf{B}}_{\text{OLS}})$. The OLS estimator $\text{vec}(\widehat{\mathbf{B}}_{\text{OLS}})$ is $\sqrt{n}$ -consistent and asymptotically normal with mean zero covariance equals to $\mathbf{U}_{\text{OLS}} =$

$\Sigma_X^{-1} \otimes \Sigma$. Since $\mathbf{P}_t = \mathbf{I}_p \otimes \mathbf{P}_{\Gamma_m} \otimes \cdots \otimes \mathbf{P}_{\Gamma_1}$ is fixed as population truth, we have proven that the envelope estimator $\widehat{\mathbf{B}}_\Gamma$ is $\sqrt{n}$ -consistent and asymptotically normal with mean zero covariance $\mathbf{U}_\Gamma = \mathbf{P}_t \mathbf{U}_{\text{OLS}} \mathbf{P}_t = \Sigma_X^{-1} \otimes \mathbf{P}_{\Gamma_m} \Sigma_m \mathbf{P}_{\Gamma_m} \otimes \cdots \otimes \mathbf{P}_{\Gamma_1} \Sigma_1 \mathbf{P}_{\Gamma_1}$.

H Proof for Theorem 4

Recall that the parameter vectors involved in this Theorem are:

$$\mathbf{h} = \begin{pmatrix} \mathbf{h}_1 \\ \mathbf{h}_2 \end{pmatrix} = \begin{pmatrix} \text{vec}(\mathbf{B}) \\ \text{vech}(\Sigma) \end{pmatrix}, \quad \phi = \begin{pmatrix} \phi_1 \\ \phi_2 \\ \vdots \\ \phi_{m+1} \end{pmatrix} = \begin{pmatrix} \text{vec}(\mathbf{B}) \\ \text{vech}(\Sigma_1) \\ \vdots \\ \text{vech}(\Sigma_m) \end{pmatrix}, \quad \xi = \begin{pmatrix} \xi_1 \\ \vdots \\ \xi_{3m+1} \end{pmatrix},$$

where $\xi_1 = \text{vec}(\Theta)$, $\{\xi_j\}_{j=2}^{m+1} = \{\text{vec}(\Gamma_k)\}_{k=1}^m$, $\{\xi_j\}_{j=m+2}^{2m+1} = \{\text{vech}(\Omega_k)\}_{k=1}^m$, $\{\xi_j\}_{j=2m+2}^{3m+1} = \{\text{vech}(\Omega_{0k})\}_{k=1}^m$. The length of vectors are monotonically decreasing from $\mathbf{h}$ to ϕ and then to ξ , because they corresponding to three nested model assumptions: (M1) $\mathbf{h}$ corresponding to the unrestricted, vectorized linear regression model, or the OLS model; (M2) ϕ is from the Kronecker covariance assumption, so letting $\mathbf{h} = \mathbf{h}(\phi)$ is essentially imposing the Kronecker matrix structure that $\Sigma = \Sigma_m \otimes \cdots \otimes \Sigma_1$; (M3) ξ is based on envelope assumptions of $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$, $k = 1, \dots, m$.

From Theorem 2, we know that $\mathbf{h}^{(\infty)} = \mathbf{h}(\phi^{(\infty)}) = \mathbf{h}(\xi^{(\infty)})$ is the MLE under the envelope assumption, i.e. model assumption (M3). Also, $\mathbf{h}^{(0)} = \mathbf{h}(\phi^{(0)})$ contains the OLS estimator and the Kronecker covariance estimator from Manceur and Dutilleul (2013). Hence it is the MLE under tensor normal assumption, i.e. model assumption (M2). Since $\mathbf{h} = \mathbf{h}(\phi) = \mathbf{h}(\xi)$ is overparameterized, from Shapiro's (1986) we see the following results: $\sqrt{n}(\mathbf{h}^{(0)} - \mathbf{h}_{\text{True}}) \rightarrow N(0, \mathbf{V}_{\mathbf{h}0})$ and $\sqrt{n}(\mathbf{h}^{(\infty)} - \mathbf{h}_{\text{True}}) \rightarrow N(0, \mathbf{V}_{\mathbf{h}\infty})$, where $\mathbf{V}_{\mathbf{h}0} = \mathbf{H}(\mathbf{H}^\top \mathbf{J}_{\mathbf{h}} \mathbf{H})^\dagger \mathbf{H}^\top$, $\mathbf{V}_{\mathbf{h}\infty} = \mathbf{K}(\mathbf{K}^\top \mathbf{J}_{\mathbf{h}} \mathbf{K})^\dagger \mathbf{K}^\top$ and the gradient matrices $\mathbf{H} = \partial \mathbf{h}(\phi) / \partial \phi$ and $\mathbf{K} = \partial \mathbf{h}(\xi) / \partial \xi$. Moreover, the Fisher information matrix $\mathbf{J}_{\mathbf{h}}$ has the same form as in the usual vector-response linear regression model,

$$\mathbf{J}_{\mathbf{h}} = \begin{pmatrix} \Sigma_X \otimes \Sigma^{-1} & 0 \\ 0 & \frac{1}{2} \mathbf{E}_d^T (\Sigma^{-1} \otimes \Sigma^{-1}) \mathbf{E}_d \end{pmatrix},$$

where the expansion matrix $\mathbf{E}_d \in \mathbb{R}^{d^2 \times d(d+1)/2}$ has the corresponding dimension $d = \prod_{k=1}^m r_k$. Then,

$$\mathbf{H}(\mathbf{H}^\top \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^\top = \mathbf{J}_h^{-1/2} \cdot \mathbf{J}_h^{1/2} \mathbf{H}(\mathbf{H}^\top \mathbf{J}_h^{1/2} \mathbf{J}_h^{1/2} \mathbf{H})^\dagger \mathbf{H}^\top \mathbf{J}_h^{1/2} \cdot \mathbf{J}_h^{-1/2} = \mathbf{J}_h^{-1/2} \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{H}} \mathbf{J}_h^{-1/2},$$

and similarly, $\mathbf{K}(\mathbf{K}^\top \mathbf{J}_h \mathbf{K})^\dagger \mathbf{K}^\top = \mathbf{J}_h^{-1/2} \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{K}} \mathbf{J}_h^{-1/2}$. By chain rule, we can write

$$\mathbf{K} = \partial \mathbf{h}(\boldsymbol{\xi}) / \partial \boldsymbol{\xi} = \partial \mathbf{h}(\boldsymbol{\phi}) / \partial \boldsymbol{\phi} \cdot \partial \boldsymbol{\phi}(\boldsymbol{\xi}) / \partial \boldsymbol{\xi} = \mathbf{H} \cdot \partial \boldsymbol{\phi}(\boldsymbol{\xi}) / \partial \boldsymbol{\xi}.$$

Therefore $\text{span}(\mathbf{K}) \subseteq \text{span}(\mathbf{H})$ and $\text{span}(\mathbf{J}_h^{1/2} \mathbf{K}) \subseteq \text{span}(\mathbf{J}_h^{1/2} \mathbf{H})$, which implies that $\mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{K}} = \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{H}} \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{K}} = \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{K}} \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{H}}$. Finally, we have arrived at

$$\begin{aligned} \mathbf{V}_{h0} - \mathbf{V}_{h\infty} &= \mathbf{H}(\mathbf{H}^\top \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^\top - \mathbf{K}(\mathbf{K}^\top \mathbf{J}_h \mathbf{K})^\dagger \mathbf{K}^\top \\ &= \mathbf{J}_h^{-1/2} \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{H}} \mathbf{J}_h^{-1/2} - \mathbf{J}_h^{-1/2} \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{K}} \mathbf{J}_h^{-1/2} \\ &= \mathbf{J}_h^{-1/2} \left(\mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{H}} - \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{K}} \right) \mathbf{J}_h^{-1/2} \\ &= \mathbf{J}_h^{-1/2} \left(\mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{H}} - \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{H}} \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{K}} \right) \mathbf{J}_h^{-1/2} \\ &= \mathbf{J}_h^{-1/2} \mathbf{P}_{\mathbf{J}_h^{1/2} \mathbf{H}} \mathbf{Q}_{\mathbf{J}_h^{1/2} \mathbf{K}} \mathbf{J}_h^{-1/2} \geq 0. \end{aligned}$$

So far, we have proved the main part of the Proposition 6. We next provide details about the gradient matrices $\mathbf{H}$ and $\mathbf{K}$. First of all, since $\mathbf{h}_1 = \boldsymbol{\phi}_1 = \text{vec}(\mathbf{B})$, we have

$$\mathbf{H} = \begin{pmatrix} \mathbf{I}_{p \prod_{k=1}^m r_k} & 0 & \cdots & 0 \\ 0 & \frac{\partial \text{vech}(\boldsymbol{\Sigma})}{\partial \text{vec}(\boldsymbol{\Sigma}_1)} & \cdots & \frac{\partial \text{vech}(\boldsymbol{\Sigma})}{\partial \text{vec}(\boldsymbol{\Sigma}_m)} \end{pmatrix}. \quad (\text{H1})$$

For a symmetric matrix $\mathbf{A} \in \mathbb{R}^{a \times a}$, $\text{vech}(\mathbf{A}) = \mathbf{C}_a \text{vec}(\mathbf{A})$ and $\text{vec}(\mathbf{A}) = \mathbf{E}_a \text{vech}(\mathbf{A})$, where $\mathbf{C}_a \in \mathbb{R}^{a(a+1)/2 \times a^2}$ is the contraction matrix and $\mathbf{E}_a \in \mathbb{R}^{a^2 \times a(a+1)/2}$ is the extraction matrix. Therefore,

$$\frac{\partial \text{vech}(\boldsymbol{\Sigma})}{\partial \text{vech}(\boldsymbol{\Sigma}_k)} = \mathbf{C}_{\prod_{j=1}^m r_j} \frac{\partial \text{vec}(\boldsymbol{\Sigma})}{\partial \text{vec}(\boldsymbol{\Sigma}_k)} \mathbf{E}_{r_k}, \quad k = 1, \dots, m.$$

We then use the Kronecker structure $\boldsymbol{\Sigma} = \boldsymbol{\Sigma}_m \otimes \cdots \otimes \boldsymbol{\Sigma}_1$ to calculate $\frac{\partial \text{vec}(\boldsymbol{\Sigma})}{\partial \text{vec}(\boldsymbol{\Sigma}_k)}$. For $k = 1$ and $k = m$, we can use the formulas from Fackler (2005) for the derivatives of $\text{vec}(\mathbf{A} \otimes \mathbf{B})$ over $\text{vec}(\mathbf{A})$ and over $\text{vec}(\mathbf{B})$. For $2 \leq k \leq m-1$, there is no compact way of writing down the matrix form but elementwise derivatives, which is straightforward but non-trivial.

Next we write $\mathbf{K}$ as

$$\mathbf{K} = \mathbf{H} \cdot \frac{\partial \phi(\boldsymbol{\xi})}{\partial \boldsymbol{\xi}} = \mathbf{H} \mathbf{R}, \quad \mathbf{R} = \left(\left(\frac{\partial \text{vec}(\mathbf{B})}{\partial \boldsymbol{\xi}} \right)^\top \quad \left(\frac{\partial \text{vech}(\boldsymbol{\Sigma}_1)}{\partial \boldsymbol{\xi}} \right)^\top \quad \dots \quad \left(\frac{\partial \text{vech}(\boldsymbol{\Sigma}_m)}{\partial \boldsymbol{\xi}} \right)^\top \right)^\top, \quad (\text{H2})$$

where we calculate each blocks of $\mathbf{R}$ in the following. Recall that $\boldsymbol{\xi} = (\boldsymbol{\xi}_1^\top, \dots, \boldsymbol{\xi}_{3m+1}^\top)^\top$ and $\boldsymbol{\xi}_1 = \text{vec}(\boldsymbol{\Theta})$, $\{\boldsymbol{\xi}_j\}_{j=2}^{m+1} = \{\text{vec}(\boldsymbol{\Gamma}_k)\}_{k=1}^m$, $\{\boldsymbol{\xi}_j\}_{j=m+2}^{2m+1} = \{\text{vech}(\boldsymbol{\Omega}_k)\}_{k=1}^m$, $\{\boldsymbol{\xi}_j\}_{j=2m+2}^{3m+1} = \{\text{vech}(\boldsymbol{\Omega}_{0k})\}_{k=1}^m$.

The first block in $\mathbf{R}$ is

$$\frac{\partial \text{vec}(\mathbf{B})}{\partial \boldsymbol{\xi}} = \begin{pmatrix} \frac{\partial \text{vec}(\mathbf{B})}{\partial \text{vec}(\boldsymbol{\Theta})} & \frac{\partial \text{vec}(\mathbf{B})}{\partial \text{vec}(\boldsymbol{\Gamma}_1)} & \dots & \frac{\partial \text{vec}(\mathbf{B})}{\partial \text{vec}(\boldsymbol{\Gamma}_m)} & 0 & \dots & 0 \end{pmatrix}, \quad (\text{H3})$$

where the zeros are because of $\mathbf{B}$ does not depend on $\boldsymbol{\Omega}_k$ or $\boldsymbol{\Omega}_{0k}$. We re-write $\text{vec}(\mathbf{B})$ as

$$\begin{aligned} \text{vec}(\mathbf{B}) &= \text{vec}(\mathbf{B}_{(m+1)}^\top) = \text{vec}\left((\boldsymbol{\Gamma}_m \otimes \dots \otimes \boldsymbol{\Gamma}_1) \boldsymbol{\Theta}_{(m+1)}^\top\right) \\ &= (\mathbf{I}_p \otimes \boldsymbol{\Gamma}_m \otimes \dots \otimes \boldsymbol{\Gamma}_1) \text{vec}(\boldsymbol{\Theta}_{(m+1)}^\top) \\ &= (\mathbf{I}_p \otimes \boldsymbol{\Gamma}_m \otimes \dots \otimes \boldsymbol{\Gamma}_1) \text{vec}(\boldsymbol{\Theta}). \end{aligned}$$

Thus,

$$\frac{\partial \text{vec}(\mathbf{B})}{\partial \text{vec}(\boldsymbol{\Theta})} = (\mathbf{I}_p \otimes \boldsymbol{\Gamma}_m \otimes \dots \otimes \boldsymbol{\Gamma}_1). \quad (\text{H4})$$

We next introduce the notation of re-arranging the vectorizations of a mode- N tensor $\mathbf{T} \in \mathbb{R}^{d_1 \times \dots \times d_N}$: squared constant matrix $\boldsymbol{\Pi}_n^\top$ satisfies: $\text{vec}(\mathbf{T}) = \boldsymbol{\Pi}_n^\top \text{vec}(\mathbf{T}_{(n)})$. Therefore, $\text{vec}(\mathbf{B}) = \boldsymbol{\Pi}_k^\top \text{vec}(\mathbf{B}_{(k)})$ and hence

$$\begin{aligned} \text{vec}(\mathbf{B}) &= \boldsymbol{\Pi}_k^\top \text{vec}\left(\boldsymbol{\Gamma}_k \boldsymbol{\Theta}_{(k)} (\boldsymbol{\Gamma}_m^\top \otimes \dots \otimes \boldsymbol{\Gamma}_{k+1}^\top \otimes \boldsymbol{\Gamma}_{k-1}^\top \otimes \dots \otimes \boldsymbol{\Gamma}_1^\top)\right) \\ &= \boldsymbol{\Pi}_k^\top \left((\boldsymbol{\Gamma}_m \otimes \dots \otimes \boldsymbol{\Gamma}_{k+1} \otimes \boldsymbol{\Gamma}_{k-1} \otimes \dots \otimes \boldsymbol{\Gamma}_1) \boldsymbol{\Theta}_{(k)}^\top \otimes \mathbf{I}_{r_k} \right) \text{vec}(\boldsymbol{\Gamma}_k), \\ \frac{\partial \text{vec}(\mathbf{B})}{\partial \text{vec}(\boldsymbol{\Gamma}_k)} &= \boldsymbol{\Pi}_k^\top \left((\boldsymbol{\Gamma}_m \otimes \dots \otimes \boldsymbol{\Gamma}_{k+1} \otimes \boldsymbol{\Gamma}_{k-1} \otimes \dots \otimes \boldsymbol{\Gamma}_1) \boldsymbol{\Theta}_{(k)}^\top \otimes \mathbf{I}_{r_k} \right). \end{aligned} \quad (\text{H5})$$

Finally for $\boldsymbol{\Sigma}_k = \boldsymbol{\Gamma}_k \boldsymbol{\Omega}_k \boldsymbol{\Gamma}_k^\top + \boldsymbol{\Gamma}_{0k} \boldsymbol{\Omega}_{0k} \boldsymbol{\Gamma}_{0k}^\top$, we have

$$\frac{\partial \text{vech}(\boldsymbol{\Sigma}_k)}{\partial \boldsymbol{\xi}} = \begin{pmatrix} 0 & \dots & 0 & \frac{\partial \text{vech}(\boldsymbol{\Sigma}_k)}{\partial \text{vec}(\boldsymbol{\Gamma}_k)} & 0 & \dots & 0 & \frac{\partial \text{vech}(\boldsymbol{\Sigma}_k)}{\partial \text{vech}(\boldsymbol{\Omega}_k)} & 0 & \dots & 0 & \frac{\partial \text{vech}(\boldsymbol{\Sigma}_k)}{\partial \text{vech}(\boldsymbol{\Omega}_{0k})} & 0 & \dots & 0 \end{pmatrix}, \quad (\text{H6})$$

where the three nonzero elements are, (analogous to Cook et al. (2010)),

$$\begin{aligned}\frac{\partial \text{vech}(\boldsymbol{\Sigma}_k)}{\partial \text{vec}(\boldsymbol{\Gamma}_k)} &= 2\mathbf{C}_{r_k}(\boldsymbol{\Gamma}_k\boldsymbol{\Omega}_k \otimes \mathbf{I}_{r_k} - \boldsymbol{\Gamma}_{r_k} \otimes \boldsymbol{\Gamma}_{0k}\boldsymbol{\Omega}_{0k}\boldsymbol{\Gamma}_{0k}^\top), \\ \frac{\partial \text{vech}(\boldsymbol{\Sigma}_k)}{\partial \text{vech}(\boldsymbol{\Omega}_k)} &= \mathbf{C}_{r_k}(\boldsymbol{\Gamma}_k \otimes \boldsymbol{\Gamma}_k)\mathbf{E}_{u_k}, \\ \frac{\partial \text{vech}(\boldsymbol{\Sigma}_k)}{\partial \text{vech}(\boldsymbol{\Omega}_{0k})} &= \mathbf{C}_{r_k}(\boldsymbol{\Gamma}_{0k} \otimes \boldsymbol{\Gamma}_{0k})\mathbf{E}_{r_k-u_k}.\end{aligned}$$

Finally the explicit gradient matrices $\mathbf{H}$ is obtained by plugging each blocks (H3)–(H6) into (H1); similarly, the explicit gradient matrices $\mathbf{K}$ is obtained by plugging each blocks (H3)–(H6) into (H2). We have thus completed the proof of this theorem.

Additional References

- Cook, R. D., Li, B. and Chiaromonte, F. (2010), Envelope models for parsimonious and efficient multivariate linear regression, *Statistica Sinica*, **20**(3), 927–960.
- Cook, R. D. and Zhang, X. (2015), Simultaneous envelope for multivariate linear regression, *Technometrics*, **57**(1), 11–25.
- Cook, R. D. and Zhang, X. (2016), Algorithms for envelope estimation, *Journal of Computational and Graphical Statistics*, **In press**.
- Fackler, P. L. (2005). Notes on matrix calculus. <http://www4.ncsu.edu/~pfackler/MatCalc.pdf>, Technical Report.
- Gupta, A. and Nagar, D. (1999), *Matrix variate distributions*, CRC Press.
- Kolda, T. G. (2006). Multilinear operators for higher-order decompositions. United States, Department of Energy, Technical Report.
- Lehmann, E. L. and Casella, G. (1998). *Theory of point estimation*, Springer Science & Business Media.

- Manceur, A. M. and Dutilleul, P. (2013), Maximum likelihood estimation for the tensor normal distribution: Algorithm, minimum sample size, and empirical bias and dispersion, *Journal of Computational and Applied Mathematics*, **239**, 37–49.
- Shapiro, A. (1986), Asymptotic theory of overparameterized structural models, *Journal of the American Statistical Association*, **81**, 142–149.