

Supplementary Materials for “Tensor Envelope Partial Least Squares Regression”

Xin Zhang and Lexin Li

Florida State University and University of California, Berkeley

1 Proofs and Technical Details

Proof of Lemma 1

Proof. From the vectorized linear model (3), we can see that $\mathbf{B}_{(m+1)}$ is the regression coefficient matrix for the multivariate linear regression of $\mathbf{Y}$ on $\text{vec}(\mathbf{X})$. Therefore $\mathbf{B}_{(m+1)} = \text{cov}^{-1}\{\text{vec}(\mathbf{X})\}\text{cov}\{\text{vec}(\mathbf{X}), \mathbf{Y}\} = (\boldsymbol{\Sigma}_m^{-1} \otimes \cdots \otimes \boldsymbol{\Sigma}_1^{-1}) \mathbf{C}_{(m+1)}$. It follows from the basic property of Tucker operator that $\mathbf{B} = \llbracket \mathbf{C}; \boldsymbol{\Sigma}_1^{-1}, \dots, \boldsymbol{\Sigma}_m^{-1}, \mathbf{I}_r \rrbracket$. $\square$

Proof of Lemma 2

Proof. Recall that in the tensor PLS regression, $\text{E}(\mathbf{Y} \mid \mathbf{X}) = \mathbf{B}_{(m+1)}\text{vec}(\mathbf{X}) = \text{E}(\mathbf{Y} \mid \mathbf{T}) = \boldsymbol{\Psi}_{(m+1)}\text{vec}(\mathbf{T})$, where $\mathbf{T} = \llbracket \mathbf{X}; \mathbf{W}_1^\top, \dots, \mathbf{W}_m^\top \rrbracket$. Therefore,

$$\begin{aligned} \mathbf{B}_{(m+1)}\text{vec}(\mathbf{X}) &= \boldsymbol{\Psi}_{(m+1)}\text{vec}(\mathbf{T}) \\ &= \boldsymbol{\Psi}_{(m+1)}\text{vec}(\llbracket \mathbf{X}; \mathbf{W}_1^\top, \dots, \mathbf{W}_m^\top \rrbracket) \\ &= \boldsymbol{\Psi}_{(m+1)} (\mathbf{W}_m^\top \otimes \cdots \otimes \mathbf{W}_1^\top) \cdot \text{vec}(\mathbf{X}) \\ &= \llbracket \boldsymbol{\Psi}; \mathbf{W}_1, \dots, \mathbf{W}_m, \mathbf{I}_r \rrbracket_{(m+1)} \cdot \text{vec}(\mathbf{X}). \end{aligned}$$

This implies that $\mathbf{B} = \llbracket \boldsymbol{\Psi}; \mathbf{W}_1, \dots, \mathbf{W}_m, \mathbf{I}_r \rrbracket$ under the PLS regression assumption, and hence $\widehat{\mathbf{B}}_{\text{PLS}} = \llbracket \widehat{\boldsymbol{\Psi}}; \widehat{\mathbf{W}}_1, \dots, \widehat{\mathbf{W}}_m, \mathbf{I}_r \rrbracket$. The rest of the proof follows from the fact that $\text{cov}\{\text{vec}(\mathbf{T})\} = (\mathbf{W}_m^\top \boldsymbol{\Sigma}_m \mathbf{W}_m) \otimes \cdots \otimes (\mathbf{W}_1^\top \boldsymbol{\Sigma}_1 \mathbf{W}_1)$ and then applying Lemma 1 on the regression of $\mathbf{Y}$ on $\mathbf{T}$ and $\square$

Proof of Proposition 1

Proof. From Lemma 2, we have

$$\begin{aligned}
\widehat{\mathbf{B}}_{\text{PLS}} &= \llbracket \widehat{\Psi}; \widehat{\mathbf{W}}_1, \dots, \widehat{\mathbf{W}}_m, \mathbf{I}_r \rrbracket \\
&= \llbracket [\widehat{\mathbf{C}}_T; (\widehat{\mathbf{W}}_1^\top \widehat{\Sigma}_1 \widehat{\mathbf{W}}_1)^{-1}, \dots, (\widehat{\mathbf{W}}_m^\top \widehat{\Sigma}_m \widehat{\mathbf{W}}_m)^{-1}, \mathbf{I}_r]; \widehat{\mathbf{W}}_1, \dots, \widehat{\mathbf{W}}_m, \mathbf{I}_r \rrbracket \\
&= \llbracket \widehat{\mathbf{C}}_T; \widehat{\mathbf{W}}_1 (\widehat{\mathbf{W}}_1^\top \widehat{\Sigma}_1 \widehat{\mathbf{W}}_1)^{-1}, \dots, \widehat{\mathbf{W}}_m (\widehat{\mathbf{W}}_m^\top \widehat{\Sigma}_m \widehat{\mathbf{W}}_m)^{-1}, \mathbf{I}_r \rrbracket \\
&= \llbracket \widehat{\mathbf{C}}; \widehat{\mathbf{W}}_1 (\widehat{\mathbf{W}}_1^\top \widehat{\Sigma}_1 \widehat{\mathbf{W}}_1)^{-1} \widehat{\mathbf{W}}_1^\top, \dots, \widehat{\mathbf{W}}_m (\widehat{\mathbf{W}}_m^\top \widehat{\Sigma}_m \widehat{\mathbf{W}}_m)^{-1} \widehat{\mathbf{W}}_m^\top, \mathbf{I}_r \rrbracket,
\end{aligned}$$

where the last equality follows from $\widehat{\mathbf{C}}_T = \llbracket \widehat{\mathbf{C}}; \widehat{\mathbf{W}}_1^\top, \dots, \widehat{\mathbf{W}}_m^\top, \mathbf{I}_r \rrbracket$. The conclusion then follows from $\widehat{\mathbf{W}}_1 (\widehat{\mathbf{W}}_1^\top \widehat{\Sigma}_1 \widehat{\mathbf{W}}_1)^{-1} \widehat{\mathbf{W}}_1^\top = \mathbf{P}_{\widehat{\mathbf{W}}_1(\widehat{\Sigma}_1)} \widehat{\Sigma}_1^{-1}$ and $\widehat{\mathbf{B}}_{\text{OLS}} = \llbracket \widehat{\mathbf{C}}; \widehat{\Sigma}_1^{-1}, \dots, \widehat{\Sigma}_m^{-1}, \mathbf{I}_r \rrbracket$. $\square$

Proof of Proposition 2

Proof. First, $\mathbf{X} \times_k \mathbf{Q}_k \perp \mathbf{X} \times_k \mathbf{P}_k$ implies $0 = \text{cov}\{\text{vec}(\mathbf{X} \times_k \mathbf{Q}_k), \text{vec}(\mathbf{X} \times_k \mathbf{P}_k)\} = \Sigma_m \otimes \dots \otimes \mathbf{Q}_k \Sigma_k \mathbf{P}_k \otimes \dots \otimes \Sigma_1$. It holds if and only if $\mathbf{Q}_k \Sigma_k \mathbf{P}_k = 0$. By the definition of a reducing subspace (Cook et al., 2010), $\mathbf{Q}_k \Sigma_k \mathbf{P}_k = 0$ if and only if $\text{span}(\mathbf{P}_k) = \text{span}(\Gamma_k)$ is a reducing subspace of Σ_k , i.e., $\Sigma_k = \Gamma_k \Omega_k \Gamma_k^\top + \Gamma_{0k} \Omega_{0k} \Gamma_{0k}^\top$.

Second, we substitute $\mathbf{X} = \mathbf{X} \times_k \mathbf{P}_k + \mathbf{X} \times_k \mathbf{Q}_k$ into (3) and get

$$\begin{aligned}
\mathbf{Y} &= \mathbf{B}_{(m+1)} \text{vec}(\mathbf{X} \times_k \mathbf{P}_k) + \mathbf{B}_{(m+1)} \text{vec}(\mathbf{X} \times_k \mathbf{Q}_k) + \varepsilon \\
&= (\mathbf{B} \times_k \mathbf{P}_k)_{(m+1)} \text{vec}(\mathbf{X} \times_k \mathbf{P}_k) + (\mathbf{B} \times_k \mathbf{Q}_k)_{(m+1)} \text{vec}(\mathbf{X} \times_k \mathbf{Q}_k) + \varepsilon.
\end{aligned}$$

Therefore, $\mathbf{Y} \perp \mathbf{X} \times_k \mathbf{Q}_k \mid \mathbf{X} \times_k \mathbf{P}_k$ implies $\mathbf{B} \times_k \mathbf{Q}_k = 0$, which is equivalent to $\mathbf{B} = \mathbf{B} \times_k \mathbf{P}_k = \mathbf{B} \times_k \Gamma_k \Gamma_k^\top$, which gives the parametrization $\mathbf{B} = \llbracket \Theta; \Gamma_1, \dots, \Gamma_m, \mathbf{I}_r \rrbracket$. $\square$

Proof of Proposition 3

Proof. We first show $\mathcal{T}_{\Sigma_X}(\mathbf{B}) \subseteq \mathcal{E}_{\Sigma_m}(\mathbf{B}_{(m)}) \otimes \dots \otimes \mathcal{E}_{\Sigma_1}(\mathbf{B}_{(1)})$. From Definition 2, this means we need to show that (a) $\mathcal{E}_{\Sigma_m}(\mathbf{B}_{(m)}) \otimes \dots \otimes \mathcal{E}_{\Sigma_1}(\mathbf{B}_{(1)})$ is a reducing subspace of Σ_X , and (b) it contains $\text{span}(\mathbf{B}_{(m+1)})$. It is straightforward to see (a) from $\Sigma_X = \Sigma_m \otimes \dots \otimes \Sigma_1$ and that $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$ is a reducing subspace of Σ_k for each k . To show (b), we can write $\mathbf{B}_{(k)} = \mathbf{P}_k \mathbf{B}_{(k)}$, where $\mathbf{P}_k$ is the projection on to $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$, because

$\text{span}(\mathbf{B}_{(k)}) \subseteq \mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$ by definition. This further implies that $\mathbf{B} = \mathbf{B} \times_k \mathbf{P}_k$, $k = 1, \dots, m$, and that $\mathbf{B} = \mathbf{B} \times_1 \mathbf{P}_1 \times_2 \dots \times_m \mathbf{P}_m$. Taking mode- $(m+1)$ matricization on both sides of the last equation, we have (b).

We next show $\mathcal{T}_{\Sigma_X}(\mathbf{B}) \supseteq \mathcal{E}_{\Sigma_m}(\mathbf{B}_{(m)}) \otimes \dots \otimes \mathcal{E}_{\Sigma_1}(\mathbf{B}_{(1)})$. By definition, we can write $\mathcal{T}_{\Sigma_X}(\mathbf{B}) = \mathcal{E}_m \otimes \dots \otimes \mathcal{E}_1$ for some $\mathcal{E}_k \subseteq \mathbb{R}^{p_k}$, $k = 1, \dots, m$. It remains to show that $\mathcal{E}_k \supseteq \mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$. We achieve this by showing (c) $\mathcal{E}_k$ is a reducing subspace of Σ_k , (d) $\mathcal{E}_k$ contains $\text{span}(\mathbf{B}_{(k)})$, and then by noticing $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$ is the smallest subspace satisfying (c) and (d). Note that (d) can be directly obtained from Proposition 2. To get (c), we recall that $\mathcal{T}_{\Sigma_X}(\mathbf{B})$ is a reducing subspace of $\Sigma_X = \Sigma_m \otimes \dots \otimes \Sigma_1$ and thus

$$\mathbf{P}_{\mathcal{E}_m \otimes \dots \otimes \mathcal{E}_1}(\Sigma_m \otimes \dots \otimes \Sigma_1) \mathbf{Q}_{\mathcal{E}_m \otimes \dots \otimes \mathcal{E}_1} = 0,$$

where $\mathbf{P}_{\mathcal{E}_m \otimes \dots \otimes \mathcal{E}_1} = \mathbf{P}_{\mathcal{E}_m} \otimes \dots \otimes \mathbf{P}_{\mathcal{E}_1}$ and $\mathbf{Q}_{\mathcal{E}_m \otimes \dots \otimes \mathcal{E}_1} = \mathbf{I} - \mathbf{P}_{\mathcal{E}_m \otimes \dots \otimes \mathcal{E}_1}$. Expand the above equation we get,

$$\mathbf{P}_{\mathcal{E}_m} \Sigma_m \otimes \dots \otimes \mathbf{P}_{\mathcal{E}_1} \Sigma_1 - \mathbf{P}_{\mathcal{E}_m} \Sigma_m \mathbf{P}_{\mathcal{E}_m} \otimes \dots \otimes \mathbf{P}_{\mathcal{E}_1} \Sigma_1 \mathbf{P}_{\mathcal{E}_1} = 0,$$

which implies the following equality by right-multiplying $\mathbf{P}_{\mathcal{E}_m} \otimes \mathbf{P}_{\mathcal{E}_2} \otimes \mathbf{I}_{p_1}$:

$$\mathbf{P}_{\mathcal{E}_m} \Sigma_m \mathbf{P}_{\mathcal{E}_m} \otimes \dots \otimes \mathbf{P}_{\mathcal{E}_2} \Sigma_2 \mathbf{Q}_{\mathcal{E}_2} \otimes \mathbf{P}_{\mathcal{E}_1} \Sigma_1 \mathbf{Q}_{\mathcal{E}_1} = 0,$$

which implies $\mathcal{E}_1$ reduces Σ_1 . Similarly we get $\mathcal{E}_k$ reduces Σ_k for all $k = 1, \dots, m$. This completes the proof. $\square$

Proof of Lemma 3

Proof. From Algorithm 4, we know that $\mathbf{w}_{sk}$ is the dominant eigenvector of $\tilde{\mathbf{C}}_{(s-1)k} \tilde{\mathbf{C}}_{(s-1)k}^\top$, which equals to $\mathbf{Q}_{(s-1)k} \tilde{\mathbf{C}}_{0k} \tilde{\mathbf{C}}_{0k}^\top \mathbf{Q}_{(s-1)k}$. The conclusion follows from noticing $\tilde{\mathbf{C}}_k = \tilde{\mathbf{C}}_{0k} \tilde{\mathbf{C}}_{0k}^\top$. $\square$

Proof of Theorem 1

Proof. From Lemma 3, we see that $\mathbf{w}_{1k}$ is the first eigenvector of $\tilde{\mathbf{C}}_k$. Then for $s > 1$, we have $\mathbf{w}_{sk}$ is the first eigenvector of $\mathbf{Q}_{(s-1)k} \tilde{\mathbf{C}}_k \mathbf{Q}_{(s-1)k}$ where $\mathbf{Q}_{(s-1)k}$ is the projection

onto the orthogonal complement of $\text{span}(\Sigma_k(\mathbf{w}_{1k}, \dots, \mathbf{w}_{(s-1)k}))$. Following the proof of Proposition 4.1 in Cook et al. (2013), we have

$$\mathcal{W}_k^{(0)} \subset \dots \subset \mathcal{W}_k^{(u_k)} = \mathcal{E}_{\Sigma_k}(\tilde{\mathbf{C}}_k) = \mathcal{W}_k^{(u_k+1)} = \dots = \mathcal{W}_k^{(p_k)},$$

where $\mathcal{E}_{\Sigma_k}(\tilde{\mathbf{C}}_k)$ is the Σ_k -envelope of $\text{span}(\tilde{\mathbf{C}}_k)$. We next need to show $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)}) = \mathcal{E}_{\Sigma_k}(\tilde{\mathbf{C}}_k)$. From Lemma 3, we see that $\text{span}(\tilde{\mathbf{C}}_k) = \text{span}(\mathbf{C}_{(k)})$. From Lemma 1, we see that $\text{span}(\mathbf{B}_{(k)}) = \text{span}(\Sigma_k^{-1}\mathbf{C}_{(k)})$. Finally, from Proposition 2.4 of Cook et al. (2010), we have $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)}) = \mathcal{E}_{\Sigma_k}(\Sigma_k^{-1}\mathbf{B}_{(k)})$ and thus $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)}) = \mathcal{E}_{\Sigma_k}(\mathbf{C}_{(k)}) = \mathcal{E}_{\Sigma_k}(\tilde{\mathbf{C}}_k)$. $\square$

Proof of Theorem 2

Proof. From Theorem 1, if d_k is chosen as u_k , the population value $\text{span}(\mathbf{W}_k) = \mathcal{W}_k^{(u_k)} = \mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$. Since the sample version of Algorithm 4 is based on eigen-decomposition and $\sqrt{n}$ -consistent sample covariance matrices $\hat{\mathbf{C}}, \hat{\Sigma}_{\mathbf{Y}}$, and $\hat{\Sigma}_k$, $k = 1, \dots, m$, it is clear that $\hat{\mathbf{W}}_k$ is $\sqrt{n}$ -consistent for $\mathbf{W}_k$. Hence $\hat{\mathbf{W}}_k$ is $\sqrt{n}$ -consistent for the envelope $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$.

For $d_k \geq u_k$, we have $\mathcal{W}_k^{(d_k)} = \mathcal{W}_k^{(u_k)} = \mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$, therefore the projection $\mathbf{P}_{\hat{\mathbf{W}}_k(\hat{\Sigma}_k)}$ is $\sqrt{n}$ -consistent for $\mathbf{P}_{\mathbf{W}_k(\Sigma_k)}$. Since $\mathbf{W}_k$ is a semi-orthogonal basis for the envelope $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$, it is a reducing subspace of Σ_k and hence $\mathbf{P}_{\mathbf{W}_k(\Sigma_k)} = \mathbf{P}_{\mathbf{W}_k}$ by definition. Therefore, $\mathbf{P}_{\hat{\mathbf{W}}_k(\hat{\Sigma}_k)}$ is $\sqrt{n}$ -consistent estimator for the projection onto the envelope $\mathcal{E}_{\Sigma_k}(\mathbf{B}_{(k)})$. Then from Proposition 1, we recall that $\hat{\mathbf{B}}_{\text{PLS}} = \hat{\mathbf{B}}_{\text{OLS}} \times_1 \mathbf{P}_{\hat{\mathbf{W}}_1(\hat{\Sigma}_1)} \times_2 \dots \times_m \mathbf{P}_{\hat{\mathbf{W}}_m(\hat{\Sigma}_m)}$. The $\sqrt{n}$ -consistency of $\hat{\mathbf{B}}_{\text{PLS}}$ then follows from the $\sqrt{n}$ -consistency of $\hat{\mathbf{B}}_{\text{OLS}}$ and $\mathbf{P}_{\hat{\mathbf{W}}_k(\hat{\Sigma}_k)}$. $\square$

2 Additional simulations

We consider an additional simulation example, with a univariate response and a 3-way tensor predictor with higher dimension and rank. Specifically, the simulation setup is similar to that in Section 5.2, with two main changes. The first is that the response is now a scalar, which allows a more direct comparison with the CP method of Zhou et al. (2013) that was designed for a univariate response. The second is that we now consider both the original setup of a predictor tensor $\mathbf{X} \in \mathbb{R}^{20 \times 20 \times 20}$ with a core tensor $\Theta \in \mathbb{R}^{2 \times 2 \times 2}$ ($p = 20, u = 2$), and a new setup of a predictor tensor $\mathbf{X} \in \mathbb{R}^{40 \times 40 \times 40}$ with

(p, u)	Model	Prediction			
		OLS	CP	TEPLS	TEPLS-CV
(20, 2)	I	84.4 (1.3)	7.2 (0.2)	1.4 (0.1)	1.4 (0.1)
	II	42.0 (0.6)	5.6 (0.1)	1.0 (0.1)	1.0 (0.1)
	III	41.0 (0.5)	5.4 (0.1)	1.0 (0.1)	1.0 (0.1)
(40, 5)	I	1504.9 (21.4)	17.7 (0.3)	4.0 (0.1)	4.4 (0.1)
	II	367.5 (4.9)	4.8 (0.1)	1.2 (0.1)	1.2 (0.1)
	III	361.9 (5.3)	4.5 (0.1)	1.0 (0.1)	1.0 (0.1)
(p, u)	Model	Estimation			
		OLS	CP	TEPLS	TEPLS-CV
(20, 2)	I	$> 10^4$	$> 10^3$	1.1 (0.1)	1.1 (0.1)
	II	60.9 (0.3)	20.4 (0.2)	1.7 (0.1)	1.7 (0.1)
	III	$> 10^6$	$> 10^5$	1.7 (0.1)	1.7 (0.1)
(40, 5)	I	$> 10^5$	$> 10^4$	5.8 (0.1)	6.2 (0.1)
	II	305.9 (1.5)	30.9 (0.3)	6.4 (0.1)	6.4 (0.1)
	III	$> 10^7$	$> 10^5$	6.3 (0.1)	6.3 (0.1)

Table S1: Univariate response and 3-way predictor. Performance under various scenarios and comparison of estimators. OLS, CP, tensor envelope PLS with true and estimated envelope dimensions. Reported are the average and standard error (in parenthesis) of the prediction mean squared error evaluated on an independent testing data, and the estimation error, all based on 100 data replications.

a core tensor $\Theta \in \mathbb{R}^{5 \times 5 \times 5}$ ($p = 40, u = 5$). The latter has a higher predictor dimension and rank, and is comparable to the dimension of the ADHD real data example in Section 6.2. The sample size for training and testing data is still fixed at $n = 200$. Table S1 summarizes the prediction and estimation results based on 100 data replications. It is again clearly seen that the proposed tensor envelope PLS method is more competitive than the alternative solutions in terms of both prediction and estimation accuracy across all model scenarios.

References

Cook, R. D., Helland, I. S., and Su, Z. (2013). Envelopes and partial least squares regression. *J. R. Stat. Soc. Ser. B. Stat. Methodol.*, 75(5):851–877.

- Cook, R. D., Li, B., and Chiaromonte, F. (2010). Envelope models for parsimonious and efficient multivariate linear regression. *Statist. Sinica*, 20(3):927–960.
- Zhou, H., Li, L., and Zhu, H. (2013). Tensor regression with applications in neuroimaging data analysis. *Journal of the American Statistical Association*, 108(502):540–552.