

Local Interaction Autoregressive Model for High-Dimensional Time Series Data

Jingyang Li¹, Yang Chen²

¹*Fudan University*; ²*University of Michigan, Ann Arbor*

Supplementary Material

This supplementary material contains proofs of the main theorems, auxiliary lemmas, and additional methodological and numerical details. It is organized as follows. Section S1 presents the proofs of the main theoretical results in the order in which they appear in the main paper, including the coefficient estimation theory, the separable estimator, auto-covariance estimation, and neighborhood-selection consistency. Section S2 collects the auxiliary probabilistic and matrix-analytic lemmas used throughout the proofs. Section S3 provides additional methodological and numerical details, including the non-identifiability issue for non-nested neighborhoods, estimation and bandwidth selection for the general lag- P model, and supplementary simulation results. Section S4 describes the tensor extension of the proposed LIAR framework and presents the tensor TEC application, including the tensor neighborhood-selection results and prediction comparison.

S1 Proofs of Main Theorems

The auxiliary lemmas used in the proofs are collected in Section S2. We present the proofs of the main theorems in the same order as they appear in the main text.

S1.1 Proof of Theorem 1

The proof is the one-block version of the martingale central limit theorem argument used in the proof of Theorem 2. Since Theorem 1 also follows immediately from the joint result in Theorem 2, we present the short marginal derivation here for completeness.

By Theorem 2,

$$\sqrt{T} \text{vecstack}\{\widehat{\mathbf{m}}_{\mathbf{i}} - \mathbf{m}_{\mathbf{i}} : \mathbf{i} \in \mathcal{S}\} \Rightarrow N(0, \mathbf{D}^{-1} \mathbf{C} \mathbf{D}^{-1}).$$

Taking the block corresponding to a fixed location $\mathbf{i}$, the limiting covariance is

$$\mathbf{\Gamma}_0(\mathcal{J}_{\mathbf{i}}; \mathcal{J}_{\mathbf{i}})^{-1} \{[\mathbf{\Sigma}]_{\mathbf{i}, \mathbf{i}} \mathbf{\Gamma}_0(\mathcal{J}_{\mathbf{i}}; \mathcal{J}_{\mathbf{i}})\} \mathbf{\Gamma}_0(\mathcal{J}_{\mathbf{i}}; \mathcal{J}_{\mathbf{i}})^{-1} = \sigma_{\mathbf{i}}^2 \mathbf{\Gamma}_0(\mathcal{J}_{\mathbf{i}}; \mathcal{J}_{\mathbf{i}})^{-1},$$

where $[\mathbf{\Sigma}]_{\mathbf{i}, \mathbf{i}} = \text{Var}([\mathbf{E}_t]_{\mathbf{i}}) = \sigma_{\mathbf{i}}^2$. Hence

$$\sqrt{T}(\widehat{\mathbf{m}}_{\mathbf{i}} - \mathbf{m}_{\mathbf{i}}) \Rightarrow N(0, \sigma_{\mathbf{i}}^2 \mathbf{\Gamma}_0(\mathcal{J}_{\mathbf{i}}; \mathcal{J}_{\mathbf{i}})^{-1}),$$

which proves Theorem 1.

S1.2 Proof of Theorem 2

Proof. For each $\mathbf{i}$, $\widehat{\mathbf{m}}_{\mathbf{i}}$ admits the following closed-form solution:

$$\widehat{\mathbf{m}}_{\mathbf{i}} = (\mathbf{Y}_{\mathbf{i}}^{\top} \mathbf{Y}_{\mathbf{i}})^{-1} \mathbf{Y}_{\mathbf{i}}^{\top} \mathbf{z}_{\mathbf{i}} = \mathbf{m}_{\mathbf{i}} + (\mathbf{Y}_{\mathbf{i}}^{\top} \mathbf{Y}_{\mathbf{i}})^{-1} \mathbf{Y}_{\mathbf{i}}^{\top} \boldsymbol{\nu}_{\mathbf{i}}.$$

Therefore,

$$\text{vecstack}\{\widehat{\mathbf{m}}_{\mathbf{i}} - \mathbf{m}_{\mathbf{i}} : \mathbf{i} \in \mathcal{S}\} = \text{vecstack}\{(\mathbf{Y}_{\mathbf{i}}^{\top} \mathbf{Y}_{\mathbf{i}})^{-1} \mathbf{Y}_{\mathbf{i}}^{\top} \boldsymbol{\nu}_{\mathbf{i}} : \mathbf{i} \in \mathcal{S}\}.$$

Let $\boldsymbol{\gamma} = \text{vecstack}\{\boldsymbol{\gamma}_{\mathbf{i}} : \mathbf{i} \in \mathcal{S}\}$, then we have

$$\begin{aligned} T^{-1/2} \boldsymbol{\gamma}^{\top} \text{vecstack}\{\mathbf{Y}_{\mathbf{i}}^{\top} \boldsymbol{\nu}_{\mathbf{i}} : \mathbf{i} \in \mathcal{S}\} &= T^{-1/2} \sum_{\mathbf{i}} \boldsymbol{\gamma}_{\mathbf{i}}^{\top} \mathbf{Y}_{\mathbf{i}}^{\top} \boldsymbol{\nu}_{\mathbf{i}} \\ &= \sum_{t=2}^T T^{-1/2} \underbrace{\sum_{\mathbf{i}} \boldsymbol{\gamma}_{\mathbf{i}}^{\top} \mathbf{x}_{t-1}^{(i)} [\mathbf{E}_t]_{\mathbf{i}}}_{:= Z_{t,T}}. \end{aligned}$$

Let $\mathcal{A}_{t-1}$ denote the σ -field generated by previous innovations. Then

$$\mathbb{E}(Z_{t,T} \mid \mathcal{A}_{t-1}) = 0,$$

and

$$\mathbb{E}(Z_{t,T}^2 \mid \mathcal{A}_{t-1}) = T^{-1} \sum_{\mathbf{i}, \mathbf{j} \in \mathcal{S}} \boldsymbol{\Sigma}_{\mathbf{i}, \mathbf{j}} \boldsymbol{\gamma}_{\mathbf{i}}^{\top} \mathbf{x}_{t-1}^{(i)} \mathbf{x}_{t-1}^{(j)\top} \boldsymbol{\gamma}_{\mathbf{j}}.$$

Denote

$$V_{TT}^2 = T^{-1} \sum_{t=2}^T \sum_{\mathbf{i}, \mathbf{j} \in \mathcal{S}} \boldsymbol{\Sigma}_{\mathbf{i}, \mathbf{j}} \boldsymbol{\gamma}_{\mathbf{i}}^{\top} \mathbf{x}_{t-1}^{(i)} \mathbf{x}_{t-1}^{(j)\top} \boldsymbol{\gamma}_{\mathbf{j}}.$$

On the other hand, we have

$$s_{TT} := \mathbb{E} \left(\sum_{t=2}^T Z_{t,T} \right)^2 = T^{-1} (T-1) \sum_{\mathbf{i}, \mathbf{j}} \boldsymbol{\Sigma}_{\mathbf{i}, \mathbf{j}} \boldsymbol{\gamma}_{\mathbf{i}}^{\top} \boldsymbol{\Gamma}_0(\mathcal{J}_{\mathbf{i}}; \mathcal{J}_{\mathbf{j}}) \boldsymbol{\gamma}_{\mathbf{j}}.$$

Then from Lemma 6 (under Regime 1), or from Lemma 2 (under Regime 2 or 3), we have $V_{TT}s_{TT}^{-1} \xrightarrow{P} 1$. Finally, for all $\epsilon > 0$, we can show using the moment condition of $[\mathbf{E}_t]_i$,

$$\lim_{T \rightarrow \infty} s_{TT}^{-2} \sum_{t=2}^T \mathbb{E}(Z_{t,T}^2 \cdot 1(|Z_{t,T}| \geq \epsilon s_{TT} | \mathcal{A}_{t-1})) = 0.$$

So we conclude from Theorem 5.3.4 in Fuller (2009),

$$T^{-1/2} \text{vecstack}\{\mathbf{Y}_i^\top \boldsymbol{\nu}_i : i \in \mathcal{S}\} \implies N(0, \mathbf{C}).$$

From Lemma 6 (under Regime 1), or from Lemma 2 (under Regime 2 or 3), we have $\frac{1}{T} \mathbf{Y}_i^\top \mathbf{Y}_i \xrightarrow{P} \boldsymbol{\Gamma}_0(\mathcal{J}_i; \mathcal{J}_i)$. So we conclude

$$\sqrt{T} \text{vecstack}\{\widehat{\mathbf{m}}_i - \mathbf{m}_i : i \in \mathcal{S}\} \Rightarrow N(0, \mathbf{D}^{-1} \mathbf{C} \mathbf{D}^{-1}).$$

where $\mathbf{D} = \text{blkdiag}\{\boldsymbol{\Gamma}_0(\mathcal{J}_i; \mathcal{J}_i) : i \in \mathcal{S}\}$. □

S1.3 Proof of Theorem 3

The proof of this theorem relies on Lemma 5. Since $\widehat{\mathbf{M}}_R^{\text{blk}}$ is the best rank R approximation of $\widehat{\mathbf{M}}^{\text{blk}}$. And we have $\|\widehat{\mathbf{M}}^{\text{blk}} - \mathbf{M}^{\text{blk}}\|_F \rightarrow 0$ as $T \rightarrow \infty$, so the condition in Lemma 5 is satisfied for sufficiently large T . So we conclude

$$\begin{aligned} \widehat{\mathbf{M}}_R^{\text{blk}} - \mathbf{M}^{\text{blk}} &= \widehat{\mathbf{M}}^{\text{blk}} - \mathbf{M}^{\text{blk}} - \mathbf{U}_\perp^{\text{blk}} \mathbf{U}_\perp^{\text{blk}\top} (\widehat{\mathbf{M}}^{\text{blk}} - \mathbf{M}^{\text{blk}}) \\ &\quad \times \mathbf{V}_\perp^{\text{blk}} \mathbf{V}_\perp^{\text{blk}\top} + \mathbf{R}, \end{aligned}$$

where $\|\mathbf{R}\|_F \leq \frac{40\|\widehat{\mathbf{M}}^{\text{blk}} - \mathbf{M}^{\text{blk}}\|_F^2}{\lambda_{\min}(\mathbf{M}^{\text{blk}})}$, and thus $\sqrt{T}\|\mathbf{R}\|_F = o_p(1)$. Next we vectorize both sides and we obtain

$$\begin{aligned} \text{vec}(\widehat{\mathbf{M}}_R^{\text{blk}} - \mathbf{M}^{\text{blk}}) &= (\mathbf{I} - \mathbf{V}_{\perp}^{\text{blk}} \mathbf{V}_{\perp}^{\text{blk}\top} \otimes \mathbf{U}_{\perp}^{\text{blk}} \mathbf{U}_{\perp}^{\text{blk}\top}) \\ &\quad \times \text{vec}(\widehat{\mathbf{M}}^{\text{blk}} - \mathbf{M}^{\text{blk}}) + \text{vec}(\mathbf{R}). \end{aligned}$$

Together with Theorem 2, we obtain the result.

S1.4 Proof of Theorem 4

Proof. The result is an immediate consequence of the first part of Lemma 2.

Indeed, Lemma 2 implies that, under Assumption 1,

$$\max_{i,j \in \mathcal{S}} \left| [\widehat{\mathbf{\Gamma}}_0]_{i,j} - [\mathbf{\Gamma}_0]_{i,j} \right| = O_p \left\{ \left(\frac{\log(M \vee N \vee T)}{T} \right)^{1/2} \mu_{2q}^2 J_{\square} (1 - \delta)^{-4} \right\}.$$

Since μ_{2q} and δ are treated as constants, and since

$$\|\widehat{\mathbf{\Gamma}}_0 - \mathbf{\Gamma}_0\|_{\ell_{\infty}} = \max_{i,j \in \mathcal{S}} \left| [\widehat{\mathbf{\Gamma}}_0]_{i,j} - [\mathbf{\Gamma}_0]_{i,j} \right|,$$

we obtain

$$\|\widehat{\mathbf{\Gamma}}_0 - \mathbf{\Gamma}_0\|_{\ell_{\infty}} = O_p \left\{ \left(\frac{\log(M \vee N \vee T)}{T} \right)^{1/2} J_{\square} \right\}.$$

This proves Theorem 4. □

S1.5 Proof of Theorem 5

Proof. For each $\mathbf{i} \in \mathcal{S}$, recall we have $\mathcal{J}_i[k_0] = \mathcal{J}_i$. And $\{\widehat{k}_i \neq k_0\} = \{\widehat{k}_i < k_0\} \cup \{\widehat{k}_i > k_0\}$. For $k < k_0$, we have

$$\begin{aligned} \text{RSS}_i(k) &= \mathbf{z}_i^\top (\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) \mathbf{z}_i \\ &= (\mathbf{Y}_i[k_0] \mathbf{m}_i + \boldsymbol{\nu}_i)^\top (\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) \\ &\quad \times (\mathbf{Y}_i[k_0] \mathbf{m}_i + \boldsymbol{\nu}_i). \end{aligned}$$

Since $k < k_0$, $\mathbf{Y}_i[k]$ is a submatrix of $\mathbf{Y}_i[k_0]$, and we can split $\mathbf{Y}_i[k_0] \mathbf{m}_i$ as

$$\mathbf{Y}_i[k_0] \mathbf{m}_i = \mathbf{Y}_i[k] \mathbf{b}_1 + \mathbf{S} \mathbf{b}_2,$$

where $\|\mathbf{b}_2\|_{\ell_2} = \|\mathbf{M}_i(\mathcal{J}_i \setminus \mathcal{J}_i[k])\|_{\text{F}}$. Let $\mathcal{I}_1 := \mathcal{J}_i[k]$ and $\mathcal{I}_2 := \mathcal{J}_i \setminus \mathcal{J}_i[k]$.

Define

$$\widehat{\mathbf{G}} := \frac{1}{T-1} \begin{bmatrix} \mathbf{Y}_i[k]^\top \mathbf{Y}_i[k] & \mathbf{Y}_i[k]^\top \mathbf{S} \\ \mathbf{S}^\top \mathbf{Y}_i[k] & \mathbf{S}^\top \mathbf{S} \end{bmatrix} := \begin{bmatrix} \widehat{\boldsymbol{\Gamma}}_0(\mathcal{I}_1; \mathcal{I}_1) & \widehat{\boldsymbol{\Gamma}}_0(\mathcal{I}_1; \mathcal{I}_2) \\ \widehat{\boldsymbol{\Gamma}}_0(\mathcal{I}_2; \mathcal{I}_1) & \widehat{\boldsymbol{\Gamma}}_0(\mathcal{I}_2; \mathcal{I}_2) \end{bmatrix}.$$

Then we denote

$$\mathbf{G} := \mathbb{E} \widehat{\mathbf{G}} = \begin{bmatrix} \boldsymbol{\Gamma}_0(\mathcal{I}_1; \mathcal{I}_1) & \boldsymbol{\Gamma}_0(\mathcal{I}_1; \mathcal{I}_2) \\ \boldsymbol{\Gamma}_0(\mathcal{I}_2; \mathcal{I}_1) & \boldsymbol{\Gamma}_0(\mathcal{I}_2; \mathcal{I}_2) \end{bmatrix}.$$

Since $(\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) \mathbf{Y}_i[k_0] = 0$, we have

$$\text{RSS}_i(k) = (\mathbf{S} \mathbf{b}_2 + \boldsymbol{\nu}_i)^\top (\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) (\mathbf{S} \mathbf{b}_2 + \boldsymbol{\nu}_i).$$

On the other hand, we have

$$\text{RSS}_i(k_0) = \boldsymbol{\nu}_i^\top (\mathbf{I} - \mathbf{Y}_i[k_0](\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top) \boldsymbol{\nu}_i.$$

Therefore,

$$\begin{aligned} \text{RSS}_i(k) - \text{RSS}_i(k_0) &\geq \mathbf{b}_2^\top \mathbf{S}^\top (\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) \mathbf{S} \mathbf{b}_2 \\ &\quad + 2\mathbf{b}_2^\top \mathbf{S}^\top (\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) \boldsymbol{\nu}_i. \end{aligned}$$

We denote $\widehat{\mathbf{M}} = \widehat{\mathbf{G}}^{-1}$, $\mathbf{M} = \mathbf{G}^{-1}$. And $\widehat{\mathbf{M}}$ is partitioned with the same shape as $\widehat{\mathbf{G}}$, $\widehat{\mathbf{M}} = \begin{bmatrix} \widehat{\mathbf{M}}_{1,1} & \widehat{\mathbf{M}}_{1,2} \\ \widehat{\mathbf{M}}_{2,1} & \widehat{\mathbf{M}}_{2,2} \end{bmatrix}$. For the first term, we have

$$\mathbf{S}^\top (\mathbf{I} - \mathbf{Y}_{ij}[k](\mathbf{Y}_{ij}[k]^\top \mathbf{Y}_{ij}[k])^{-1} \mathbf{Y}_{ij}[k]^\top) \mathbf{S} = (T-1)(\widehat{\mathbf{M}}_{2,2})^{-1}.$$

And

$$\|\widehat{\mathbf{M}} - \mathbf{M}\| = \|\widehat{\mathbf{G}}^{-1} - \mathbf{G}^{-1}\| \leq \|\widehat{\mathbf{G}}^{-1}\| \cdot \|\widehat{\mathbf{G}} - \mathbf{G}\| \cdot \|\mathbf{G}^{-1}\|. \quad (\text{S1.1})$$

Notice from Lemma 2, with probability tending to 1,

$$\begin{aligned} \|\widehat{\mathbf{G}} - \mathbf{G}\| &\leq |\mathcal{J}_i| \cdot \|\widehat{\mathbf{G}} - \mathbf{G}\|_{\ell_\infty} \\ &= O_p((T^{-1} \log(M \vee N \vee T))^{1/2} |\mathcal{J}_i| J_\square \mu_{2q}^2 (1-\delta)^{-4}). \end{aligned}$$

Since $\mathbf{G}$ is a submatrix of $\mathbf{\Gamma}_0$, we have $\lambda_{\min}(\mathbf{G}) \geq \lambda_{\min}(\mathbf{\Gamma}_0) \geq \kappa_1$. Also, as long as $(T^{-1} \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_\square (1-\delta)^{-4} \lesssim \kappa_1$, we have $\lambda_{\min}(\widehat{\mathbf{G}}) \geq$

$\lambda_{\min}(\mathbf{G}) - \frac{\kappa_1}{2} \geq \frac{\kappa_1}{2}$. Therefore $\|\widehat{\mathbf{M}} - \mathbf{M}\| \leq c_0 \kappa_1^{-1}$ for some $c_0 \in (0, 1)$. So

$$\begin{aligned} & \lambda_{\min}\{\mathbf{S}^\top (\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) \mathbf{S}\} \\ &= T \lambda_{\min}((\widehat{\mathbf{M}}_{2,2})^{-1}) = T \lambda_{\max}^{-1}(\widehat{\mathbf{M}}_{2,2}). \end{aligned}$$

Since $\widehat{\mathbf{M}}_{2,2}$ is a submatrix of $\widehat{\mathbf{M}}$, we have $\lambda_{\max}(\widehat{\mathbf{M}}_{2,2}) \leq \lambda_{\max}(\widehat{\mathbf{M}}) \leq \lambda_{\max}(\mathbf{M}) + \|\widehat{\mathbf{M}} - \mathbf{M}\| \leq \kappa_1^{-1} + c_0 \kappa_1^{-1}$. And thus

$$\lambda_{\min}(\mathbf{S}^\top (\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) \mathbf{S}) \geq (1 + c_0)^{-1} T \kappa_1,$$

which also implies

$$\mathbf{b}_2^\top \mathbf{S}^\top (\mathbf{I} - \mathbf{Y}_i[k](\mathbf{Y}_i[k]^\top \mathbf{Y}_i[k])^{-1} \mathbf{Y}_i[k]^\top) \mathbf{S} \mathbf{b}_2 \geq (1 + c_0)^{-1} T \kappa_1 \|\mathbf{b}_2\|_{\ell_2}^2.$$

On the other hand, we have from Cauchy-Schwarz inequality and second part in Lemma 2,

$$\begin{aligned} & 2\mathbf{b}_2^\top \mathbf{S}^\top (\mathbf{I} - \mathbf{Y}_{ij;k}(\mathbf{Y}_{ij;k}^\top \mathbf{Y}_{ij;k})^{-1} \mathbf{Y}_{ij;k}^\top) \boldsymbol{\epsilon}_{ij} \\ & \leq \frac{1}{2} (1 + c_0)^{-1} T \kappa_1 \|\mathbf{b}_2\|_{\ell_2}^2 + C_1 \kappa_1^{-1} \mu_{2q}^4 (1 - \delta)^{-4} J_\square^2 \log(M \vee N \vee T). \end{aligned}$$

From Lemma 3, we have $\text{RSS}_i(k_0) \lesssim T \mu_{2q}^2$ with probability tending to 1.

Therefore, for sufficiently large T , we have

$$\begin{aligned} \text{RSS}_i(k) - \text{RSS}_i(k_0) & \geq \frac{1}{2} (1 + c_0)^{-1} T \kappa_1 \|\mathbf{b}_2\|_{\ell_2}^2 \\ & \quad - C_1 \kappa_1^{-1} \mu_{2q}^4 (1 - \delta)^{-4} J_\square^2 \log(M \vee N \vee T) \\ & \geq c_1 \text{RSS}_i(k_0) \frac{\kappa_1}{\mu_{2q}^2} \|\mathbf{b}_2\|_{\ell_2}^2 \end{aligned}$$

for some $c_1 \in (0, 1)$. And using $\log(1+x) \geq \frac{1}{2}x$ for $x \in (0, 2)$, we have

$$\log \text{RSS}_i(k) - \log \text{RSS}_i(k_0) \geq \min\{c_2 \frac{\kappa_1}{\mu_{2q}^2} \|\mathbf{b}_2\|_{\ell_2}^2, \log 3\}.$$

Therefore for sufficiently large T , we have

$$\begin{aligned} \text{BIC}_i(k) - \text{BIC}_i(k_0) &\geq \min\{c_2 \frac{\kappa_1}{\mu_{2q}^2} \|\mathbf{b}_2\|_{\ell_2}^2, \log 3\} \\ &\quad - D_0 \frac{|\mathcal{J}_i|}{T} \log(M \vee N \vee T) > 0, \end{aligned}$$

under the margin condition that

$$\|\mathbf{b}_2\|_{\ell_2} = \|\mathbf{M}_i(\mathcal{J}_i \setminus \mathcal{J}_i[k])\|_F \geq \|\mathbf{M}_i(\mathcal{J}_i \setminus \mathcal{J}_i[k_0 - 1])\|_F.$$

Now we show, for all $\mathbf{i} \in \mathcal{S}$ and $k > k_0$, $\mathbb{P}(\text{BIC}_i(k) < \text{BIC}_i(k_0)) \rightarrow 0$.

We have

$$\text{RSS}_i(k) = \inf_{\boldsymbol{\gamma}} \|\mathbf{z}_i - \mathbf{Y}_i[k]\boldsymbol{\gamma}\|_{\ell_2}^2.$$

We can write $\mathbf{Y}_i[k]\boldsymbol{\gamma} = \mathbf{Y}_i[k_0]\boldsymbol{\gamma}_1 + \mathbf{S}\boldsymbol{\gamma}_2$. Denote

$$\tilde{\mathbf{S}} = (I - \mathbf{Y}_i[k_0](\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top) \mathbf{S}.$$

Then

$$\begin{aligned} \text{RSS}_i(k) &= \inf_{\boldsymbol{\gamma}} \|\mathbf{z}_i - \mathbf{Y}_i[k]\boldsymbol{\gamma}\|_{\ell_2}^2 = \inf_{\boldsymbol{\gamma}_1, \boldsymbol{\gamma}_2} \|\mathbf{z}_i - \mathbf{Y}_i[k_0]\boldsymbol{\gamma}_1 - \mathbf{S}\boldsymbol{\gamma}_2\|_{\ell_2}^2 \\ &= \inf_{\boldsymbol{\gamma}_1, \boldsymbol{\gamma}_2} \|\mathbf{z}_i - \mathbf{Y}_i[k_0] \left(\boldsymbol{\gamma}_1 + (\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top \boldsymbol{\gamma}_2 \right) - \tilde{\mathbf{S}}\boldsymbol{\gamma}_2\|_{\ell_2}^2 \\ &= \inf_{\boldsymbol{\gamma}_1, \boldsymbol{\gamma}_2} \|\mathbf{z}_i - \mathbf{Y}_i[k_0]\boldsymbol{\gamma}_1 - \tilde{\mathbf{S}}\boldsymbol{\gamma}_2\|_{\ell_2}^2. \end{aligned}$$

An explicit minimizer is

$$\gamma_1 = (\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top \mathbf{z}_i, \quad \gamma_2 = (\tilde{\mathbf{S}}^\top \tilde{\mathbf{S}})^{-1} \tilde{\mathbf{S}}^\top \mathbf{z}_i.$$

Therefore,

$$\text{RSS}_i(k) = \text{RSS}_i(k_0) - \|\tilde{\mathbf{S}}(\tilde{\mathbf{S}}^\top \tilde{\mathbf{S}})^{-1} \tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i\|_{\ell_2}^2.$$

Next we bound $\|\tilde{\mathbf{S}}(\tilde{\mathbf{S}}^\top \tilde{\mathbf{S}})^{-1} \tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i\|_{\ell_2}^2$:

$$\|\tilde{\mathbf{S}}(\tilde{\mathbf{S}}^\top \tilde{\mathbf{S}})^{-1} \tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i\|_{\ell_2}^2 = \boldsymbol{\nu}_i^\top \tilde{\mathbf{S}}(\tilde{\mathbf{S}}^\top \tilde{\mathbf{S}})^{-1} \tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i \leq \lambda_{\max}((\tilde{\mathbf{S}}^\top \tilde{\mathbf{S}})^{-1}) \|\tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i\|_{\ell_2}^2.$$

Since $\tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i \in \mathbb{R}^{|\mathcal{J}_i[k]| - |\mathcal{J}_i|}$, we have

$$\|\tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i\|_{\ell_2}^2 \leq (|\mathcal{J}_i[k]| - |\mathcal{J}_i|) \|\tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i\|_{\ell_\infty}^2.$$

Recall $\tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i = \mathbf{S}^\top (\mathbf{I} - \mathbf{Y}_i[k_0](\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top) \boldsymbol{\nu}_i$. Let an arbitrary row of $\mathbf{S}^\top$ be $\mathbf{x}^\top$. Then

$$\begin{aligned} & \mathbf{x}^\top (\mathbf{I} - \mathbf{Y}_i[k_0](\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top) \boldsymbol{\nu}_i \\ &= \mathbf{x}^\top \boldsymbol{\nu}_i - \mathbf{x}^\top \mathbf{Y}_i[k_0](\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top \boldsymbol{\nu}_i. \end{aligned} \quad (\text{S1.2})$$

From Lemma 2, $\mathbf{x}^\top \boldsymbol{\nu}_i = O_p((T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_\square^{1/2} (1 - \delta)^{-3/2})$.

Moreover,

$$\begin{aligned} & \mathbf{x}^\top \mathbf{Y}_i[k_0](\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top \boldsymbol{\nu}_i \\ &= \mathbf{x}^\top \mathbf{Y}_i[k_0] \left((\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} - (\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \right) \mathbf{Y}_i[k_0]^\top \boldsymbol{\nu}_i \\ & \quad + \mathbf{x}^\top \mathbf{Y}_i[k_0] (\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top \boldsymbol{\nu}_i. \end{aligned}$$

Since $\frac{1}{T}\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0]$ is a submatrix of $\mathbf{\Gamma}_0$, we have

$$\lambda_{\max}((\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1}) = \lambda_{\min}^{-1}(\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0]) \leq T^{-1} \kappa_1^{-1}.$$

Therefore, from Lemma 2 and $\max_i [\mathbf{\Gamma}_0]_{i,i} \leq \|[\mathbf{E}_t]_i\|_{2q}^2 \leq \mu_{2q}^2$, we have

$$\begin{aligned} & |\mathbf{x}^\top \mathbf{Y}_i[k_0] (\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \mathbf{Y}_i[k_0]^\top \boldsymbol{\nu}_i| \\ & \lesssim (T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 |\mathcal{J}_i| J_\square^{1/2} (1 - \delta)^{-2} \kappa_1^{-1} \mu_{2q}^2. \end{aligned}$$

Next we bound $\|(\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} - (\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1}\|$ similarly as in

(S1.1):

$$\begin{aligned} & \|(\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} - (\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1}\| \\ & \lesssim T^{-1} (T^{-1} \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 |\mathcal{J}_i| J_\square (1 - \delta)^{-4} \kappa_1^{-2}. \end{aligned}$$

Thus

$$\begin{aligned} & \mathbf{x}^\top \mathbf{Y}_i[k_0] \left((\mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} - (\mathbb{E} \mathbf{Y}_i[k_0]^\top \mathbf{Y}_i[k_0])^{-1} \right) \mathbf{Y}_i[k_0]^\top \boldsymbol{\nu}_i \\ & \lesssim T^{-1} (T^{-1} \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 |\mathcal{J}_i| J_\square (1 - \delta)^{-3} \kappa_1^{-2} \cdot |\mathcal{J}_i| \cdot T \mu_{2q}^2 \\ & \quad \cdot (T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_\square^{1/2} (1 - \delta)^{-2} \\ & = J_\square^{3/2} |\mathcal{J}_i|^2 \log(M \vee N \vee T) \mu_{2q}^6 \kappa_1^{-2} (1 - \delta)^{-6}. \end{aligned}$$

Returning to (S1.2), we conclude

$$\|\tilde{\mathbf{S}}^\top \boldsymbol{\nu}_i\|_{\ell_\infty} = O_p((T \log(M \vee N \vee T))^{1/2} \mu_{2q}^4 |\mathcal{J}_i| J_\square^{1/2} (1 - \delta)^{-2} \kappa_1^{-1}).$$

Hence

$$\begin{aligned} \text{RSS}_i(k_0) - \text{RSS}_i(k) &\lesssim (|\mathcal{J}_i[k]| - |\mathcal{J}_i|) \log(M \vee N \vee T) \mu_{2q}^8 |\mathcal{J}_i|^2 \\ &\quad \times J_\square (1 - \delta)^{-4} \kappa_1^{-3}. \end{aligned}$$

From Lemma 3, $\text{RSS}_i(k) \gtrsim T \mu_{2q}^2$. Therefore

$$\begin{aligned} \frac{\text{RSS}_i(k_0) - \text{RSS}_i(k)}{\text{RSS}_i(k)} &\lesssim \frac{|\mathcal{J}_i[k]| - |\mathcal{J}_i|}{T} \log(M \vee N \vee T) \mu_{2q}^6 |\mathcal{J}_i|^2 \\ &\quad \times J_\square (1 - \delta)^{-4} \kappa_1^{-3}. \end{aligned}$$

Using $\log(1 + x) \leq x$ for $x > 0$, we have

$$\begin{aligned} &\log \text{RSS}_i(k_0) - \log \text{RSS}_i(k) \\ &\lesssim \frac{|\mathcal{J}_i[k]| - |\mathcal{J}_i|}{T} \log(M \vee N \vee T) \mu_{2q}^6 |\mathcal{J}_i|^2 \\ &\quad \times J_\square (1 - \delta)^{-4} \kappa_1^{-3}. \end{aligned}$$

Therefore

$$\begin{aligned} \text{BIC}_i(k) - \text{BIC}_i(k_0) &\geq D_0 \frac{|\mathcal{J}_i[k]| - |\mathcal{J}_i|}{T} \log(M \vee N \vee T) \\ &\quad + \log \text{RSS}_i(k) - \log \text{RSS}_i(k_0) \\ &> 0, \end{aligned}$$

as long as we choose $D_0 \gtrsim \mu_{2q}^6 |\mathcal{J}_i|^2 J_\square (1 - \delta)^{-4} \kappa_1^{-3}$.

□

S2 Auxiliary Lemmas

The auxiliary results in this section are used to support the proofs in Section S1. Lemma 1 provides the functional-dependence concentration inequality that is used as the main probabilistic tool. Lemma 2 applies this concentration result to obtain entrywise bounds for the sample autocovariance matrix and the noise-design cross terms; these bounds are used in the proofs of Theorems 2, 4, and 5. Lemma 3 controls the residual sum of squares over candidate neighborhoods and is used in the proof of Theorem 5. Lemma 4 records a locality-induced sparsity property of powers of the autoregressive coefficient matrix, which is used in the proof of Lemma 2. Lemma 5 is the perturbation result for the best rank- R projection and is used in the proof of Theorem 3. Finally, Lemma 6 provides the fixed-dimensional law of large numbers used under Regime 1 in the proof of Theorem 2.

We start with some notations. Let $\mathbf{\Gamma} \in \mathbb{R}^{MN \times MN}$. For multi-indices $\mathbf{i} = (i_1, i_2), \mathbf{j} = (j_1, j_2) \in \mathcal{S} = [M] \times [N]$, define

$$[\mathbf{\Gamma}]_{\mathbf{i}, \mathbf{j}} := [\mathbf{\Gamma}]_{(i_2-1)M+i_1, (j_2-1)M+j_1}, \quad (\text{S2.3})$$

where the last equality uses column-major linear indexing. The sub-matrix $\mathbf{\Gamma}(\mathcal{J}_1; \mathcal{J}_2)$ is then formed by extracting the elements of $\mathbf{\Gamma}$ corresponding to

the multi-indices in $\mathcal{J}_1$ and $\mathcal{J}_2$, arranged using column-major ordering.

Following Wu (2005), we introduce the following functional dependent measure. Let z_i be a stationary process of the form $z_i = g(\mathcal{F}_i)$, where g is a measurable function and $\mathcal{F}_i = (\cdots, e_{-1}, e_0, e_1, \cdots)$ with independent and identically distributed random variables $\{e_i\}$. The functional dependent measure is defined as

$$\theta_{i,q} = \|z_i - z_i^*\|_q = \|g(\mathcal{F}_i) - g(\mathcal{F}_i^*)\|_q,$$

where $z_i^* = g(\mathcal{F}_i^*)$ is the coupled process of z_i , $\mathcal{F}_i^* = (\cdots, e_{-1}, e_0^*, e_1, \cdots)$ with $\{e_0, e_0^*\}$ being independent and identically distributed. Here $\|\cdot\|_q := (\mathbb{E}|\cdot|)^{1/q}$ for some $q \geq 1$. Then we have the following from Theorem 2 in Liu et al. (2013) that plays the crucial role in our proof.

Lemma 1. *Let $S_n = \sum_{i=1}^n z_n$ and $\Theta_{m,q} = \sum_{i=m}^\infty \theta_{i,q}$. Assume for each m , $\Theta_{m,q} = O(m^{-\alpha})$ with $\alpha > \frac{1}{2} - \frac{1}{q}$ and $q > 2$. Then we have for all $x > 0$,*

$$\mathbb{P}(|S_n| \geq x) \leq C_1 \frac{\Theta_{0,q}^q n}{x^q} + C_3 \exp(-C_2 \frac{x^2}{\Theta_{0,q}^2 n}),$$

for $C_1, C_2, C_3 > 0$ depending only on q .

Recall $[\widehat{\mathbf{\Gamma}}_0]_{i,j} = \frac{1}{T} \sum_{t=1}^T [\mathbf{X}_t]_i [\mathbf{X}_t]_j$ and thus $\mathbf{\Gamma}_0 = \mathbb{E}\widehat{\mathbf{\Gamma}}_0$. We denote

$$\mathbf{x}_i = ([\mathbf{X}_1]_i, \cdots, [\mathbf{X}_{T-1}]_i)^\top,$$

and recall $\boldsymbol{\nu}_i = ([\mathbf{E}_2]_i, \cdots, [\mathbf{E}_T]_i)^\top$. We first establish the following concentration inequality.

Lemma 2. *Under Assumption 1, there exists absolute constant $C > 0$, such that*

1. *For $\mathbf{i}, \mathbf{j} \in \mathcal{S}$, we have for $x > 0$,*

$$\begin{aligned} \mathbb{P}\left(|[\widehat{\mathbf{\Gamma}}_0]_{\mathbf{i}, \mathbf{j}} - [\mathbf{\Gamma}_0]_{\mathbf{i}, \mathbf{j}}| \geq x\right) &\lesssim CT \left(\frac{\mu_{2q}^2 J_{\square}(1-\delta)^{-4}}{Tx}\right)^q \\ &\quad + C \exp\left(-\frac{Tx^2}{(\mu_{2q}^2 J_{\square}(1-\delta)^{-4})^2}\right). \end{aligned}$$

And this leads to

$$\max_{\mathbf{i}, \mathbf{j} \in \mathcal{S}} |[\widehat{\mathbf{\Gamma}}_0]_{\mathbf{i}, \mathbf{j}} - [\mathbf{\Gamma}_0]_{\mathbf{i}, \mathbf{j}}| = O_p((T^{-1} \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_{\square}(1-\delta)^{-4}).$$

2. *For $\mathbf{i}, \mathbf{j} \in \mathcal{S}$, we have for $x > 0$,*

$$\begin{aligned} \mathbb{P}(|\boldsymbol{\nu}_{\mathbf{i}}^{\top} \mathbf{x}_{\mathbf{j}}| \geq x) &\leq CT \left(\frac{\mu_{2q}^2 J_{\square}^{1/2}(1-\delta)^{-2}}{x}\right)^q \\ &\quad + C \exp\left(-\frac{x^2}{T \cdot (\mu_{2q}^2 J_{\square}^{1/2}(1-\delta)^{-2})^2}\right). \end{aligned}$$

And this leads to

$$\max_{\mathbf{i}, \mathbf{j} \in \mathcal{S}} |\boldsymbol{\nu}_{\mathbf{i}}^{\top} \mathbf{x}_{\mathbf{j}}| = O_p((T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_{\square}^{1/2}(1-\delta)^{-2}).$$

Proof. We shall use Lemma 1 to prove this. Denote $\mu_q = \max_{\mathbf{i}} \|[\mathbf{E}_0]_{\mathbf{i}}\|_q$.

Using innovation representation, we have

$$\mathbf{x}_t = \sum_{l=0}^{\infty} \mathbf{M}^l \boldsymbol{\epsilon}_{t-l}.$$

As a result, for each $\mathbf{i} = (i_1, i_2) \in \mathcal{S}$,

$$\begin{aligned} (\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{x}_t &= \sum_{l=0}^{\infty} (\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{M}^l \boldsymbol{\epsilon}_{t-l} \\ &= \sum_{l=0}^{\infty} \sum_j (\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{M}^l (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1}) (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1})^\top \boldsymbol{\epsilon}_{t-l}. \end{aligned}$$

Following the definition of $J_\square$ and Lemma 4, we have $\|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{M}^l\|_{\ell_0} \leq \min\{J_\square l^2, MN\}$. And thus

$$\begin{aligned} \max_{\mathbf{i}} \|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{x}_t\|_{2q} &\leq \max_{\mathbf{i}} \left\| \sum_{l=0}^{\infty} \sum_j (\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{M}^l (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1}) \right. \\ &\quad \left. \times (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1})^\top \boldsymbol{\epsilon}_{t-l} \right\|_{2q} \\ &\leq \max_{\mathbf{i}} \sum_{l=0}^{\infty} \sum_j |[\mathbf{M}^l]_{\mathbf{i}, \mathbf{j}}| \cdot \|[\boldsymbol{\epsilon}_{t-l}]_{\mathbf{j}}\|_{2q} \\ &\leq \max_{\mathbf{i}} \sum_{l=0}^{\infty} J_\square^{1/2} l \|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{M}^l\|_{\ell_2} \cdot \mu_{2q} \\ &\leq J_\square^{1/2} \mu_{2q} \sum_{l \geq 0} l \delta^l \leq C J_\square^{1/2} \mu_{2q} (1 - \delta)^{-2} \\ &< +\infty, \end{aligned} \tag{S2.4}$$

where in the second-to-last line, we use the fact $\|\mathbf{M}^l\| \leq \delta^l$.

On the other hand, we can bound $\|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{x}_t - (\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{x}_t^*\|_{2q}$

as follows using the innovation representation:

$$\begin{aligned}
& \|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{x}_t - (\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{x}_t^*\|_{2q} \\
&= \left\| \sum_j (\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{M}^t (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1}) (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1})^\top (\boldsymbol{\epsilon}_0 - \boldsymbol{\epsilon}_0^*) \right\|_{2q} \\
&\leq C \mu_{2q} J_\square^{1/2} t \delta^t.
\end{aligned} \tag{S2.5}$$

Now we can bound $\|[\mathbf{X}_t]_i [\mathbf{X}_t]_j - [\mathbf{X}_t]_i^* [\mathbf{X}_t]_j^*\|_q$ using (S2.4) and (S2.5):

$$\begin{aligned}
& \|[\mathbf{X}_t]_i [\mathbf{X}_t]_j - [\mathbf{X}_t]_i^* [\mathbf{X}_t]_j^*\|_q \\
&\leq \|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{x}_t \cdot (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1})^\top (\mathbf{x}_t - \mathbf{x}_t^*)\|_q \\
&\quad + \|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top (\mathbf{x}_t - \mathbf{x}_t^*) \cdot (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1})^\top \mathbf{x}_t^*\|_q \\
&\leq \|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{x}_t\|_{2q} \cdot \|(\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1})^\top (\mathbf{x}_t - \mathbf{x}_t^*)\|_{2q} \\
&\quad + \|(\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1})^\top \mathbf{x}_t^*\|_{2q} \cdot \|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top (\mathbf{x}_t - \mathbf{x}_t^*)\|_{2q} \\
&\leq C \mu_{2q}^2 J_\square (1 - \delta)^{-2} t \delta^t.
\end{aligned}$$

Therefore

$$\begin{aligned}
\Theta_{m,q} &:= \sum_{t=m}^{\infty} \|[\mathbf{X}_t]_i [\mathbf{X}_t]_j - [\mathbf{X}_t]_i^* [\mathbf{X}_t]_j^*\|_q \\
&\leq C \mu_{2q}^2 J_\square (1 - \delta)^{-4} \delta^m = O(m^{-\alpha})
\end{aligned}$$

for any $\alpha > 1$.

From Lemma 1, we conclude

$$\begin{aligned} \mathbb{P}\left(\left|\widehat{\Gamma}_0\right]_{i,j} - \left[\Gamma_0\right]_{i,j}\right| \geq x\right) &\lesssim T\left(\frac{\mu_{2q}^2 J_{\square}(1-\delta)^{-4}}{Tx}\right)^q \\ &\quad + \exp\left(-\frac{Tx^2}{(\mu_{2q}^2 J_{\square}(1-\delta)^{-4})^2}\right). \end{aligned}$$

By setting $x = T^{-1/2} \mu_{2q}^2 J_{\square}(1-\delta)^{-4} (\log(M \vee N \vee T))^{1/2}$, we conclude

$$\max_{i,j \in \mathcal{S}} \left| \widehat{\Gamma}_0\right]_{i,j} - \left[\Gamma_0\right]_{i,j} \right| = O_p\left(T^{-1/2} \mu_{2q}^2 J_{\square}(1-\delta)^{-4} (\log(M \vee N \vee T))^{1/2}\right).$$

Similarly, we can show

$$\begin{aligned} \mathbb{P}\left(\left|\boldsymbol{\nu}_i^{\top} \mathbf{x}_j\right| \geq x\right) &\lesssim T\left(\frac{\mu_{2q}^2 J_{\square}^{1/2}(1-\delta)^{-2}}{x}\right)^q \\ &\quad + \exp\left(-\frac{x^2}{T \cdot (\mu_{2q}^2 J_{\square}^{1/2}(1-\delta)^{-2})^2}\right), \end{aligned}$$

and

$$\max_{i,j \in \mathcal{S}} \left|\boldsymbol{\nu}_i^{\top} \mathbf{x}_j\right| = O_p\left((T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_{\square}^{1/2}(1-\delta)^{-2}\right).$$

□

Lemma 3. *Under Assumption 1, we have for each $\mathbf{i}$ and $k \geq k_0$,*

$$\max_{\mathbf{i}} \left| \text{RSS}_{\mathbf{i}}(k) - (T-1) \boldsymbol{\Sigma}_{\mathbf{i},\mathbf{i}} \right| = O_p\left((T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2\right)$$

holds as $T \rightarrow \infty$ and $\frac{\max_{\mathbf{i}} \|\mathcal{J}_{\mathbf{i}}[K_0]\|^2 \cdot J_{\square}^2 \log(M \vee N \vee T)}{T} \leq \tilde{c}$ for some $\tilde{c} > 0$ depending only on $\kappa_1, \delta, \mu_{2q}$.

Proof. When $k \geq k_0$, we can decompose $\text{RSS}_{\mathbf{i}}(k)$ as

$$\text{RSS}_{\mathbf{i}}(k) = \boldsymbol{\nu}_{\mathbf{i}}^{\top} \boldsymbol{\nu}_{\mathbf{i}} - \boldsymbol{\nu}_{\mathbf{i}}^{\top} \mathbf{Y}_{\mathbf{i}}[k] (\mathbf{Y}_{\mathbf{i}}[k]^{\top} \mathbf{Y}_{\mathbf{i}}[k])^{-1} \mathbf{Y}_{\mathbf{i}}[k]^{\top} \boldsymbol{\nu}_{\mathbf{i}}.$$

We consider the event

$$\mathcal{E}_1 = \left\{ \max_{i,j \in \mathcal{S}} |[\widehat{\mathbf{\Gamma}}_0]_{i,j} - [\mathbf{\Gamma}_0]_{i,j}| \lesssim (T^{-1} \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_{\square} (1 - \delta)^{-4}, \right. \\ \left. \max_{i,j \in \mathcal{S}} |\boldsymbol{\nu}_i^{\top} \mathbf{x}_j| \lesssim (T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_{\square}^{1/2} (1 - \delta)^{-2} \right\}.$$

From Lemma 2, $\mathbb{P}(\mathcal{E}_1) \rightarrow 0$. We shall continue our proof under $\mathcal{E}_1$. From

Lemma 1, we obtain

$$\max_i |\boldsymbol{\nu}_i^{\top} \boldsymbol{\nu}_i - (T - 1) \sigma_i^2| = O_p((T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2). \quad (\text{S2.6})$$

Recall $\mathcal{J}_i[k]$ is the index set corresponding with the columns of $\mathbf{Y}_{ij;k}$, then

for all i, j , under $\mathcal{E}_1$,

$$\begin{aligned} & (T - 1)^{-1} \lambda_{\min}(\mathbf{Y}_i[k]^{\top} \mathbf{Y}_i[k]) \\ & \geq \lambda_{\min}(\mathbf{\Gamma}_0(\mathcal{J}_i[k]; \mathcal{J}_i[k])) \\ & \quad - \|(T - 1)^{-1} \mathbf{Y}_i[k]^{\top} \mathbf{Y}_i[k] - \mathbf{\Gamma}_0(\mathcal{J}_i[k]; \mathcal{J}_i[k])\| \\ & \geq \kappa_1 - |\mathcal{J}_i[k]| (T^{-1} \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_{\square} (1 - \delta)^{-4} \\ & \geq \kappa_1 (1 - c_1), \end{aligned}$$

for some $c_1 \in (0, 1)$ as long as

$$|\mathcal{J}_i[k]| (T^{-1} \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_{\square} (1 - \delta)^{-4} \leq c_1 \kappa_1.$$

And under $\mathcal{E}_1$,

$$\|\mathbf{Y}_i[k_0]^{\top} \boldsymbol{\nu}_i\|_{\ell_{\infty}} \lesssim (T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2 J_{\square}^{1/2} (1 - \delta)^{-2}.$$

Therefore, with

$$\Pi_{\mathbf{i}, k_0} = \mathbf{Y}_{\mathbf{i}}[k_0](\mathbf{Y}_{\mathbf{i}}[k_0]^\top \mathbf{Y}_{\mathbf{i}}[k_0])^{-1} \mathbf{Y}_{\mathbf{i}}[k_0]^\top,$$

$$|\boldsymbol{\nu}_{\mathbf{i}}^\top \Pi_{\mathbf{i}, k_0} \boldsymbol{\nu}_{\mathbf{i}}| \lesssim |\mathcal{J}_{\mathbf{i}}| J_{\square} (1 - \delta)^{-4} \mu_{2q}^4 \kappa_1^{-1} \log(M \vee N \vee T).$$

Together with (S2.6), and

$$|\mathcal{J}_{\mathbf{i}}| J_{\square} (1 - \delta)^{-4} \mu_{2q}^4 \kappa_1^{-1} \log(M \vee N \vee T) \lesssim (T \log(M \vee N \vee T))^{1/2} \mu_{2q}^2,$$

we finish the proof. $\square$

Lemma 4. *For any $p > 0$, we have*

$$\|(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{M}^p\|_{\ell_0} \leq \min\{J_{\square} p^2, MN\}.$$

Proof. We prove this for $p = 2$ and it can be similarly extended to larger p .

Recall all the support $\mathcal{J}_{\mathbf{i}}$ are included in the same rectangle $\{i_1 - k_1, \dots, i_1 + k_1\} \times \{i_2 - k_2, \dots, i_2 + k_2\}$ of size $J_{\square} = (2k_1 + 1)(2k_2 + 1)$. We denote $\mathbf{M}$ by its columns: $\mathbf{M} = \begin{bmatrix} \widetilde{\mathbf{m}}_{11} & \dots & \widetilde{\mathbf{m}}_{MN} \end{bmatrix}$. Then

$$(\mathbf{e}_{i_2} \otimes \mathbf{e}_{i_1})^\top \mathbf{M}^2 (\mathbf{e}_{j_2} \otimes \mathbf{e}_{j_1}) = \mathbf{m}_{\mathbf{i}}^\top \mathbf{m}_{\mathbf{j}} = \sum_{\mathbf{k}} [\mathbf{M}]_{\mathbf{i}, \mathbf{k}} [\mathbf{M}]_{\mathbf{k}, \mathbf{j}}.$$

The summation is non-vanishing only when $\mathcal{J}_{\mathbf{i}} \cap \mathcal{J}_{\mathbf{j}} \neq \emptyset$. When $|i_1 - j_1| > 2k_1$ or $|i_2 - j_2| > 2k_2$, the intersection is empty. And thus there are at most $(4k_1 + 1)(4k_2 + 2) \leq 4J_{\square}$ $\mathbf{j}$ s making the term non-vanishing. $\square$

Lemma 5. *Let $\mathbf{M} \in \mathbb{R}^{P_1 \times P_2}$ be a rank R matrix. Let $\mathbf{M} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top$ be its compact singular value decomposition and denote $\lambda_{\min} > 0$ the smallest non-zero singular value of $\mathbf{M}$. Then for any matrix $\widehat{\mathbf{M}}$ such that $\|\widehat{\mathbf{M}} - \mathbf{M}\|_{\text{F}} \leq \frac{\lambda_{\min}}{8}$, we have that the best rank R approximation of $\widehat{\mathbf{M}}$ under Frobenius norm (denoted by $\text{SVD}_R(\widehat{\mathbf{M}})$) satisfies*

$$\text{SVD}_R(\widehat{\mathbf{M}}) - \mathbf{M} = \mathbf{Z} - \mathbf{U}_\perp \mathbf{U}_\perp^\top \mathbf{Z} \mathbf{V}_\perp \mathbf{V}_\perp^\top + \mathbf{R},$$

where $\mathbf{Z} = \widehat{\mathbf{M}} - \mathbf{M}$ and $\mathbf{U}_\perp$ ($\mathbf{V}_\perp$ resp.) is the orthogonal complement of $\mathbf{U}$ ($\mathbf{V}$ resp.), and the remainder term $\mathbf{R}$ satisfies

$$\|\mathbf{R}\|_{\text{F}} \leq 40 \frac{\|\widehat{\mathbf{M}} - \mathbf{M}\|_{\text{F}}^2}{\lambda_{\min}}.$$

Proof. See the proof of Lemma 18 in Shen et al. (2025). □

Lemma 6. *Under Regime 1, suppose $\rho(\mathbf{M}) < 1$. And let $\mathbf{E}_t$ be i.i.d. with mean zero and $\mathbb{E}[\|\mathbf{E}_t\|_{\text{F}}^{2+\tilde{q}}] < +\infty$ for each $\mathbf{i} \in \mathcal{S}$ and some $\tilde{q} > 0$. Then we have*

$$\frac{1}{T-1} \sum_{t=1}^{T-1} \text{vec}(\mathbf{X}_t) \text{vec}(\mathbf{X}_t)^\top \xrightarrow{p} \mathbf{\Gamma}_0$$

as $T \rightarrow \infty$.

Proof. See Theorem 8.2.3 in Fuller (2009). □

S3 Additional Methodological and Numerical Details

This section collects additional methodological details and numerical results.

S3.1 Non-identifiability under non-nested neighborhoods

In this section, we give the formal statement of the ambiguity for neighborhood selection within non-nested candidates.

Proposition 1. *Let $K_1 > 0, K_2 \geq 0$. Then there exists a distribution $\mathcal{P}$, and $\mathbf{X}_t \stackrel{i.i.d.}{\sim} \mathcal{P}$, such that there exists $k_1 < K_1, k_2 > K_2$, and $\widehat{\mathbf{M}} \in \mathbb{R}^S$, the following holds:*

- $(2k_1 + 1)(2k_2 + 1) < (2K_1 + 1)(2K_2 + 1)$,
- $\text{supp}(\widehat{\mathbf{M}}) = \widetilde{\mathcal{J}}$,
- $y_t = \langle \mathbf{X}_t, \widehat{\mathbf{M}} \rangle$.

Proof. We set $k_1 = 0, k_2 = K_2 + 1$ (see Figure 1). Let $\mathbf{X}$ be supported on only two entries $\{(i_0 - K_1, j_0), (i_0, j_0 - K_2 - 1)\}$ and that $\mathbf{X}_t(i_0 - K_1, j_0) = \mathbf{X}_t(i_0, j_0 - K_2 - 1) \sim \mathcal{P}_0$, where $\mathcal{P}_0$ can be an arbitrary distribution. And let $\widehat{\mathbf{M}} = \mathbf{M}(i_0 - K_1, j_0)\mathbf{e}_{i_0}\mathbf{e}_{j_0 - K_2 - 1}^\top$. Then we have $y_t = \langle \mathbf{X}_t, \mathbf{M} \rangle = \mathbf{X}_t(i_0 - K_1, j_0) \cdot \mathbf{M}(i_0 - K_1, j_0) = \langle \mathbf{X}_t, \widehat{\mathbf{M}} \rangle$. $\square$

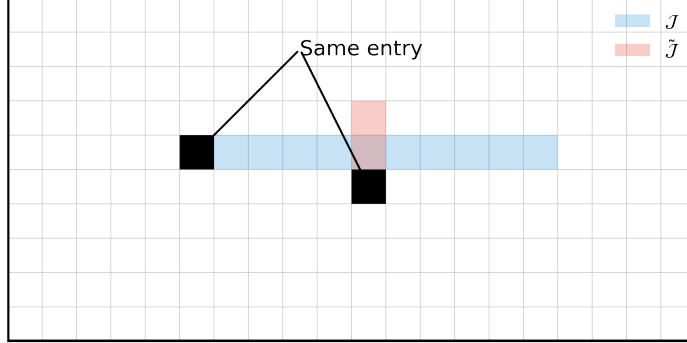


Figure 1: An example with the ground-truth bandwidths $K_1 = 5, K_2 = 0$ (in black), and $k_1 = 0, k_2 = 1$ (in red).

S3.2 Extension to the general lag-P model

This subsection collects the extension of LIAR to the general lag- P setting.

S3.2.1 Estimation for the general lag-P model

We introduce the estimation for the general lag P model. We work on the following model:

$$\mathbf{X}_t = \sum_{\mathbf{i} \in \mathcal{S}} \sum_{p \in [P]} \langle \mathbf{X}_{t-p}, \mathbf{M}_{p,\mathbf{i}} \rangle \mathbf{e}_{i_1} \mathbf{e}_{i_2}^\top + \mathbf{E}_t, \quad (\text{S3.7})$$

For each $\mathbf{i} \in \mathcal{S}$, we denote $\mathbf{x}_t^{(i)} := \mathbf{X}_t(\mathcal{J}_i)$, $\mathbf{m}_p^{(i)} = \mathbf{M}_{p,i}(\mathcal{J}_i) \in \mathbb{R}^{|\mathcal{J}_i|}$. Define the stacked coefficient and per-time design vectors

$$\mathbf{m}_i := \begin{bmatrix} \mathbf{m}_1^{(i)} \\ \vdots \\ \mathbf{m}_P^{(i)} \end{bmatrix}, \quad \mathbf{r}_t^{(i)} := \begin{bmatrix} \mathbf{x}_{t-1}^{(i)} \\ \vdots \\ \mathbf{x}_{t-P}^{(i)} \end{bmatrix} \in \mathbb{R}^{P|\mathcal{J}_i|}$$

Then the scalar regression for entry $\mathbf{i}$ reads, for $t = P + 1, \dots, T$, $[\mathbf{X}_t]_i = \langle \mathbf{r}_t^{(i)}, \mathbf{m}_i \rangle + [\mathbf{E}_t]_i$. Stacking over t gives the standard linear model $\mathbf{z}_i = \mathbf{Y}_i \mathbf{m}_i + \boldsymbol{\nu}_i$, where

$$\begin{aligned} \mathbf{z}_i &= \begin{bmatrix} [\mathbf{X}_{P+1}]_i \\ \vdots \\ [\mathbf{X}_T]_i \end{bmatrix}, & \boldsymbol{\nu}_i &= \begin{bmatrix} [\mathbf{E}_{P+1}]_i \\ \vdots \\ [\mathbf{E}_T]_i \end{bmatrix} \in \mathbb{R}^{T-P}, \\ \mathbf{Y}_i &= \begin{bmatrix} \mathbf{r}_{P+1}^{(i)\top} \\ \vdots \\ \mathbf{r}_T^{(i)\top} \end{bmatrix} \in \mathbb{R}^{(T-P) \times P|\mathcal{J}_i|}. \end{aligned} \tag{S3.8}$$

The least square estimator is

$$\widehat{\mathbf{m}}_i = \arg \min_{\mathbf{m}} \frac{1}{2} \|\mathbf{z}_i - \mathbf{Y}_i \mathbf{m}\|_{\ell_2}^2 = (\mathbf{Y}_i^\top \mathbf{Y}_i)^{-1} \mathbf{Y}_i^\top \mathbf{z}_i.$$

Notice this procedure can be done in parallel for different $\mathbf{i} \in \mathcal{S}$. We summarize the procedure in Algorithm 1.

Algorithm 1 Parallel Least Square for General Lag P

Input: Data $\{\mathbf{X}_t\}_{t=1}^T$, local neighborhood $\{\mathcal{J}_i\}_{i \in \mathcal{S}}$, lag parameter P

ParFor $i \in \mathcal{S}$

Collect the covariate $\mathbf{Y}_i$ and response $\mathbf{z}_i$ defined in (S3.8)

Solve least square $\widehat{\mathbf{m}}_i = (\mathbf{Y}_i^\top \mathbf{Y}_i)^{-1} \mathbf{Y}_i^\top \mathbf{z}_i$

End ParFor

Output: $\{\widehat{\mathbf{m}}_i\}_{i \in \mathcal{S}}$

Estimation of Coefficients for SP-LIAR. For SP-LIAR, $\mathcal{J}_i$ are rectangles by construction, and we denote $\mathbf{M}_p^{[i]} := \mathbf{M}_{p,i}[\mathcal{J}_i]$ as a sub-matrix of $\mathbf{M}_{p,i}$.

Arrange these matrices into the block matrix

$$\mathbf{M}_p^{\text{blk}} = \begin{bmatrix} \mathbf{M}_p^{[1,1]} & \cdots & \mathbf{M}_p^{[1,N]} \\ \vdots & & \vdots \\ \mathbf{M}_p^{[M,1]} & \cdots & \mathbf{M}_p^{[M,N]} \end{bmatrix}.$$

Then under SP-LIAR, $\mathbf{M}_p^{\text{blk}}$ has rank at most R . This motivates us to

consider the best rank R approximation of the following estimator of $\mathbf{M}_p^{\text{blk}}$,

where $\widehat{\mathbf{m}}_i = \begin{bmatrix} \widehat{\mathbf{m}}_1^{(i)\top} & \cdots & \widehat{\mathbf{m}}_P^{(i)\top} \end{bmatrix}^\top$:

$$\widehat{\mathbf{M}}_p^{\text{blk}} = \begin{bmatrix} \text{mat}(\widehat{\mathbf{m}}_p^{(1,1)}) & \cdots & \text{mat}(\widehat{\mathbf{m}}_p^{(1,N)}) \\ \vdots & & \vdots \\ \text{mat}(\widehat{\mathbf{m}}_p^{(M,1)}) & \cdots & \text{mat}(\widehat{\mathbf{m}}_p^{(M,N)}) \end{bmatrix},$$

where $\mathbf{mat}$ is the inverse operation of $\mathbf{vec}$. Our algorithm is summarized in Algorithm 2.

Algorithm 2 Parallel Least Square for SP-LIAR

Input: Local neighborhood $\mathcal{J}_i$, data $\{\mathbf{X}_t\}_{t=1}^T$, lag parameter P , rank R

Run Algorithm 1 for $\{\widehat{\mathbf{m}}_i\}_{i \in \mathcal{S}}$

Perform rank R SVD approximation on $\widehat{\mathbf{M}}_p^{\text{blk}}$ and get: $\widehat{\mathbf{M}}_{p,R}^{\text{blk}} = \text{SVD}_R(\widehat{\mathbf{M}}_p^{\text{blk}})$

Output: $\{\widehat{\mathbf{M}}_{p,R}^{(i)}\}_{p \in [P], i \in \mathcal{S}}$, where $\widehat{\mathbf{M}}_{p,R}^{(i)}$ be the (i_1, i_2) -th block of $\widehat{\mathbf{M}}_{p,R}^{\text{blk}}$

S3.2.2 Bandwidth selection for the general lag-P model

We also restrict the search to a nested sequence indexed by k for each location $\mathbf{i} = (i_1, i_2)$:

$$\mathcal{J}_i[1] \subsetneq \cdots \subsetneq \mathcal{J}_i[k_0] \subsetneq \cdots \subsetneq \mathcal{J}_i[K_0] \subset \mathcal{S},$$

where $\mathcal{J}_i[k_0] = \mathcal{J}_i$, and K_0 is a prescribed cap. Write $s_k = |\mathcal{J}_i[k]|$ (depends on k , and implicitly on $\mathbf{i}$). For any level k , let $\mathbf{x}_t[k] := \mathbf{X}_t(\mathcal{J}_i[k]) \in \mathbb{R}^{s_k}$ and stack P lags into one vector $\mathbf{r}_t[k] = \begin{bmatrix} \mathbf{x}_t[k]^\top & \cdots & \mathbf{x}_{t-P+1}[k]^\top \end{bmatrix}^\top \in \mathbb{R}^{Ps_k}$. Then we form the design matrix $\mathbf{Y}[k] \in \mathbb{R}^{(T-P) \times Ps_k}$:

$$\mathbf{Y}[k] = \begin{bmatrix} \mathbf{r}_{P+1}[k] & \cdots & \mathbf{r}_T[k] \end{bmatrix}^\top.$$

Also recall $\mathbf{z} = \begin{bmatrix} [\mathbf{X}_{P+1}]_{\mathbf{i}} & \cdots & [\mathbf{X}_T]_{\mathbf{i}} \end{bmatrix}^\top$. The residual sum of squares is

$$\text{RSS}_i(k) := \|\mathbf{z} - \mathbf{Y}[k](\mathbf{Y}[k]^\top \mathbf{Y}[k])^{-1} \mathbf{Y}[k]^\top \mathbf{z}\|_{\ell_2}^2.$$

And we use the Bayesian information criterion for the bandwidth selection.

$$\text{BIC}_{\mathbf{i}}(k) = \log \text{RSS}_{\mathbf{i}}(k) + D_0 \frac{|\mathcal{J}_{\mathbf{i}}[k]| \cdot P}{T} \log(M \vee N \vee T). \quad (\text{S3.9})$$

For each $\mathbf{i} \in \mathcal{S}$, we select $\hat{k}_{\mathbf{i}}$ by $\hat{k}_{\mathbf{i}} = \arg \min_{k \leq K_0} \text{BIC}_{\mathbf{i}}(k)$.

S3.3 Additional Simulation Results

S3.3.1 Visualization of the CLT

To illustrate the asymptotic normality result in Theorem 1 of the main paper, we conduct an oracle diagnostic under the same data-generating mechanism as the estimation-accuracy experiment in Section 5. Specifically, we consider the lag-one LIAR model with $P = 1$, $(M, N) = (10, 10)$, and a square neighborhood with radius $K = 3$. The local coefficients are generated as in Section 5 and the resulting transition matrix is rescaled to have spectral radius smaller than one. The innovations are Gaussian with the same noise level as in Section 5, and the first 500 observations are discarded as burn-in.

We focus on the interior location $\mathbf{i}_0 = (5, 5)$. Its true neighborhood $\mathcal{J}_{\mathbf{i}_0}$ is not truncated, so $|\mathcal{J}_{\mathbf{i}_0}| = (2K + 1)^2 = 49$. Let j_0 denote the coordinate of $\mathbf{m}_{\mathbf{i}_0}$ corresponding to the self-lag entry at $\mathbf{i}_0$. Under the column-major ordering used for $\mathcal{J}_{\mathbf{i}_0}$, this coordinate is $j_0 = 25$ in one-based indexing.

For each sample size $T \in \{500, 1000, 2000, 5000\}$, we generate $B = 500$

independent replications. In each replication, we estimate $\mathbf{m}_{i_0}$ by least squares using the oracle true neighborhood $\mathcal{J}_{i_0}$, rather than a BIC-selected neighborhood. This construction isolates the least-squares CLT from the additional effect of neighborhood selection. We then compute the studentized statistic

$$Z_{i_0, j_0} = \frac{\widehat{\mathbf{m}}_{i_0, j_0} - \mathbf{m}_{i_0, j_0}}{\widehat{\sigma}_{i_0} \left\{ [(\mathbf{Y}_{i_0}^\top \mathbf{Y}_{i_0})^{-1}]_{j_0, j_0} \right\}^{1/2}},$$

where

$$\widehat{\sigma}_{i_0}^2 = \frac{\|\mathbf{z}_{i_0} - \mathbf{Y}_{i_0} \widehat{\mathbf{m}}_{i_0}\|_{\ell_2}^2}{T - 1 - |\mathcal{J}_{i_0}|}.$$

Figure 2 compares the empirical distributions of Z_{i_0, j_0} with the standard normal benchmark across sample sizes. The distributions are centered around zero and become increasingly close to the standard normal density as T increases.

S3.3.2 Additional numerical results for neighborhood selection

After the oracle CLT diagnostic above, we next consider a larger-sample version of the BIC selection experiment in Section 5. The data-generating mechanism is unchanged: the true neighborhood radius is $K = 3$, and BIC selects over $k \in \{0, 1, 2, 3, 4, 5\}$. Only the sample size is increased to $T \in \{4000, 5000, 6000\}$.

Figure 3 reports the per-pixel success rates for $M = N \in \{10, 15, 20\}$.

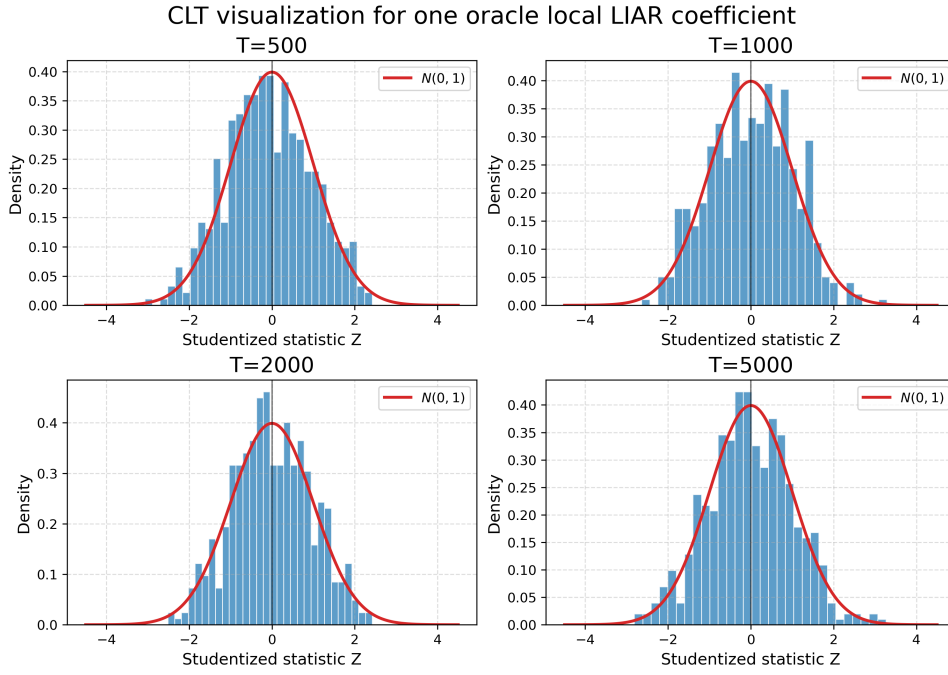


Figure 2: Visualization of the CLT for the representative coefficient $\mathbf{m}_{i_0, j_0}$. The empirical distributions of Z_{i_0, j_0} are compared with the standard normal benchmark across sample sizes.

Compared with the moderate-sample results in the main paper, the success rates further improve as T increases. These results are consistent with the neighborhood-selection consistency result in Theorem 5.

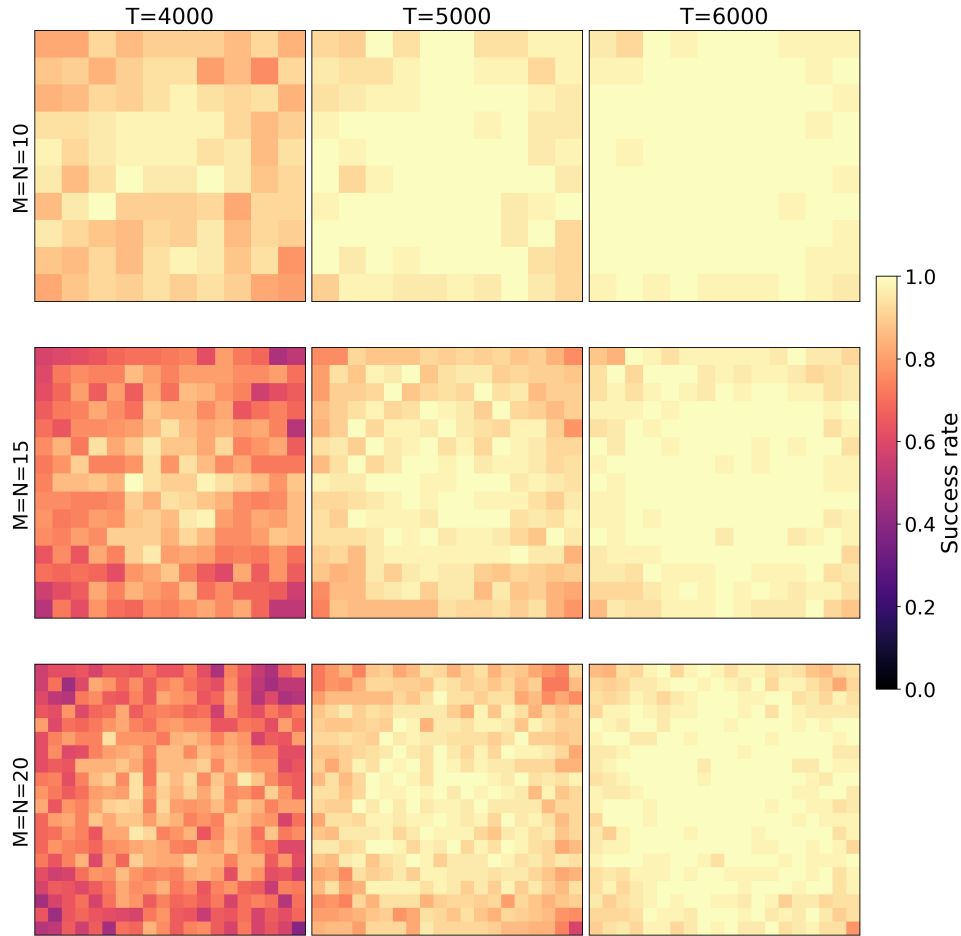


Figure 3: Additional neighborhood-selection results for $M = N \in \{10, 15, 20\}$ and $T \in \{4000, 5000, 6000\}$.

S4 Tensor Local Interaction Autoregression and TEC

Application

This section collects the tensor extension of LIAR and the tensor TEC application. We first formulate the tensor local interaction autoregression and its entrywise least-squares estimator, and then apply the same construction to TEC data organized by two spatial modes and one within-hour temporal mode.

S4.1 Tensor LIAR formulation and estimation

We consider a local interaction autoregressive model for tensor-valued time series. The modeling idea mirrors the matrix case: each tensor entry evolves from past values in a local neighborhood, while the neighborhood itself may vary across entries and across tensor modes.

We use calligraphic-font bold-face letters, such as $\mathcal{M}$ and $\mathcal{X}$, to denote tensors. An order- d tensor is a d -way array; for example, $\mathcal{X} \in \mathbb{R}^{N_1 \times \dots \times N_d}$ has size N_j along its j th mode. Let $\mathcal{S}_d = [N_1] \times \dots \times [N_d]$ be the tensor index set. For a tensor time series $\{\mathcal{X}_t\}_{t \geq 0} \subset \mathbb{R}^{\mathcal{S}_d}$, the tensor LIAR takes the form

$$\mathcal{X}_t = \sum_{i \in \mathcal{S}_d} \langle \mathcal{X}_{t-1}, \mathcal{M}_i \rangle \cdot e_{i_1} \circ \dots \circ e_{i_d} + \mathcal{E}_t,$$

where $\mathbf{i} = (i_1, \dots, i_d) \in \mathcal{S}_d$, and $\mathcal{M}_{\mathbf{i}}$ is supported on a local neighborhood $\mathcal{J}_{\mathbf{i}}$ of $\mathbf{i}$. For example, a rectangular neighborhood with radius vector $(k_{1,\mathbf{i}}, \dots, k_{d,\mathbf{i}})$ contains indices whose ℓ th coordinate lies within $k_{\ell,\mathbf{i}}$ of i_ℓ , truncated at the tensor boundary.

Vectorizing both sides gives

$$\mathbf{x}_t = \underbrace{\sum_{\mathbf{i} \in \mathcal{S}_d} (\mathbf{e}_{i_d} \otimes \dots \otimes \mathbf{e}_{i_1}) \text{vec}(\mathcal{M}_{\mathbf{i}})^\top}_{=: \mathbf{M}} \mathbf{x}_{t-1} + \boldsymbol{\epsilon}_t,$$

where $\mathbf{x}_t = \text{vec}(\mathcal{X}_t)$ and $\boldsymbol{\epsilon}_t = \text{vec}(\mathcal{E}_t)$. Denote the local covariate vector and coefficient vector by

$$\mathbf{x}_t^{(i)} = \mathcal{X}_t(\mathcal{J}_{\mathbf{i}}) \in \mathbb{R}^{|\mathcal{J}_{\mathbf{i}}|}, \quad \mathbf{m}_{\mathbf{i}} = \mathcal{M}_{\mathbf{i}}(\mathcal{J}_{\mathbf{i}}) \in \mathbb{R}^{|\mathcal{J}_{\mathbf{i}}|}.$$

Then, for each entry $\mathbf{i}$ and $t = 2, \dots, T$, the tensor model reduces to the scalar regression

$$[\mathcal{X}_t]_{\mathbf{i}} = \langle \mathbf{x}_{t-1}^{(i)}, \mathbf{m}_{\mathbf{i}} \rangle + [\mathcal{E}_t]_{\mathbf{i}}.$$

Stacking over time gives the standard linear model

$$\mathbf{z}_{\mathbf{i}} = \mathbf{Y}_{\mathbf{i}} \mathbf{m}_{\mathbf{i}} + \boldsymbol{\nu}_{\mathbf{i}},$$

where

$$\begin{aligned} \mathbf{z}_{\mathbf{i}} &= \begin{bmatrix} [\mathcal{X}_2]_{\mathbf{i}} & \dots & [\mathcal{X}_T]_{\mathbf{i}} \end{bmatrix}^\top, \quad \boldsymbol{\nu}_{\mathbf{i}} = \begin{bmatrix} [\mathcal{E}_2]_{\mathbf{i}} & \dots & [\mathcal{E}_T]_{\mathbf{i}} \end{bmatrix}^\top \in \mathbb{R}^{T-1}, \\ \mathbf{Y}_{\mathbf{i}} &= \begin{bmatrix} \mathbf{x}_1^{(i)} & \dots & \mathbf{x}_{T-1}^{(i)} \end{bmatrix}^\top \in \mathbb{R}^{(T-1) \times |\mathcal{J}_{\mathbf{i}}|}. \end{aligned}$$

The least-squares estimator is

$$\widehat{\mathbf{m}}_i = \arg \min_{\mathbf{m}} \frac{1}{2} \|\mathbf{z}_i - \mathbf{Y}_i \mathbf{m}\|_{\ell_2}^2 = (\mathbf{Y}_i^\top \mathbf{Y}_i)^{-1} \mathbf{Y}_i^\top \mathbf{z}_i.$$

Algorithm 3 summarizes the parallel implementation.

Algorithm 3 Parallel least squares for tensor LIAR

Input: Local neighborhoods $\{\mathcal{J}_i\}_{i \in \mathcal{S}_d}$, data $\{\mathbf{x}_t\}_{t=1}^T$

ParFor $i \in \mathcal{S}_d$

Collect the covariate $\mathbf{Y}_i$ and response $\mathbf{z}_i$

Solve least squares $\widehat{\mathbf{m}}_i = (\mathbf{Y}_i^\top \mathbf{Y}_i)^{-1} \mathbf{Y}_i^\top \mathbf{z}_i$

End ParFor

Output: $\{\widehat{\mathbf{m}}_i\}_{i \in \mathcal{S}_d}$

Theoretical results for coefficient estimation, auto-covariance estimation, and neighborhood selection extend to this tensor setting by the same proof strategy as in the matrix case. The framework also allows tensor autoregressions with no locality along some modes, by choosing $\mathcal{J}_i$ to span the full range along those dimensions while remaining local along others.

S4.2 Tensor TEC application

The TEC application uses the tensor formulation above to separate coarse spatial variation from within-hour temporal variation. Starting from the same GNSS-derived TEC product (Sun et al., 2023), we first aggregate the spatial grid to $4^\circ \times 4^\circ$, yielding a 46×91 spatial grid. We then stack,

for each clock hour, the twelve 5-minute frames into a third mode, which represents the within-hour slot. This produces hourly tensors of size $46 \times 91 \times 12$. Indexing these tensors by hour gives a tensor time series of shape $46 \times 91 \times 12 \times 720$ for the 30-day period in June 2019.

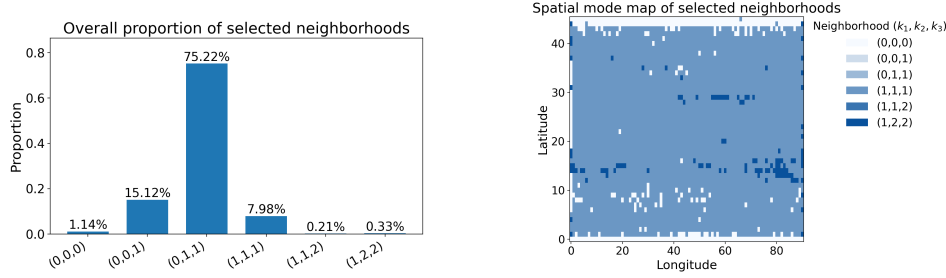
We use the first 80% of hours for training and the remaining 20% for testing. Test RMSE is computed from one-step-ahead full-tensor predictions over the test period, averaging squared prediction errors over all test hours and tensor entries. Neighborhood sizes are selected entrywise by BIC over the candidate set

$$\{(0, 0, 0), (0, 0, 1), (0, 1, 1), (1, 1, 1), (1, 1, 2), (1, 2, 2), (2, 2, 2)\}.$$

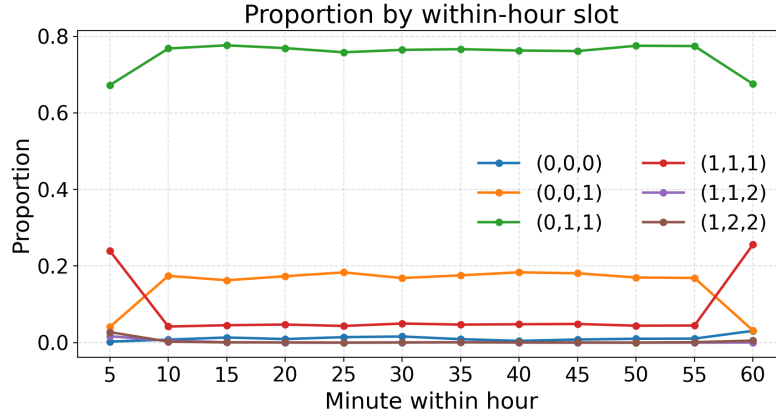
Figure 4a reports the selection frequencies across all locations and within-hour slots. The selections concentrate on small neighborhoods, with $(0, 1, 1)$ most common and $(0, 0, 1)$ the primary alternative. This concentration suggests that the hourly TEC tensor dynamics are driven mainly by short-range interactions along the spatial and within-hour dimensions, rather than by dense full-tensor dependence.

To examine when and where these selected neighborhoods arise, Figure 4c shows the proportion of each selected neighborhood versus within-hour slot, from 5 to 60 minutes. The slot-wise proportions are fairly stable across the hour, with only mild changes near hour boundaries. This stabil-

ity suggests that the learned local dependence pattern is persistent within an hour, with only modest changes near transitions between adjacent hourly blocks. Figure 4b displays the spatial mode map, defined as the most frequently selected neighborhood per pixel over the 12 within-hour slots.



(a) Selection frequencies of neighborhoods (k_1, k_2, k_3) over all locations and slots. (b) Spatial mode map: per pixel, the most frequently selected neighborhood over 12 slots.



(c) Proportion of each neighborhood versus within-hour slot, from 5 to 60 minutes.

Figure 4: Neighborhood selection in the tensor TEC series using entrywise BIC. Top: (a) overall selection frequencies and (b) spatial mode map. Bottom: (c) proportions by within-hour slot.

As a benchmark, we fit the pixel-wise autoregression LIAR-P. Its test RMSE is 1.29, whereas LIAR with BIC-selected neighborhoods attains a lower test RMSE of 1.12. Thus, the tensor neighborhoods selected by BIC capture spatio-temporal interactions that are useful for one-step-ahead full-tensor prediction.

Bibliography

- Fuller, W. A. (2009). *Introduction to statistical time series*. John Wiley & Sons.
- Liu, W., H. Xiao, and W. B. Wu (2013). Probability and moment inequalities under dependence. *Statistica Sinica*, 1257–1272.
- Shen, Y., J. Li, J.-F. Cai, and D. Xia (2025). Computationally efficient and statistically optimal robust high-dimensional linear regression. *The Annals of Statistics* 53(1), 374–399.
- Sun, H., Y. Chen, S. Zou, J. Ren, Y. Chang, Z. Wang, and A. Coster (2023). Complete global total electron content map dataset based on a video imputation algorithm vista. *Scientific Data* 10(1), 236.
- Wu, W. B. (2005). Nonlinear system theory: Another look at dependence. *Proceedings of the National Academy of Sciences* 102(40), 14150–14154.