

Supplementary material for “Significance testing for canonical correlation analysis in high dimensions”

BY IAN W. McKEAGUE

Department of Biostatistics, Columbia University, Room R639, 722 West 168th Street, New York, NY 10032, USA
 im2131@columbia.edu

AND XIN ZHANG

Department of Statistics, Florida State University, 214 OSB, 117 N. Woodward Ave. Tallahassee, FL 32306, USA
 xzhang8@fsu.edu

S.1. GREEDY ALGORITHM

S.1.1. The increments in the sample Pillai trace

Recall that, Fig. 1 gives the results from a toy example showing that the proposed greedy algorithm for finding the maximal sample root-Pillai trace under varying sparsity constraints provides almost perfect agreement with the full search. The data were generated from Model A1 used in the simulation study (in Section 4), with $p = q = 10$, the true numbers of active variables $s_x^* = s_y^* = 3$, $\tau_{\max} = 0.8$, the true number of non-zero canonical correlations $K = 1$, and $n = 500$. The result of the greedy search agrees with the full search at all sparsity levels except at $s = s_x = s_y = 2$.

The left panel of Fig. S.1 shows the successive increments in the sample Pillai trace when adding one variable at a time (either X_k or Y_j), using the same simulation model as the first panel except with $p = q = 5000$. This plot is analogous to the scree plot used in principal component analysis and factor analysis, providing intuition and graphical diagnostics for how sparse the true model might be. It is clear from the plot that $s_x + s_y = 6$ gives the most reasonable terminating point; also, at that point we have $s_x = s_y = 3$ (not shown in the plot), agreeing with the true numbers of active variables. The right panel of Fig. S.1 shows the corresponding values of the studentized $\hat{\tau}_{\max}$ used as the test statistic for the proposed test. As the sparsity level s increases, the test statistic monotonically increases until s reaches the true value $s_x^* = s_y^* = 3$, but the accuracy of the sample coherence matrix $C_{X_K Y_J}$ used in the test statistic decreases as its dimension ($s_x \times s_y$) grows, causing a decrease in power at larger values of s . In our experience, the proposed test performs well for relatively small sparsity levels, as in this simulation example, but degrades when s_x^* and s_y^* become large (say $s_x^* = s_y^* = 20$).

Since the true sparsity levels are $s_x^* = s_y^* = 5$, we have satisfactory results (p-value less than 0.01) when $s \geq 4$. When fewer variables are selected, $s \leq 3$, the test procedure is not powerful enough to reject the null at $\alpha = 0.05$ level. It is also interesting to note that slightly over-specified $s = 6$ or 7 produces more significant test results than at the true sparsity level.

S.1.2. Submodularity

To gain some further insight into the performance the proposed greedy search algorithm, we show in a simulation example that the root-Pillai trace comes close to satisfying the submodular

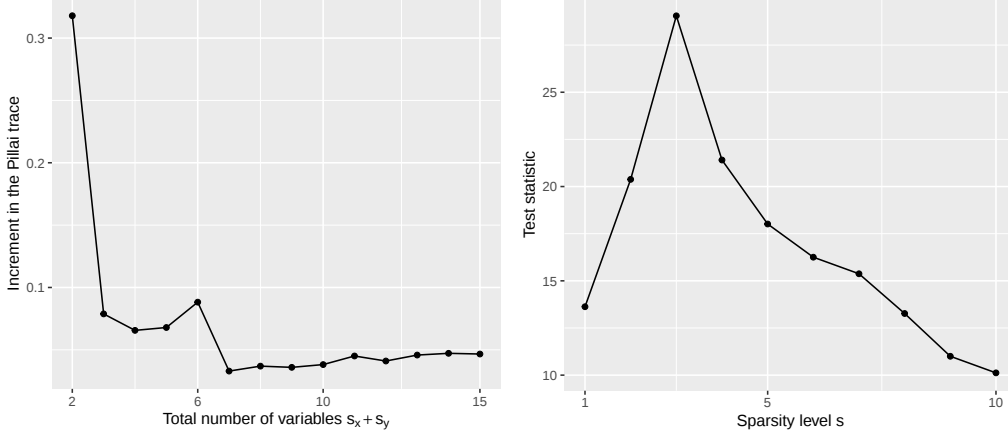


Fig. S.1. Results based on single samples generated under Model A1 with $n = 500$, $s_x^* = s_y^* = 3$, $\tau_{\max} = 0.8$. Left: scree plot of the increment in the sample Pillai trace, $\|C_{E_j|\mathcal{J}}X_{\mathcal{K}}\|_F^2$ or $\|C_{Y_{\mathcal{J}}R_{k|\mathcal{K}}}\|_F^2$, when adding one variable at a time in the greedy search, for $p = q = 5000$. Right: corresponding values of the studentized $\hat{\tau}_{\max}$ (test statistic for $\tau_{\max} = 0$) as a function of s .

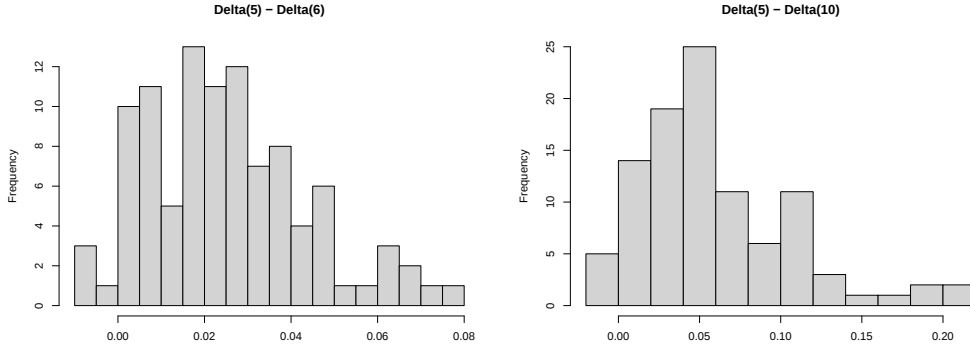


Fig. S.2. Histograms used for checking submodularity of the root-Pillai trace as a function of $S = (\mathcal{J}, \mathcal{K})$ for the GBM data set. The plotted values of $\Delta(e_i | S_1) - \Delta(e_i | S_2)$ should be non-negative if the function is submodular. The histograms are based on 100 randomly sampled elements $e_i \notin S_2$. The set S_2 is randomly chosen with $|\mathcal{K}| = |\mathcal{J}| = 6$ on the left, and with $|\mathcal{K}| = |\mathcal{J}| = 10$ on the right. For $S_1 \subset S_2$ we take the first 5 elements of S_2 .

property (which, if true, would give a guarantee of finding the maximum to within a factor of $1/e$). Maximization of the root-Pillai trace over $\mathcal{K}$ and $\mathcal{J}$ is a discrete combinatorial optimization problem for a set function $f : 2^V \mapsto \mathbb{R}$, where the finite set $V = \{1, \dots, p\} \times \{1, \dots, q\}$ and the utility function is $f(S) = \|C_{X_{\mathcal{K}}Y_{\mathcal{J}}}\|_F$ for $S = (\mathcal{J}, \mathcal{K}) \subseteq V$. Submodularity in set functions is the discrete analogy of convexity in continuous functions. Fast greedy algorithms with theoretical guarantees have been developed for submodular function maximization (see Nemhauser & Wolsey, 1978; Krause & Golovin, 2014; Khim et al., 2016, for example), when the util-

ity function is also monotonic. Specifically, the set function f is monotonic if for any two sets $S_1 \subseteq S_2 \subseteq V$, we have $f(S_1) \leq f(S_2)$. Based on Lemma 1, it is not difficult to see that $f(S) = \|C_{X_K Y_J}\|_F$ is monotonic.

We now briefly review the definition of submodularity and then numerically demonstrate that our set function f can be close to submodular, although not exactly so. A key concept related to submodularity is the discrete derivative. The discrete derivative of f at S with respect to a new element $e \in V$ is defined as $\Delta(e | S) = f(S \cup \{e\}) - f(S)$. Then f is called submodular if, for any $S_1 \subseteq S_2 \subset V$ and $e \in V \setminus S_2$, we have $\Delta(e | S_1) \geq \Delta(e | S_2)$. We consider a simple numerical experiment using the real data in Section 5, where $(p, q, n) = (1000, 534, 397)$. First, we randomly sampled $S_2 \subset V$ with size $s_x = s_y = 6$ (or 10) and defined the first 5 elements in S_2 to be S_1 . Then we computed the histogram of the differences $\Delta(e_i | S_1) - \Delta(e_i | S_2)$ for 100 randomly selected elements $e_i \notin S_2$. The results displayed in Fig. S.2 indicate a close approximation to submodularity since there are very few negative values in each histogram. In contrast, the Pillai trace is readily seen to violate the submodular property: the scree plot in Fig. S.1 is not monotonically decreasing.

S.2. COMPUTATION TIME

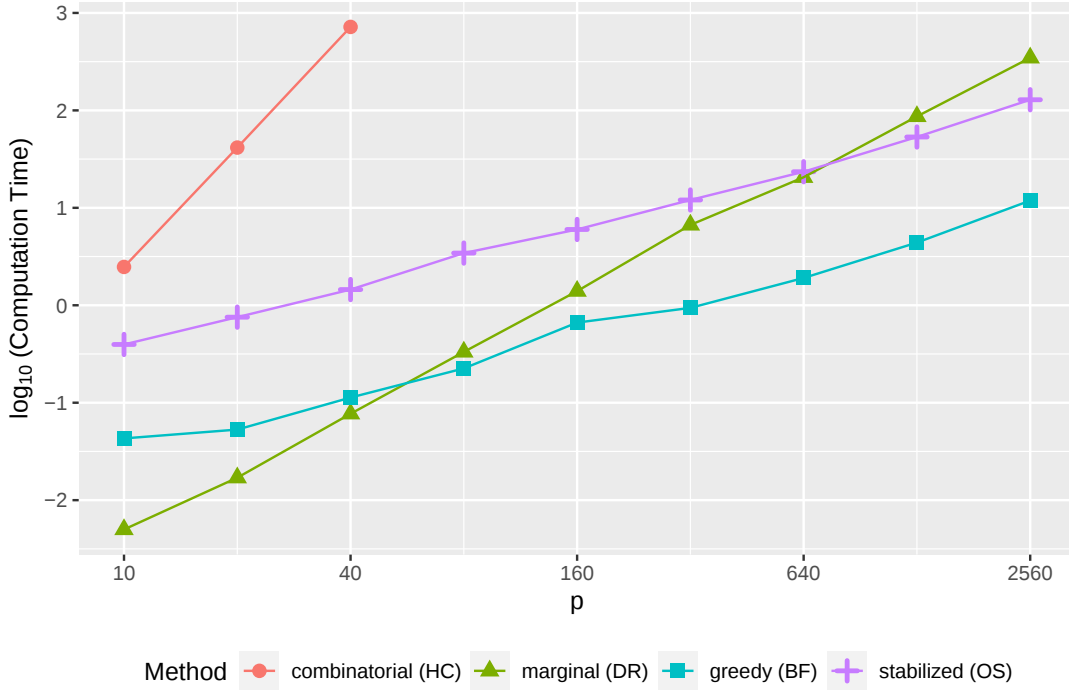


Fig. S.3. Computation time comparison. B-DR and BH-DR have roughly the same computation time and we only report BH-DR.

In this section, we report the computation time of each of the following methods: Higher Criticism (HC) (Donoho & Jin, 2004, 2015) based on p-values computed from the F-test for all $\binom{p}{s_x} \binom{q}{s_y}$ combinations of variables; the normal approximation test (B-DR or BH-DR) described in DiCiccio & Romano (2017) and applied to all pq pairs of marginal correlations $\text{corr}(X_k, Y_j)$;

the F-test on selected variables from the greedy search algorithm and with Bonferroni correction (BF); the proposed testing procedure using the stabilized one-step estimator (OS) with the greedy search algorithm.

For the simulation example of Model N, Figure S.3 shows the computation time for one data replication. We set $s_x = s_y = 2$ so that HC is feasible for $p = q = 40$. All computation was done on a MacOS laptop computer with 3.1 GHz Intel Core i5 processor and 16 GB memory. Recall that $n = 500$ and we vary $p = q \in \{10, 20, \dots, 2560\}$. From this plot, we clearly see that the proposed greedy search algorithm, applied to either the F-test on selected variables with Bonferroni correction or the stabilized one-step estimator, scales better with high dimensionality than methods based on combinatorial full search over $\binom{p}{s_x} \binom{q}{s_y}$ test statistics or based on marginal selection over all pq pairs of variables. For very large p and q , we find that the methods based on our greedy search algorithm (i.e., BF and OS) are computationally tractable while other methods are not.

S.3. GLIOBLASTOMA MULTIFORME (GBM) DATA ANALYSIS

Table S.1 lists the most correlated variables under sparsity level $s = 3$. Interestingly, the first two microRNA measurements (*hsa.miR.219* and *hsa.miR.222*) also appear in a reported dependency network of important microRNAs obtained by precision matrix estimation (Wang, 2015, Figure 1). In Fig. S.5, we repeat the analysis in the paper for 10 random re-orderings of the data, and compare the results to the ones without random re-ordering. The random ordering of the samples has little effects on the significance test result. The two plots are very similar, and it is clear that the number of important variables should be specified as less than 25. In Table S.2, we report the selected top-25 variables for X (gene expression) and Y (microRNA). The selection was based on the proposed greedy algorithm (Algorithm 1).

For $s = s_x = s_y = 1$, we calculated the 95% confidence intervals based on 100 random re-ordering of the original data. The endpoints of the CIs are displayed in Fig. S.4, where each point in the scatterplot represents one CI. All the CIs have very similar widths, and are far away from zero, which is consistent with the finding of very small p-values.

Table S.1. *GBM data analysis. Marginal correlations between the selected genes (each row) and microRNAs (each column) at sparsity level $s = s_x = s_y = 3$. The stabilized one-step estimator $\hat{\tau}_{\max} = 0.931$ with standard error 0.085*

	hsa.miR.219	hsa.miR.222	hsa.miR.138
SPIRE2	0.76	-0.27	0.50
FGF9	0.76	-0.06	0.33
ZNF553	-0.24	0.61	-0.17

S.4. CANONICAL GRADIENT

S.4.1. Some matrix notation and properties

The Frobenius norm of a matrix $A \in \mathbb{R}^{p \times q}$ is $\|A\|_F = (\sum_{i,j} A_{ij}^2)^{1/2}$. The maximum norm $\|A\|_{\max} = \max_{i,j} |A_{ij}|$ and the ℓ_1 -norm $\|A\|_1 = \max_j \sum_i |A_{ij}|$. The ℓ_2 -norm (or operator norm) $\|A\|$ is its largest singular value, and the ℓ_∞ -norm is $\|A\|_\infty = \max_i \sum_j |A_{ij}|$. For a symmetric matrix A , $\|A\|_1 = \|A\|_\infty$. For any matrices $A \in \mathbb{R}^{p \times q}$ and $B \in \mathbb{R}^{r \times s}$,

Table S.2. *Top-25 selected variables based on Algorithm 1 for GBM Data Set. The sequentially selected variables are ordered from top left to bottom right*

		Selected Variables				
gene expression		SPIRE2	FGF9	ZNF553	CPD	UCK2
		KIAA1632	SCGB3A2	SAMD9	KBTBD5	AURKB
		ERCC1	FAM113B	CCL19	DLX4	GCG
		LRRTM3	SCN10A	SMC1A	SBNO2	CENTG2
		RHOV	COX7A1	ROGDI	SUHW2	SNX1
microRNA		hsa-miR-219	hsa-miR-222	hsa-miR-138	hsa-miR-124a	hsa-miR-130b
		hsa-miR-92b	hsa-miR-423	hsa-miR-9*	hsa-miR-199a*	hsa-miR-22
		hsa-miR-195	hsa-miR-21	hsa-miR-204	hsa-miR-30e-3p	hsa-miR-181d
		hsa-miR-186	hsa-miR-484	hsa-miR-210	kshv-miR-K12-12	hsa-miR-376a
		hsa-miR-106b	hsa-miR-9	d1.hsa-miR-15b	hsa-miR-181a	hsa-miR-217

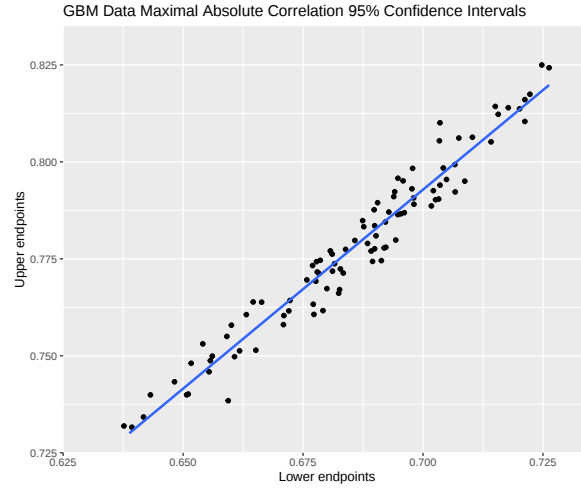


Fig. S.4. GBM data analysis. The lower and upper endpoints of the 95% confidence interval for $\tau_{\max}$ when $s = 1$, for 100 random re-orderings of the data.

we have $\|A\| \leq \|A\|_F$, $p^{-1/2}\|A\|_1 \leq \|A\| \leq q^{1/2}\|A\|_1$, $\|AB\| \leq \|A\|\|B\| = \|A \otimes B\|$ and $\|A\| \leq (\|A\|_1\|A\|_\infty)^{1/2}$. For a symmetric matrix $A \in \mathbb{R}^{p \times p}$, we have $\|A\| \leq \|A\|_1 = \|A\|_\infty \leq p\|A\|_{\max}$ and also $\|A\| \leq \|A\|_F \leq p\|A\|_{\max}$. 100

We will also be using the following matrix inequalities repeatedly in the proofs below. For matrices A, B, X, Y with compatible dimensions, we have

$$\begin{aligned}
 |\text{tr}(AXBY)| &= |\text{vec}^T(Y)(B^T \otimes A)\text{vec}(X)| \\
 &\leq \|\text{vec}(Y)\| \cdot \|(B^T \otimes A)\text{vec}(X)\| \\
 &\leq \|\text{vec}(Y)\| \cdot \|(B^T \otimes A)\| \cdot \|\text{vec}(X)\| \\
 &= \|A\|\|B\| \cdot \|\text{vec}(X)\| \cdot \|\text{vec}(Y)\| \\
 &= \|A\|\|B\| \cdot \|X\|_F \cdot \|Y\|_F.
 \end{aligned}$$

The operator norm (largest singular value of a matrix) satisfies $\|AX\| \leq \|A\|\|X\| = \|A \otimes X\|$. 105

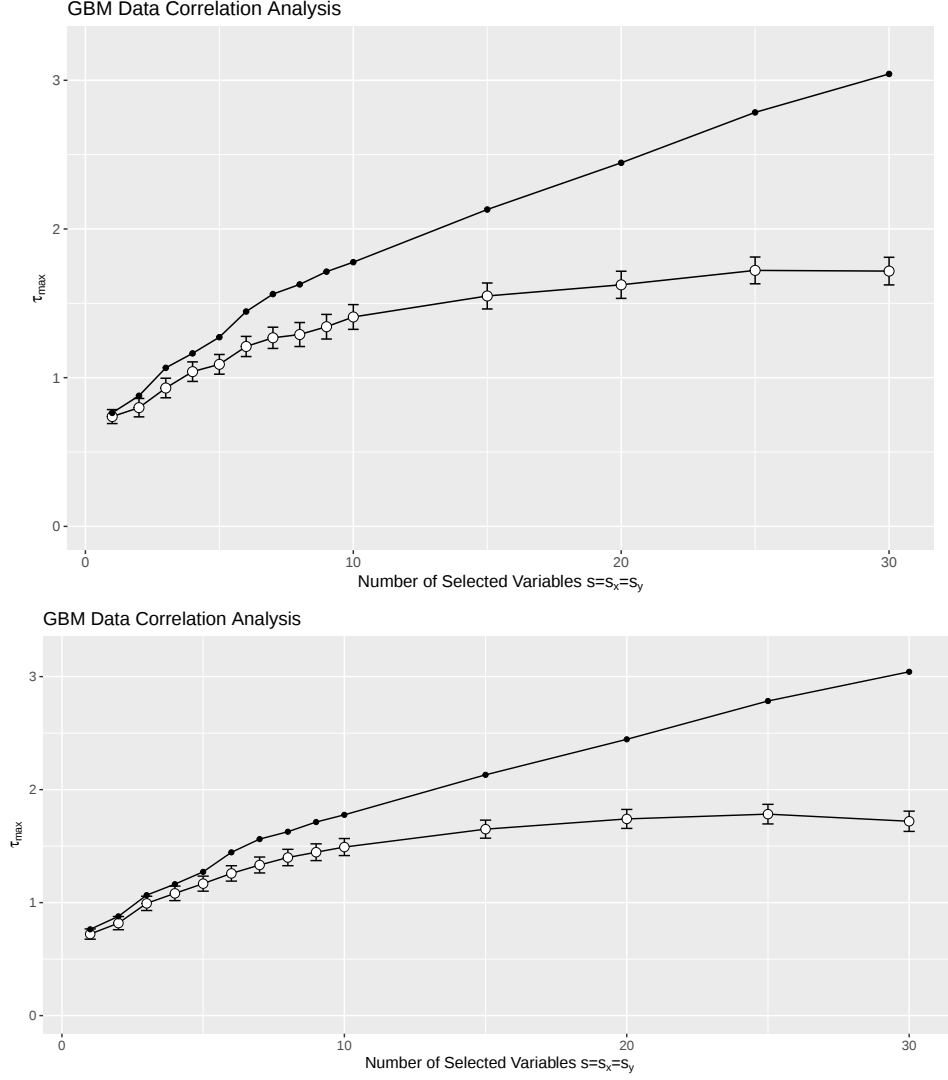


Fig. S.5. GBM data analysis. Estimates of the maximal root-Pillai trace $\tau_{\max}$ are plotted against $s = s_x = s_y$ varying from 1–30, with the selected variables at each step found using Algorithm 1. The black dots are estimated values without adjusting for post-selection. The white dots and associated 95% confidence intervals are based on the stabilized one-step estimator. The top panel is without random re-ordering of the data; the bottom panel uses 10 random re-orderings and reports the averaged point estimates and averaged lower/upper CI endpoints over these re-orderings.

S.4.2. Proof of Theorem 1 (Derivation of the canonical gradient)

In the following derivation, we write

$$\Phi(P) = \Phi^d(P) = \|\Lambda_{X_K Y_{\mathcal{J}}}\|_F^2 = \|\Lambda_{XY}\|_F^2 = \text{tr}(\Lambda_{XY} \Lambda_{XY}^T) = \text{tr}(\Sigma_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} \Sigma_{XY}^T).$$

Using standard results from matrix calculus,

$$\begin{aligned}
d\Phi(P) &= d\text{tr}(\Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\Sigma_{XY}^T) \\
&= \frac{\partial\text{tr}(\Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\Sigma_{XY}^T)}{\partial\Sigma_X} \cdot d\Sigma_X + \frac{\partial\text{tr}(\Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\Sigma_{XY}^T)}{\partial\Sigma_Y} \cdot d\Sigma_Y \\
&\quad + \frac{\partial\text{tr}(\Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\Sigma_{XY}^T)}{\partial\Sigma_{XY}} \cdot d\Sigma_{XY} \\
&= -\Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\Sigma_{XY}^T\Sigma_X^{-1} \cdot d\Sigma_X - \Sigma_Y^{-1}\Sigma_{XY}^T\Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1} \cdot d\Sigma_Y \\
&\quad + 2\Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1} \cdot d\Sigma_{XY},
\end{aligned}$$

where the dot indicates inner product of two matrices of same dimension: $A \cdot dB = \langle A, dB \rangle = \text{vec}^T(A)\text{vec}(dB) = \text{tr}(A^T dB)$. Hence

$$\begin{aligned}
\frac{d\Phi(P_\epsilon)}{d\epsilon} &= -\text{tr} \left\{ \Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\Sigma_{XY}^T\Sigma_X^{-1} \frac{d\Sigma_X(P_\epsilon)}{d\epsilon} \right\} \\
&\quad - \text{tr} \left\{ \Sigma_Y^{-1}\Sigma_{XY}^T\Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1} \frac{d\Sigma_Y(P_\epsilon)}{d\epsilon} \right\} + 2\text{tr} \left\{ \Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1} \frac{d\Sigma_{XY}(P_\epsilon)}{d\epsilon} \right\},
\end{aligned}$$

recalling that $P_\epsilon = (1 - \epsilon)P + \epsilon\delta_o$, where δ_o is the Dirac measure at $o = (x^T, y^T)^T$. Expressing $\Sigma_X(P_\epsilon)$ as 110

$$(1 - \epsilon)P(XX^T) + \epsilon xx^T - (1 - \epsilon)^2 P(X)P(X^T) - \epsilon(1 - \epsilon)\{P(X)x^T + xP(X^T)\} - \epsilon^2 xx^T,$$

by direct calculation we then have

$$\left. \frac{d\Sigma_X(P_\epsilon)}{d\epsilon} \right|_{\epsilon=0} = \lim_{\epsilon \downarrow 0} \frac{\Sigma_X(P_\epsilon) - \Sigma_X(P)}{\epsilon} = \{x - P(X)\}\{x^T - P(X^T)\}.$$

Similarly,

$$\left. \frac{d\Sigma_{XY}(P_\epsilon)}{d\epsilon} \right|_{\epsilon=0} = \{x - P(X)\}\{y^T - P(Y^T)\}, \quad \left. \frac{d\Sigma_Y(P_\epsilon)}{d\epsilon} \right|_{\epsilon=0} = \{y - P(Y)\}\{y^T - P(Y^T)\}.$$

Plugging-in these expressions, and including the relevant sets of indices $\mathcal{K}$ and $\mathcal{J}$, we obtain canonical gradient of $\Phi^d(P)$ as stated in (11):

$$\begin{aligned}
\left. \frac{d\Phi^d(P_\epsilon)}{d\epsilon} \right|_{\epsilon=0} &= -\{x_{\mathcal{K}} - \mathbb{E}_P(X_{\mathcal{K}})\}^T \Sigma_{X_{\mathcal{K}}}^{-1} \Sigma_{X_{\mathcal{K}}Y_{\mathcal{J}}} \Sigma_{Y_{\mathcal{J}}}^{-1} \Sigma_{Y_{\mathcal{J}}X_{\mathcal{K}}} \Sigma_{X_{\mathcal{K}}}^{-1} \{x_{\mathcal{K}} - \mathbb{E}_P(X_{\mathcal{K}})\} \\
&\quad - \{y_{\mathcal{J}} - \mathbb{E}_P(Y_{\mathcal{J}})\}^T \Sigma_{Y_{\mathcal{J}}}^{-1} \Sigma_{Y_{\mathcal{J}}X_{\mathcal{K}}} \Sigma_{X_{\mathcal{K}}}^{-1} \Sigma_{X_{\mathcal{K}}Y_{\mathcal{J}}} \Sigma_{Y_{\mathcal{J}}}^{-1} \{y_{\mathcal{J}} - \mathbb{E}_P(Y_{\mathcal{J}})\} \\
&\quad + 2\{y_{\mathcal{J}} - \mathbb{E}_P(Y_{\mathcal{J}})\}^T \Sigma_{Y_{\mathcal{J}}}^{-1} \Sigma_{Y_{\mathcal{J}}X_{\mathcal{K}}} \Sigma_{X_{\mathcal{K}}}^{-1} \{x_{\mathcal{K}} - \mathbb{E}_P(X_{\mathcal{K}})\}.
\end{aligned}$$

Then, as discussed in Section 2.1, the canonical gradient of $\Psi^d(P) = \|\Lambda_{X_{\mathcal{K}}Y_{\mathcal{J}}}\|_F$ is obtained via the relationship $\{\Psi^d(P)\}^2 = \Phi^d(P)$, resulting in

$$D^d(P)(o) = \left. \frac{d\Psi^d(P_\epsilon)}{d\epsilon} \right|_{\epsilon=0} = \frac{1}{2\Psi^d(P)} \left. \frac{d\Phi^d(P_\epsilon)}{d\epsilon} \right|_{\epsilon=0}$$

when $\Psi^d(P) \neq 0$.

To prove $\mathbb{E}_P\{D^d(P)(O)\} = 0$ it suffices to show that the canonical gradient of $\Phi^d(P)$ has zero mean. We apply the following property of trace and expectation operators. For any random vector $X \in \mathbb{R}^p$ with entries having finite second moments and any non-stochastic matrix $M \in$ 115

$\mathbb{R}^{p \times p}$,

$$\mathbb{E}(X^T M X) = \mathbb{E}\{\text{tr}(X^T M X)\} = \mathbb{E}\{\text{tr}(M X X^T)\} = \text{tr}\{M \mathbb{E}(X X^T)\}. \quad (\text{S.1})$$

Therefore, direct calculation shows that

$$\begin{aligned} \mathbb{E} \left\{ \frac{d\Phi^d(P_\epsilon)}{d\epsilon} \Big|_{\epsilon=0} \right\} &= \text{tr} \left(\Sigma_{X_K}^{-1} \Sigma_{X_K Y_{\mathcal{J}}} \Sigma_{Y_{\mathcal{J}}}^{-1} \Sigma_{Y_{\mathcal{J}} X_K} \Sigma_{X_K}^{-1} \cdot \mathbb{E}_P [\{X_K - \mathbb{E}_P(X_K)\} \{X_K - \mathbb{E}_P(X_K)\}^T] \right) \\ &\quad - \text{tr} \left(\Sigma_{Y_{\mathcal{J}}}^{-1} \Sigma_{Y_{\mathcal{J}} X_K} \Sigma_{X_K}^{-1} \Sigma_{X_K Y_{\mathcal{J}}} \Sigma_{Y_{\mathcal{J}}}^{-1} \cdot \mathbb{E}_P [\{Y_{\mathcal{J}} - \mathbb{E}_P(Y_{\mathcal{J}})\} \{Y_{\mathcal{J}} - \mathbb{E}_P(Y_{\mathcal{J}})\}^T] \right) \\ &\quad + 2 \text{tr} \left(\Sigma_{Y_{\mathcal{J}}}^{-1} \Sigma_{Y_{\mathcal{J}} X_K} \Sigma_{X_K}^{-1} \cdot \mathbb{E}_P [\{X_K - \mathbb{E}_P(X_K)\} \{Y_{\mathcal{J}} - \mathbb{E}_P(Y_{\mathcal{J}})\}^T] \right) \\ &= -\Phi^d(P) - \Phi^d(P) + 2\Phi^d(P) = 0. \end{aligned}$$

120 When $\Psi^d(P) = 0$, using instead the expression (12) for the canonical gradient and again applying (S.1) we obtain

$$\begin{aligned} \mathbb{E}_P \{D^d(P)(O)\} &= \mathbb{E}_P \left[\{Y_{\mathcal{J}} - \mathbb{E}_P(Y_{\mathcal{J}})\}^T \Sigma_{Y_{\mathcal{J}}}^{-1/2} L \Sigma_{X_K}^{-1/2} \{X_K - \mathbb{E}_P(X_K)\} \right] \\ &= \text{tr} \left(L \Sigma_{X_K}^{-1/2} \mathbb{E}_P [\{X_K - \mathbb{E}_P(X_K)\} \{Y_{\mathcal{J}} - \mathbb{E}_P(Y_{\mathcal{J}})\}^T] \Sigma_{Y_{\mathcal{J}}}^{-1/2} \right) \\ &= \text{tr}(L \Lambda_{X_K Y_{\mathcal{J}}}) = 0 \end{aligned}$$

since $\Lambda_{X_K Y_{\mathcal{J}}} = 0$ when $\Psi^d(P) = 0$.

S.4.3. Derivation of the remainder term

The second order remainder term is defined as,

$$\text{Rem}(P) = \Psi^d(P) - \Psi^d(P_0) + \int D^d(P)(o) dP_0(o), \quad (\text{S.2})$$

125 where we later use P to denote the sample empirical distribution and P_0 to denote the population distribution. In the following calculations, we omit the superscript d since it is fixed. To simplify the calculation, it is convenient to study $\Phi(P) = \text{tr}(S_X^{-1} S_{XY} S_Y^{-1} S_{XY}^T)$ and $\Phi(P_0) = \text{tr}(\Sigma_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} \Sigma_{XY}^T)$ instead of doing the calculation in the $\Psi = \Phi^{1/2}$ scale. As such, we have

$$\begin{aligned} \int D(P)(o) dP_0(o) &= -\text{tr} \{S_X^{-1} S_{XY} S_Y^{-1} S_{YX} S_X^{-1} P_0 [\{x - P(X)\} \{x^T - P(X^T)\}]\} \\ &\quad - \text{tr} \{S_Y^{-1} S_{YX} S_X^{-1} S_{XY} S_Y^{-1} P_0 [\{y - P(Y)\} \{y^T - P(Y^T)\}]\} \\ &\quad + 2 \text{tr} \{S_Y^{-1} S_{YX} S_X^{-1} P_0 [\{x - P(X)\} \{y^T - P(Y^T)\}]\}. \end{aligned}$$

We can substitute the $P_0[\dots]$ terms as follows,

$$P_0 [\{x - P(X)\} \{y^T - P(Y^T)\}] = \Sigma_{XY} - (\mu_{X,0} - \mu_X)(\mu_{Y,0} - \mu_Y)^T,$$

130 where we simplified notation as $\mu_X = P(X)$, $\mu_Y = P(Y)$, $\mu_{X,0} = P_0(X)$ and $\mu_{Y,0} = P_0(Y)$. We decompose the term $\int D(P)(o) dP_0(o) = T_1 + T_2$ into two parts accordingly. The first part is $O_P(1/n)$ if μ_X and μ_Y are close to $\mu_{X,0}$ and $\mu_{Y,0}$ in the order of $O_p(n^{-1/2})$:

$$\begin{aligned} T_1 &= 2 \text{tr} \{S_Y^{-1} S_{YX} S_X^{-1} (\mu_{X,0} - \mu_X)(\mu_{Y,0} - \mu_Y)^T\} \\ &\quad - \text{tr} \{S_X^{-1} S_{XY} S_Y^{-1} S_{YX} S_X^{-1} (\mu_{X,0} - \mu_X)(\mu_{X,0} - \mu_X)^T\} \\ &\quad - \text{tr} \{S_Y^{-1} S_{YX} S_X^{-1} S_{XY} S_Y^{-1} (\mu_{Y,0} - \mu_Y)(\mu_{Y,0} - \mu_Y)^T\}. \end{aligned}$$

The second part is small when Σ 's are close to Σ_0 's, but not small enough (i.e. $O_P(n^{-1/2})$ instead of $O_P(n^{-1})$) and will be combined with $\Phi(P) - \Phi(P_0)$ later:

$$\begin{aligned} T_2 &= 2\text{tr} (S_Y^{-1} S_{YX} S_X^{-1} \Sigma_{XY}) \\ &\quad - \text{tr} (S_X^{-1} S_{XY} S_Y^{-1} S_{YX} S_X^{-1} \Sigma_X) - \text{tr} (S_Y^{-1} S_{YX} S_X^{-1} S_{XY} S_Y^{-1} \Sigma_Y) \\ &= 2\text{tr} \{S_Y^{-1} S_{YX} S_X^{-1} (\Sigma_{XY} - S_{XY})\} + \text{tr} \{S_X^{-1} S_{XY} S_Y^{-1} S_{YX} S_X^{-1} (S_X - \Sigma_X)\} \\ &\quad + \text{tr} \{S_Y^{-1} S_{YX} S_X^{-1} S_{XY} S_Y^{-1} (S_Y - \Sigma_Y)\}. \end{aligned}$$

In $\text{Rem}(P)$, we express $\Phi(P) - \Phi(P_0)$ as follows,

135

$$\begin{aligned} \Phi(P) - \Phi(P_0) &= \text{tr}(S_X^{-1} S_{XY} S_Y^{-1} S_{YX} - \Sigma_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} \Sigma_{YX}) \\ &= \text{tr} \{ \Sigma_{XY} \Sigma_Y^{-1} \Sigma_{YX} (S_X^{-1} - \Sigma_X^{-1}) \} + \text{tr} \{ \Sigma_{YX} S_X^{-1} \Sigma_{XY} (S_Y^{-1} - \Sigma_Y^{-1}) \} \\ &\quad + \text{tr} \{ S_X^{-1} \Sigma_{XY} S_Y^{-1} (S_{YX} - \Sigma_{YX}) \} + \text{tr} \{ S_Y^{-1} S_{YX} S_X^{-1} (S_{XY} - \Sigma_{XY}) \}, \end{aligned}$$

where the first two part can be re-organized by noticing $(S_X^{-1} - \Sigma_X^{-1}) = S_X^{-1} (\Sigma_X - S_X) \Sigma_X^{-1}$. Hence, we have

$$\begin{aligned} \Phi(P) - \Phi(P_0) &= \text{tr} \{ \Sigma_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} \Sigma_{YX} S_X^{-1} (\Sigma_X - S_X) \} \\ &\quad + \text{tr} \{ S_Y^{-1} \Sigma_{YX} S_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} (\Sigma_Y - S_Y) \} \\ &\quad + \text{tr} \{ S_X^{-1} \Sigma_{XY} S_Y^{-1} (S_{YX} - \Sigma_{YX}) \} + \text{tr} \{ S_Y^{-1} S_{YX} S_X^{-1} (S_{XY} - \Sigma_{XY}) \}. \end{aligned}$$

Adding the various terms, we have

$$\begin{aligned} \text{Rem}(P) &= T_1 + T_2 + \Phi(P) - \Phi(P_0) \\ &= 2\text{tr} \{ S_Y^{-1} S_{YX} S_X^{-1} (\mu_{X,0} - \mu_X) (\mu_{Y,0} - \mu_Y)^T \} \\ &\quad - \text{tr} \{ S_X^{-1} S_{XY} S_Y^{-1} S_{YX} S_X^{-1} (\mu_{X,0} - \mu_X) (\mu_{X,0} - \mu_X)^T \} \\ &\quad - \text{tr} \{ S_Y^{-1} S_{YX} S_X^{-1} S_{XY} S_Y^{-1} (\mu_{Y,0} - \mu_Y) (\mu_{Y,0} - \mu_Y)^T \} \\ &\quad + \text{tr} \{ S_Y^{-1} (\Sigma_{YX} - S_{YX}) S_X^{-1} (S_{XY} - \Sigma_{XY}) \} \\ &\quad + \text{tr} \{ (\Sigma_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} \Sigma_{YX} - S_X^{-1} S_{XY} S_Y^{-1} S_{YX}) S_X^{-1} (\Sigma_X - S_X) \} \\ &\quad + \text{tr} \{ S_Y^{-1} (\Sigma_{YX} S_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} - S_{YX} S_X^{-1} S_{XY} S_Y^{-1}) (\Sigma_Y - S_Y) \}, \end{aligned}$$

which leads to the following final expression after expanding the last two terms,

$$\begin{aligned}
\text{Rem}(P) &= T_1 + T_2 + \Phi(P) - \Phi(P_0) \\
&= 2\text{tr} \{ S_Y^{-1} S_{YX} S_X^{-1} (\mu_{X,0} - \mu_X) (\mu_{Y,0} - \mu_Y)^T \} \\
&\quad - \text{tr} \{ S_X^{-1} S_{XY} S_Y^{-1} S_{YX} S_X^{-1} (\mu_{X,0} - \mu_X) (\mu_{X,0} - \mu_X)^T \} \\
&\quad - \text{tr} \{ S_Y^{-1} S_{YX} S_X^{-1} S_{XY} S_Y^{-1} (\mu_{Y,0} - \mu_Y) (\mu_{Y,0} - \mu_Y)^T \} \\
&\quad + \text{tr} \{ S_Y^{-1} (\Sigma_{YX} - S_{YX}) S_X^{-1} (S_{XY} - \Sigma_{XY}) \} \\
&\quad + \text{tr} \{ \Sigma_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} (\Sigma_{YX} - S_{YX}) S_X^{-1} (\Sigma_X - S_X) \} \\
&\quad + \text{tr} \{ \Sigma_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} (S_Y - \Sigma_Y) S_Y^{-1} S_{YX} S_X^{-1} (\Sigma_X - S_X) \} \\
&\quad + \text{tr} \{ \Sigma_X^{-1} (\Sigma_{XY} - S_{XY}) S_Y^{-1} S_{YX} S_X^{-1} (\Sigma_X - S_X) \} \\
&\quad + \text{tr} \{ \Sigma_X^{-1} (S_X - \Sigma_X) S_X^{-1} S_{XY} S_Y^{-1} S_{YX} S_X^{-1} (\Sigma_X - S_X) \} \\
&\quad + \text{tr} \{ S_Y^{-1} \Sigma_{YX} S_X^{-1} \Sigma_{XY} \Sigma_Y^{-1} (S_Y - \Sigma_Y) S_Y^{-1} (\Sigma_Y - S_Y) \} \\
&\quad + \text{tr} \{ S_Y^{-1} \Sigma_{YX,0} S_X^{-1} (\Sigma_{XY} - S_{XY}) S_Y^{-1} (\Sigma_Y - S_Y) \} \\
&\quad + \text{tr} \{ S_Y^{-1} (\Sigma_{YX} - S_{YX}) S_X^{-1} S_{XY} S_Y^{-1} (\Sigma_Y - S_Y) \}. \tag{S.3}
\end{aligned}$$

S.4.4. Proof of Lemma 1

Let $\mathbb{X} \in \mathbb{R}^{n \times s_x}$, $\mathbb{Y} \in \mathbb{R}^{n \times s_y}$ and $\mathbb{Z} \in \mathbb{R}^{n \times 1}$ be the centered data matrices of $X_{\mathcal{K}}$, $Y_{\mathcal{J}}$ and X_k . That is, the i -th row of $\mathbb{X}$ is $(X_{\mathcal{K},i} - \bar{X}_{\mathcal{K}})^T$. Note that $S_{X_{\mathcal{K}}} = \mathbb{X}\mathbb{X}^T/n$ is the sample covariance matrix of $X_{\mathcal{K}}$, which is assumed to be positive definite. Then we can write $C_{Y_{\mathcal{J}}X_{\mathcal{K}}} = n^{-1}S_{Y_{\mathcal{J}}}^{-1/2}\mathbb{Y}^T\mathbb{X}S_{X_{\mathcal{K}}}^{-1/2}$ and $C_{Y_{\mathcal{J}}X_{\mathcal{K}}}C_{Y_{\mathcal{J}}X_{\mathcal{K}}}^T = n^{-1}S_{Y_{\mathcal{J}}}^{-1/2}\mathbb{Y}^T\mathbb{X}(\mathbb{X}^T\mathbb{X})^{-1}\mathbb{X}^T\mathbb{Y}S_{Y_{\mathcal{J}}}^{-1/2}$, where $\mathbb{X}(\mathbb{X}^T\mathbb{X})^{-1}\mathbb{X}^T \equiv P_{\mathbb{X}} \in \mathbb{R}^{n \times n}$ is the projection matrix onto the s_x -dimensional subspace spanned by the columns of $\mathbb{X}$. Then,

$$\|C_{Y_{\mathcal{J}}X_{\mathcal{K}}}\|_F^2 = \text{tr}(C_{Y_{\mathcal{J}}X_{\mathcal{K}}}C_{Y_{\mathcal{J}}X_{\mathcal{K}}}^T) = n^{-1}\text{tr}(S_{Y_{\mathcal{J}}}^{-1/2}\mathbb{Y}^TP_{\mathbb{X}}\mathbb{Y}S_{Y_{\mathcal{J}}}^{-1/2}), \tag{S.4}$$

$$\|C_{Y_{\mathcal{J}}X_{\mathcal{K} \cup k}}\|_F^2 - \|C_{Y_{\mathcal{J}}X_{\mathcal{K}}}\|_F^2 = n^{-1}\text{tr}\{S_{Y_{\mathcal{J}}}^{-1/2}\mathbb{Y}^T(P_{(\mathbb{X},\mathbb{Z})} - P_{\mathbb{X}})\mathbb{Y}S_{Y_{\mathcal{J}}}^{-1/2}\}, \tag{S.5}$$

where $P_{(\mathbb{X},\mathbb{Z})} \in \mathbb{R}^{n \times n}$ is the projection matrix onto the $(s_x + 1)$ -dimensional subspace of $\mathbb{R}^n$ spanned by the columns of $(\mathbb{X}, \mathbb{Z})$. That is,

$$\begin{aligned}
P_{(\mathbb{X},\mathbb{Z})} &= (\mathbb{X} \mathbb{Z}) \begin{pmatrix} \mathbb{X}^T\mathbb{X} & \mathbb{X}^T\mathbb{Z} \\ \mathbb{Z}^T\mathbb{X} & \mathbb{Z}^T\mathbb{Z} \end{pmatrix}^{-1} \begin{pmatrix} \mathbb{X}^T \\ \mathbb{Z}^T \end{pmatrix} \\
&= n^{-1} (\mathbb{X} \mathbb{Z}) \begin{pmatrix} S_{X_{\mathcal{K}}}^{-1} + S_{X_{\mathcal{K}}}^{-1} S_{X_{\mathcal{K}}X_k} S_{R_k}^{-1} S_{X_kX_{\mathcal{K}}} S_{X_{\mathcal{K}}}^{-1} - S_{X_{\mathcal{K}}}^{-1} S_{X_{\mathcal{K}}X_k} S_{R_k}^{-1} \\ -S_{R_k}^{-1} S_{X_kX_{\mathcal{K}}} S_{X_{\mathcal{K}}}^{-1} & S_{R_k}^{-1} \end{pmatrix} \begin{pmatrix} \mathbb{X}^T \\ \mathbb{Z}^T \end{pmatrix} \\
&= n^{-1} \mathbb{X} S_{X_{\mathcal{K}}}^{-1} \mathbb{X}^T + n^{-1} (\mathbb{Z} - \mathbb{X} S_{X_{\mathcal{K}}}^{-1} S_{X_{\mathcal{K}}X_k}) S_{R_k}^{-1} (\mathbb{Z}^T - S_{X_kX_{\mathcal{K}}} S_{X_{\mathcal{K}}}^{-1} \mathbb{X}^T), \\
&= P_{\mathbb{X}} + P_{\mathbb{Z}|\mathbb{X}},
\end{aligned}$$

where $S_{R_k} = n^{-1}\{\mathbb{Z}^T\mathbb{Z} - \mathbb{Z}^T\mathbb{X}(\mathbb{X}^T\mathbb{X})^{-1}\mathbb{X}^T\mathbb{Z}\}$ is the sample covariance of the fitted residual $(\mathbb{Z} - \mathbb{X}S_{X_{\mathcal{K}}}^{-1}S_{X_{\mathcal{K}}X_k})$. Specifically, $(\mathbb{Z} - \mathbb{X}S_{X_{\mathcal{K}}}^{-1}S_{X_{\mathcal{K}}X_k})^T(\mathbb{Z} - \mathbb{X}S_{X_{\mathcal{K}}}^{-1}S_{X_{\mathcal{K}}X_k}) = nS_{R_k}$. Therefore, $P_{\mathbb{Z}|\mathbb{X}} \in \mathbb{R}^{n \times n}$ is indeed the projection matrix onto the one-dimensional subspace spanned by $(\mathbb{Z} - \mathbb{X}S_{X_{\mathcal{K}}}^{-1}S_{X_{\mathcal{K}}X_k})$. The conclusion follows from (S.4), (S.5) and $P_{(\mathbb{X},\mathbb{Z})} = P_{\mathbb{X}} + P_{\mathbb{Z}|\mathbb{X}}$.

S.5. STABILIZED ONE-STEP ESTIMATOR

S.5.1. Recursive computation of the one-step estimator 155

In the implementation, given a new observation O_{j+1} , we need to update the weight $w_j = \bar{\sigma}_n / \hat{\sigma}_j$ and the variance of the canonical gradient $\hat{\sigma}_j^2$ in order to update the current estimate ψ_j of the target parameter $\tau_{\max}$. Since we only need to know these quantities when all observations are included, an efficient approach is to exploit the following recursive properties of $\psi_j / \bar{\sigma}_j$ and $\bar{\sigma}_j^{-1}$: 160

$$\frac{\psi_{j+1}}{\bar{\sigma}_{j+1}} = \frac{1}{j+1} \left\{ j \cdot \frac{\psi_j}{\bar{\sigma}_j} + \frac{\Psi^{d_{nj}}(P_j) + D^{d_j}(P_j)(O_{j+1})}{\hat{\sigma}_{nj}} \right\}, \quad \bar{\sigma}_{j+1}^{-1} = \frac{1}{j+1} \left(j \cdot \bar{\sigma}_j^{-1} + \hat{\sigma}_{nj}^{-1} \right).$$

We then obtain the final estimate when $j = n$:

$$\hat{\tau}_{\max} \equiv \psi_n = \frac{\bar{\sigma}_n}{n - \ell_n} \sum_{j=\ell_n}^{n-1} \left\{ \frac{\Psi^{d_{nj}}(P_j) + D^{d_j}(P_j)(O_{j+1})}{\hat{\sigma}_{nj}} \right\}, \quad \bar{\sigma}_n^{-1} = \frac{1}{n - \ell_n} \sum_{j=\ell_n}^{n-1} \hat{\sigma}_{nj}^{-1}.$$

S.5.2. Proofs of Theorems 2 and 3

The proofs are analogous to those of Theorems A.10 and A.11 in Luedtke & van der Laan (2018), which will henceforth be abbreviated to LL2018, establishing the asymptotic validity of the confidence interval (extending Theorem 1 of LL2018). In what follows, we first state various lemmas that together give all the ingredients we need. 165

Recall that in Theorems 2 and 3 we assume each variable X_k , $k = 1, \dots, p$, or Y_j , $j = 1, \dots, q$, takes values within $[-1, 1]$. The other conditions are as follows. The canonical gradient satisfies for some constant $\gamma > 0$

$$\inf_{n \geq 2} \min_{d \in \mathcal{D}_n} \text{var}_{P_0} \{D^d(P_0)(O)\} \geq \gamma.$$

To ensure that the canonical gradient is uniformly bounded for all d , we assume that, for some $\delta > 0$,

$$\sup_{d \in \mathcal{D}_n} \max \{ \|\Sigma_{X_K}^{-1}\|, \|\Sigma_{Y_J}^{-1}\| \} < \delta^{-1},$$

where $\|M\|$ denotes the operator norm. We treat δ , γ and $s = s_x = s_y$ as fixed when $n \rightarrow \infty$, and thus omit the dependence on δ , γ and s in the asymptotic statements. On the other hand, we allow both the dimensions p and q to grow quickly with the sample size n . We define $\beta_n^2 = n^{-1/2} \cdot \log \max\{n, p, q\}$ and, for any $\epsilon \in (0, 2)$, assume

$$\frac{\log \max(n, p, q)}{\ell_n} \rightarrow 0, \quad \beta_n^2 \log \frac{n}{\ell_n} \rightarrow 0, \quad \limsup_{n \rightarrow \infty} \frac{\ell_n}{n} < 1, \quad (\text{S.6})$$

where

$$\ell_n = \max \left\{ (\log \max(n, p, q))^{1+\epsilon}, n \exp(-\beta_n^{-2+\epsilon}) \right\}.$$

For Theorem 3, we further assume the following margin condition. For some sequence $t_n \rightarrow \infty$, there exists a sequence of nonempty subsets $\mathcal{D}_n^* \subseteq \mathcal{D}_n$ such that, for all n , 170

$$\sup_{d_1, d_2 \in \mathcal{D}_n^*} \left\{ \Psi^{d_1}(P_0) - \Psi^{d_2}(P_0) \right\} = o(n^{-1/2}), \quad \inf_{d \in \mathcal{D}_n^*} \Psi^d(P_0) - \sup_{d \in \mathcal{D}_n \setminus \mathcal{D}_n^*} \Psi^d(P_0) \geq t_n n^{-1/2} \beta_n.$$

Replacing X and Y above by X_K and Y_J , the remainder term $\text{Rem}^d(P)$ given by (S.3) can be bounded as follows.

LEMMA S.1. For $d \in \mathcal{D}_n$, define

$$M_1(P, d) \equiv \max_{j \in \mathcal{J}, k \in \mathcal{K}} \{|\mathbb{E}_P(Y_j) - \mathbb{E}_{P_0}(Y_j)|, |\mathbb{E}_P(X_k) - \mathbb{E}_{P_0}(X_k)|\};$$

$$M_2(P, d) \equiv \max_{j \in \mathcal{J}, k \in \mathcal{K}} \{|\text{var}_P(Y_j) - \text{var}_{P_0}(Y_j)|, |\text{var}_P(X_k) - \text{var}_{P_0}(X_k)|, |\text{cov}_P(X_k, Y_j) - \text{cov}_{P_0}(X_k, Y_j)|\}.$$

175 Suppose $\tilde{\delta}^{-1}$ and λ are some positive constants such that, for all $d \in \mathcal{D}_n$,

$$\max \left(\|S_{X\mathcal{K}}^{-1}\|, \|S_{Y\mathcal{J}}^{-1}\| \right) < \tilde{\delta}^{-1}, \quad \max \left\{ \|\Sigma_{Y\mathcal{J}}^{1/2} S_{Y\mathcal{J}}^{-1} \Sigma_{Y\mathcal{J}}^{1/2}\|, \|\Sigma_{X\mathcal{K}}^{1/2} S_{X\mathcal{K}}^{-1} \Sigma_{X\mathcal{K}}^{1/2}\| \right\} < \lambda. \quad (\text{S.7})$$

Then the second-order remainder term (S.2) is bounded by

$$|\text{Rem}^d(P)| < \tilde{\delta}^{-1} 2(s_x + s_y) M_1(P, d) + \{2\tilde{\delta}^{-1} + (4 + 2\lambda^2)(\delta\tilde{\delta})^{-1}\} s_x s_y M_2^2(P, d) \\ \leq \tilde{\delta}^{-1} 4s M_1(P, d) + \{2\tilde{\delta}^{-1} + (4 + 2\lambda^2)(\delta\tilde{\delta})^{-1}\} s^2 M_2^2(P, d).$$

Let $\mathcal{F}_n$ denote the following class of functions mapping $\mathbb{R}^\infty \times \mathbb{R}^\infty$ into the real line:

$$\left\{ (x, y) \mapsto x_k^a y_j^b : 0 \leq a, b \leq 4, a + b \leq 4, k \in \{1, \dots, p\}, j \in \{1, \dots, q\} \right\}. \quad (\text{S.8})$$

Then $|\mathcal{F}_n| \lesssim pq$, where $\lesssim$ denotes less than or equal to up to a universal multiplicative constant.

We define the empirical process for $f \in \mathcal{F}_n$ and based on the first $j \in \{1, \dots, n\}$ observations

180 as

$$\mathbb{G}_{nj} \equiv j^{-1/2} \sum_{i=1}^j \{f(O_i) - P_0 f\} = j^{1/2} (P_j - P_0) f,$$

where P_j is the empirical distribution of $O_1, \dots, O_j$ and $Pf = \mathbb{E}_P\{f(O)\}$. Let $\|\mathbb{G}_{nj}\|_{\mathcal{F}_n} \equiv \sup_{f \in \mathcal{F}_n} |\mathbb{G}_{nj}|$, then by van der Vaart & Wellner (1996, Theorem 2.14.1)

$$\mathbb{E}\|\mathbb{G}_{nj}\|_{\mathcal{F}_n} \lesssim (\log pq)^{1/2} = (\log p + \log q)^{1/2} \leq (2 \log \max\{p, q\})^{1/2}.$$

For all $j = 1, \dots, n$, define the events

$$\mathcal{A}_{nj} = \left\{ \max_{f \in \mathcal{F}_n} |(P_j - P_0)f| \leq CK_{nj} \right\}, \quad \mathcal{A}_n = \bigcap_{j=1}^n \mathcal{A}_{nj}, \quad K_{nj} = j^{-1/2} (\log \max\{n, p, q\})^{1/2}, \quad (\text{S.9})$$

where C in the definition is the smallest universal constant satisfying

$$\mathbb{E}\|\mathbb{G}_{nj}\|_{\mathcal{F}_n} \leq (C - 1)(2 \log \max\{p, q\})^{1/2}.$$

LEMMA S.2. For any sample size n , the event $\mathcal{A}_n$ occurs with probability at least $1 - 1/n$.

185 Next we show that the various estimators are close to their population versions when $\mathcal{A}_n$ occurs.

LEMMA S.3. For any fixed sample size $n \geq 2$, the event $\mathcal{A}_n$ implies the following statements for $j = 2, \dots, n$, and for any $d \in \mathcal{D}_n$,

$$\max \{M_1(P_j, d), M_2(P_j, d)\} \lesssim K_{nj}, \quad (\text{S.10})$$

where $M_1(P, d)$ and $M_2(P, d)$ are defined in Lemma S.1.

190 LEMMA S.4. Let C be the smallest universal constant in (S.10) and let n be any natural number satisfying $n \geq \lceil 4C^2 \delta^{-2} s^2 \log \max\{n, p, q\} \rceil \equiv J(n, \delta)$. The occurrence of event $\mathcal{A}_n$

implies that, for all $j = J(n, \delta), \dots, n$, for any $d \in \mathcal{D}_n$,

$$\max \left\{ \|\Sigma_{Y_{\mathcal{J}}}^{1/2} S_{Y_{\mathcal{J}}}^{-1} \Sigma_{Y_{\mathcal{J}}}^{1/2}\|, \|\Sigma_{X_{\mathcal{K}}}^{1/2} S_{X_{\mathcal{K}}}^{-1} \Sigma_{X_{\mathcal{K}}}^{1/2}\| \right\} \leq 2; \quad (\text{S.11})$$

$$\max \left(\|S_{X_{\mathcal{K}}}^{-1}\|, \|S_{Y_{\mathcal{J}}}^{-1}\|, \|S_{X_{\mathcal{K}}}^{-1} S_{X_{\mathcal{K}} Y_{\mathcal{J}}} S_{Y_{\mathcal{J}}}^{-1}\|, \|S_{X_{\mathcal{K}}}^{-1} S_{X_{\mathcal{K}} Y_{\mathcal{J}}} S_{Y_{\mathcal{J}}}^{-1} S_{Y_{\mathcal{J}} X_{\mathcal{K}}} S_{X_{\mathcal{K}}}^{-1}\| \right) \leq 2\delta^{-1}; \quad (\text{S.12})$$

$$|\widehat{D}_{nj}(o)| \leq 2(s_x^{1/2} + s_y^{1/2})^2 \delta^{-1} \leq 8\delta^{-1}s; \quad (\text{S.13})$$

and the estimator $\tau^2(P_j)$ of the Pillai trace satisfies

$$|\tau^2 - \tau^2(P_j)| \equiv |\Psi_n(P_0) - \Psi^{d_{nj}}(P_j)| \lesssim \delta^{-2} s^2 K_{nj}. \quad (\text{S.14})$$

LEMMA S.5. Under the same conditions as in Lemma S.4, the occurrence of event $\mathcal{A}_n$ implies that, for all $j = J(n, \delta), \dots, n$, for any $d \in \mathcal{D}_n$,

$$|\widehat{\sigma}_{nj}^2 - \sigma_{nj}^2| \leq \delta^{-2} s^2 K_{nj}; \quad (\text{S.15})$$

$$|\widehat{\sigma}_{nj}^2 - \text{var}_{P_0}\{D_{nj}(P_0)(O)\}| \leq \delta^{-4} s^2 K_{nj} + \left(\widehat{\text{Rem}}_{nj}\right)^2; \quad (\text{S.16})$$

$$|\widehat{\text{Rem}}_{nj}| \equiv |\text{Rem}^{d_{nj}}(P_j)| \lesssim \delta^{-1}(s_x + s_y) K_{nj}^2 + \delta^{-2} s_x s_y K_{nj}^2 \lesssim (\delta^{-1} s + \delta^{-2} s^2) K_{nj}^2. \quad (\text{S.17})$$

LEMMA S.6. Let $\gamma > 0$ be as defined in (14). For a constant $C(\delta, \gamma, s) > 0$ relying on γ, δ and s only, the occurrence of event $\mathcal{A}_n$ implies that, for all $j = \lceil C(\delta, \gamma, s) \log \max(n, p, q) \rceil, \dots, n$,

$$\widehat{\sigma}_{nj}^2 \geq \gamma/2. \quad (\text{S.18})$$

In what follows, we treat δ, γ and s as fixed and omit them in the asymptotic statements. Recall that

$$\beta_n^2 = n^{-1/2} \log(p + q), \quad \ell_n = \max \left\{ (\log \max(n, p, q))^{1+\epsilon}, n \exp(-\beta_n^{-2+\epsilon}) \right\}.$$

LEMMA S.7. The following properties hold as $n \rightarrow \infty$

(C1) There exists $M < \infty$ such that

$$\frac{1}{n - \ell_n} \sum_{j=\ell_n}^{n-1} P_0 \left(\frac{|\widehat{D}_{nj}(O)|}{\widehat{\sigma}_{nj}^2} < M \mid O_0, \dots, O_{j-1} \right) \rightarrow 1$$

in probability.

(C2) $(n - \ell_n)^{-1} \sum_{j=\ell_n}^{n-1} \left| \sigma_{nj}^2 / \widehat{\sigma}_{nj}^2 - 1 \right| \rightarrow 0$ in probability.

(C3) $(n - \ell_n)^{-1/2} \sum_{j=\ell_n}^{n-1} \widehat{\sigma}_{nj}^{-1} \widehat{\text{Rem}}_{nj} \rightarrow 0$ in probability, where $\widehat{\text{Rem}}_{nj}$ is from Lemma S.5.

S.5.3. Proof of Lemma S.1

Under the assumption (15), we have the following immediate implications for any $\mathcal{K}$ and $\mathcal{J}$:

1. $\|\Sigma_{X_{\mathcal{K}}}^{-1}\| < \delta^{-1}$ and $\|\Sigma_{Y_{\mathcal{J}}}^{-1}\| < \delta^{-1}$;
2. $\|\Sigma_{X_{\mathcal{K}}}^{-1/2}\| < \delta^{-1/2}$ and $\|\Sigma_{Y_{\mathcal{J}}}^{-1/2}\| < \delta^{-1/2}$;
3. $\|\Sigma_{X_{\mathcal{K}}}^{-1} \Sigma_{X_{\mathcal{K}} Y_{\mathcal{J}}} \Sigma_{Y_{\mathcal{J}}}^{-1}\| < \delta^{-1}$;

$$4. \|\Sigma_{Y\mathcal{J}}^{-1}\Sigma_{Y\mathcal{J}X\mathcal{K}}\Sigma_{X\mathcal{K}}^{-1}\Sigma_{X\mathcal{K}Y\mathcal{J}}\Sigma_{Y\mathcal{J}}^{-1}\| < \delta^{-1}.$$

For the last two statements, we notice that $\|\Sigma_{X\mathcal{K}}^{-1/2}\Sigma_{X\mathcal{K}Y\mathcal{J}}\Sigma_{Y\mathcal{J}}^{-1/2}\| \leq 1$ because the singular values of $\Lambda_{XY} = \Sigma_X^{-1/2}\Sigma_{XY}\Sigma_Y^{-1/2}$ are no larger than 1.

The first three terms in (S.3), and the remainder term $\text{Rem}^d(P)$, can be bounded as follows.

215 For notational convenience, we omit the subscripts in $X_{\mathcal{K}}$ and $Y_{\mathcal{J}}$. Let $\delta_X = S_X^{-1/2}(\mu_{X,0} - \mu_X)$ and $\delta_Y = S_Y^{-1/2}(\mu_{Y,0} - \mu_Y)$, and note that the singular values of $C_{XY} = S_X^{-1/2}S_{XY}S_Y^{-1/2}$ are less than 1. Then the absolute value of the first three terms in (S.3) becomes

$$\begin{aligned} |2\delta_X^T C_{XY} \delta_Y - \delta_Y^T C_{XY}^T C_{XY} \delta_Y - \delta_X^T C_{XY} C_{XY}^T \delta_X| &\leq 2|\delta_X^T \delta_Y| + \|\delta_Y\|^2 + \|\delta_X\|^2 \\ &\leq 2\|\delta_X\|\|\delta_Y\| + \|\delta_Y\|^2 + \|\delta_X\|^2 \\ &< \tilde{\delta}^{-1} \cdot (s_x + s_y + 2s_x^{1/2}s_y^{1/2}) \cdot \{M_1(P, d)\}^2 \\ &\leq \tilde{\delta}^{-1} \cdot 4s \cdot \{M_1(P, d)\}^2. \end{aligned}$$

To bound the fourth term in (S.3),

$$\begin{aligned} |\text{tr}\{S_Y^{-1}(\Sigma_{YX} - S_{YX})S_X^{-1}(S_{XY} - \Sigma_{XY})\}| &\leq \|S_X^{-1}\|\|S_Y^{-1}\|\|S_{XY} - \Sigma_{XY}\|_F^2 \\ &< \tilde{\delta}^{-2}s_x s_y \cdot \{M_2(P, d)\}^2 \\ &\leq \tilde{\delta}^{-2}s^2 \cdot \{M_2(P, d)\}^2. \end{aligned}$$

Similarly, for the other terms in (S.3), we have

$$\begin{aligned} \|S_X^{-1} \otimes \Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\| &< (\tilde{\delta}\delta)^{-1}, \\ \|S_Y^{-1}S_{YX}S_X^{-1} \otimes \Sigma_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\| &< (\tilde{\delta}\delta)^{-1}, \\ \|S_Y^{-1}S_{YX}S_X^{-1} \otimes \Sigma_X^{-1}\| &< (\tilde{\delta}\delta)^{-1}, \\ \|S_X^{-1}S_{XY}S_Y^{-1}S_{YX}S_X^{-1} \otimes \Sigma_X^{-1}\| &< (\tilde{\delta}\delta)^{-1}, \\ \|S_X^{-1}S_{XY}S_Y^{-1} \otimes S_Y^{-1}\| &< \tilde{\delta}^{-2}. \end{aligned}$$

220 Further, $\|S_Y^{-1} \otimes S_Y^{-1}\Sigma_{YX}S_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\|$ is bounded above by $\tilde{\delta}^{-1}\|S_Y^{-1}\Sigma_{YX}S_X^{-1}\Sigma_{XY}\Sigma_Y^{-1}\|$, which equals

$$\begin{aligned} &\tilde{\delta}^{-1}\|\Sigma_Y^{-1/2}S_Y^{-1}\Sigma_Y^{1/2} \cdot \Sigma_Y^{-1/2}\Sigma_{YX}\Sigma_X^{-1/2} \cdot \Sigma_X^{1/2}S_X^{-1}\Sigma_X^{1/2} \cdot \Sigma_X^{-1/2}\Sigma_{XY}\Sigma_Y^{-1/2}\| \\ &\leq \tilde{\delta}^{-1}\|\Sigma_Y^{-1/2}S_Y^{-1}\Sigma_Y^{1/2}\| \cdot \|\Sigma_X^{1/2}S_X^{-1}\Sigma_X^{1/2}\| \\ &\leq \tilde{\delta}^{-1}\delta^{-1}\|\Sigma_Y^{1/2}S_Y^{-1}\Sigma_Y^{1/2}\| \cdot \|\Sigma_X^{1/2}S_X^{-1}\Sigma_X^{1/2}\| \\ &< (\tilde{\delta}\delta)^{-1}\lambda^2. \end{aligned}$$

The term $\|S_Y^{-1} \otimes S_Y^{-1}\Sigma_{YX}S_X^{-1}\|$ is bounded above by

$$\begin{aligned} \tilde{\delta}^{-1}\|S_Y^{-1}\Sigma_Y^{1/2} \cdot \Sigma_Y^{-1/2}\Sigma_{YX}\Sigma_X^{-1/2} \cdot \Sigma_X^{1/2}S_X^{-1}\| &\leq \tilde{\delta}^{-1}\|S_Y^{-1}\Sigma_Y^{1/2}\| \cdot \|\Sigma_X^{1/2}S_X^{-1}\| \\ &\leq \tilde{\delta}^{-1}\delta^{-1}\|\Sigma_Y^{1/2}S_Y^{-1}\Sigma_Y^{1/2}\| \cdot \|\Sigma_X^{1/2}S_X^{-1}\Sigma_X^{1/2}\| \\ &< (\tilde{\delta}\delta)^{-1}\lambda^2. \end{aligned}$$

Adding all the terms together completes the proof.

S.5.4. Proof of Lemma S.2

Omitted since it is analogous to the proof of Lemma A.4 in LL2018.

S.5.5. Proof of Lemma S.3

First recall the definitions of $M_1(P_j, d)$ and $M_2(P_j, d)$ given in Lemma S.1. By definition,

$$\max_{1 \leq k \leq p} |\mathbb{E}_{P_j}(X_k) - \mathbb{E}_{P_0}(X_k)| \leq \max_{f \in \mathcal{F}_n} |(P_j - P_0)f| \lesssim K_{nj}$$

holds for all $j = 2, \dots, n$. Similarly,

$$\max_{1 \leq i \leq q} |\mathbb{E}_{P_j}(Y_i) - \mathbb{E}_{P_0}(Y_i)| \lesssim K_{nj}, \quad \max_{1 \leq k \leq p} |\text{var}_{P_j}(X_k) - \text{var}_{P_0}(X_k)| \lesssim K_{nj}$$

and $\max_{1 \leq i \leq q} |\text{var}_{P_j}(Y_i) - \text{var}_{P_0}(Y_i)| \lesssim K_{nj}$. Therefore $M_1(P_j, d) \lesssim K_{nj}$ for all j and d . To see $M_2(P_j, d) \lesssim K_{nj}$, we notice that $|\text{cov}_{P_j}(X_k, Y_i) - \text{cov}_{P_0}(X_k, Y_i)|$ is bounded above by

$$|\mathbb{E}_{P_j}(X_k Y_i) - \mathbb{E}_{P_0}(X_k Y_i)| + |\mathbb{E}_{P_j}(X_k)| \cdot |(\mathbb{E}_{P_j} - \mathbb{E}_{P_0})Y_i| + |\mathbb{E}_{P_0}(Y_i)| \cdot |(\mathbb{E}_{P_j} - \mathbb{E}_{P_0})X_k| \lesssim K_{nj}.$$

S.5.6. Proof of Lemma S.4

First we show $\|\Sigma_{X_K}^{1/2} S_{X_K}^{-1} \Sigma_{X_K}^{1/2}\| \leq 2$. Note that

230

$$\begin{aligned} \|\Sigma_{X_K}^{1/2} S_{X_K}^{-1} \Sigma_{X_K}^{1/2}\| &= \left\| \Sigma_{X_K}^{-1/2} (\Sigma_{X_K} + S_{X_K} - \Sigma_{X_K}) \Sigma_{X_K}^{-1/2} \right\|^{-1} \\ &= \left\| I_{s_n} + \Sigma_{X_K}^{-1/2} (S_{X_K} - \Sigma_{X_K}) \Sigma_{X_K}^{-1/2} \right\|^{-1} \\ &= \left\{ 1 - \|\Sigma_{X_K}^{-1/2} (S_{X_K} - \Sigma_{X_K}) \Sigma_{X_K}^{-1/2}\| \right\}^{-1} \\ &\leq \left\{ 1 - \|\Sigma_{X_K}^{-1}\| \|S_{X_K} - \Sigma_{X_K}\| \right\}^{-1}, \end{aligned} \quad (\text{S.19})$$

where $\|\Sigma_{X_K}^{-1}\| < \delta^{-1}$ by (15) and $\|S_{X_K} - \Sigma_{X_K}\| \leq \|S_{X_K} - \Sigma_{X_K}\|_1 \leq s_n C K_{nj}$ from (S.10). Moreover, $s_n C K_{nj} \leq s_n C j^{-1/2} (\log \max\{n, p, q\})^{1/2}$. Then for

$$j \geq J(n, \delta) \geq 4C^2 \delta^{-2} \{\max(s_n, t_n)\}^2 \log \max\{n, p, q\}$$

we have

$$\|\Sigma_{X_K}^{-1}\| \|S_{X_K} - \Sigma_{X_K}\| \leq \delta^{-1} s_n C j^{-1/2} (\log \max\{n, p, q\})^{1/2} \leq 1/2,$$

which immediately implies $\|\Sigma_{X_K}^{1/2} S_{X_K}^{-1} \Sigma_{X_K}^{1/2}\| \leq 2$ by (S.19). A similar argument shows that $\|\Sigma_{Y_J}^{1/2} S_{Y_J}^{-1} \Sigma_{Y_J}^{1/2}\| \leq 2$.

Next, to see $\|S_{X_K}^{-1}\| \leq 2\delta^{-1}$, note that

$$\|S_{X_K}^{-1}\| \leq \|\Sigma_{X_K}^{-1} \Sigma_{X_K}\| \|\Sigma_{X_K}^{-1}\| = \|\Sigma_{X_K}^{1/2} S_{X_K}^{-1} \Sigma_{X_K}^{1/2}\| \|\Sigma_{X_K}^{-1}\| \leq 2\delta^{-1}.$$

Similarly, $\|S_{Y_J}^{-1}\| \leq 2\delta^{-1}$. Then, $\|\Sigma_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1}\| \leq \|\Sigma_{X_K}^{-1/2}\| \cdot \|\Sigma_{X_K}^{-1/2} S_{X_K Y_J} S_{Y_J}^{-1/2}\| \cdot$ 235
 $\|S_{Y_J}^{-1/2}\| \leq \|\Sigma_{X_K}^{-1/2}\| \cdot \|\Sigma_{Y_J}^{-1/2}\| \leq 2\delta^{-1}$, and similarly, $\|\Sigma_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1} S_{Y_J X_K} S_{X_K}^{-1}\| \leq 2\delta^{-1}$.

To show $|\hat{D}_{nj}(o)| \leq 2(s_n^{1/2} + t_n^{1/2})^2 \delta^{-1}$, recall from (11) that

$$\begin{aligned} \hat{D}_{nj}(o) &= -\{x_K - \mathbb{E}_P(X_K)\}^T S_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1} S_{Y_J X_K} S_{X_K}^{-1} \{x_K - \mathbb{E}_P(X_K)\} \\ &\quad - \{y_J - \mathbb{E}_P(Y_J)\}^T S_{Y_J}^{-1} S_{Y_J X_K} S_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1} \{y_J - \mathbb{E}_P(Y_J)\} \\ &\quad + 2\{y_J - \mathbb{E}_P(Y_J)\}^T S_{Y_J}^{-1} S_{Y_J X_K} S_{X_K}^{-1} \{x_K - \mathbb{E}_P(X_K)\}. \end{aligned}$$

Because of the bounded support of X_k and Y_i , we have that, under the event $\mathcal{A}_n$,

$$\begin{aligned}
|\widehat{D}_{nj}(o)| &\lesssim s_x \|S_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1} S_{Y_J X_K} S_{X_K}^{-1}\| + s_y \|S_{Y_J}^{-1} S_{Y_J X_K} S_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1}\| \\
&\quad + 2s_x^{1/2} s_y^{1/2} \|S_{Y_J}^{-1} S_{Y_J X_K} S_{X_K}^{-1}\| \\
&\leq \left(s_x + s_y + 2s_x^{1/2} s_y^{1/2} \right) 2\delta^{-1} \\
&\leq 8\delta^{-1} s.
\end{aligned}$$

Turning now to the Pillai trace, we write

$$\begin{aligned}
|\tau_0^2 - \tau^2(P_j)| &= |\text{tr}(\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} \Sigma_{Y_J}^{-1} \Sigma_{Y_J X_K} - S_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1} S_{Y_J X_K})| \\
&\leq |\text{tr}\{\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} \Sigma_{Y_J}^{-1} (\Sigma_{Y_J X_K} - S_{Y_J X_K})\}| \\
&\quad + |\text{tr}\{\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} (\Sigma_{Y_J}^{-1} - S_{Y_J}^{-1}) S_{Y_J X_K}\}| \\
&\quad + |\text{tr}\{\Sigma_{X_K}^{-1} (\Sigma_{Y_J X_K} - S_{Y_J X_K}) S_{Y_J}^{-1} S_{Y_J X_K}\}| \\
&\quad + |\text{tr}\{(\Sigma_{X_K}^{-1} - S_{X_K}^{-1}) S_{X_K Y_J} S_{Y_J}^{-1} S_{Y_J X_K}\}|.
\end{aligned}$$

240 We bound each of the four terms on the right hand side of the above inequality. For the first term,

$$\begin{aligned}
|\text{tr}\{\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} \Sigma_{Y_J}^{-1} (\Sigma_{Y_J X_K} - S_{Y_J X_K})\}| &\leq \|\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} \Sigma_{Y_J}^{-1}\|_F \cdot \|\Sigma_{Y_J X_K} - S_{Y_J X_K}\|_F \\
&\lesssim \min(s_n, t_n) \|\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} \Sigma_{Y_J}^{-1}\| \cdot s_n^{1/2} t_n^{1/2} K_{nj} \\
&\leq \{\max(s_n, t_n)\}^2 \delta^{-1} K_{nj}.
\end{aligned}$$

For the second term,

$$\begin{aligned}
|\text{tr}\{\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} (\Sigma_{Y_J}^{-1} - S_{Y_J}^{-1}) S_{Y_J X_K}\}| &= |\text{tr}\{\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} \Sigma_{Y_J}^{-1} (S_{Y_J} - \Sigma_{Y_J}) S_{Y_J}^{-1} S_{Y_J X_K}\}| \\
&\leq \|\Sigma_{X_K}^{-1} \Sigma_{X_K Y_J} \Sigma_{Y_J}^{-1}\| \cdot \|S_{Y_J}^{-1}\| \cdot \|S_{Y_J} - \Sigma_{Y_J}\|_F \cdot \|S_{Y_J X_K}\|_F \\
&\lesssim \delta^{-1} \cdot 2\delta^{-1} \cdot t_n K_{nj} \cdot s_n^{1/2} t_n^{1/2}, \\
&\lesssim \delta^{-2} \cdot \{\max(s_n, t_n)\}^2 \cdot K_{nj},
\end{aligned}$$

where $\|S_{Y_J}^{-1}\| \leq 2\delta^{-1}$ by (S.12), and $\|S_{Y_J X_K}\|_F \lesssim s_n^{1/2} t_n^{1/2}$ is by the boundedness of X_k and Y_i . For the third term,

$$\begin{aligned}
|\text{tr}\{\Sigma_{X_K}^{-1} (\Sigma_{Y_J X_K} - S_{Y_J X_K}) S_{Y_J}^{-1} S_{Y_J X_K}\}| &\leq \|\Sigma_{X_K}^{-1}\| \cdot \|S_{Y_J}^{-1}\| \cdot \|\Sigma_{Y_J X_K} - S_{Y_J X_K}\|_F \cdot \|S_{Y_J X_K}\|_F \\
&\lesssim \delta^{-1} \cdot 2\delta^{-1} \cdot s_n^{1/2} t_n^{1/2} K_{nj} \cdot s_n^{1/2} t_n^{1/2} \\
&\lesssim \delta^{-2} \{\max(s_n, t_n)\}^2 K_{nj}.
\end{aligned}$$

For the fourth term,

$$\begin{aligned}
|\text{tr}\{(\Sigma_{X_K}^{-1} - S_{X_K}^{-1}) S_{X_K Y_J} S_{Y_J}^{-1} S_{Y_J X_K}\}| &= |\text{tr}\{\Sigma_{X_K}^{-1} (S_{X_K} - \Sigma_{X_K}) S_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1} S_{Y_J X_K}\}| \\
&\leq \|\Sigma_{X_K}^{-1}\| \cdot \|S_{X_K} - \Sigma_{X_K}\|_F \cdot \|S_{X_K}^{-1} S_{X_K Y_J} S_{Y_J}^{-1} S_{Y_J X_K}\|_F \\
&= \|\Sigma_{X_K}^{-1}\| \cdot \|S_{X_K} - \Sigma_{X_K}\|_F \cdot \tau(P_j) \\
&\lesssim \delta^{-1} \cdot \{\max(s_n, t_n)\}^2 \cdot K_{nj},
\end{aligned}$$

245 where the last inequality is by the boundedness of the Pillai trace: $\tau^2(P_j) \leq \text{rank}\{C_{XY}(P_j)\} \leq \min(s_n, t_n)$. Finally, we have

$$|\tau_0^2 - \tau^2(P_j)| \lesssim (\delta^{-1} + \delta^{-2}) \{\max(s_n, t_n)\}^2 K_{nj} \lesssim \delta^{-2} \{\max(s_n, t_n)\}^2 K_{nj}.$$

S.5.7. Proof of Lemma S.5

The proof of $|\hat{\sigma}_{nj}^2 - \sigma_{nj}^2| \leq \delta^{-2} \{\max(s_n, t_n)\}^2 K_{nj}$ follows from the proof of Lemma A.6. in LL2018. The only difference lies in the factor $\{\max(s_n, t_n)\}^2 = \max(s_n t_n, s_n^2, t_n^2)$, which comes from the trace terms in $|(P_j - P_0)\hat{D}_{nj}|$. 250

To show $|\hat{\sigma}_{nj}^2 - \text{var}_{P_0}\{D_{nj}(P_0)(O)\}| \leq \delta^{-4} \{\max(s_n, t_n)\}^2 K_{nj} + \left(\widehat{\text{Rem}}_{nj}\right)^2$, notice that $P_0 \hat{D}_{nj} = \tau_0^2 - \tau^2(P_j) + \widehat{\text{Rem}}_{nj}$ and therefore

$$\left(P_0 \hat{D}_{nj}\right)^2 \leq \max_d \{\tau_0^2 - \tau^2(P_j)\}^2 + \left(\widehat{\text{Rem}}_{nj}\right)^2.$$

The rest of the proof follows from (S.14) and the proof of Lemma A.6. in LL2018. Finally for $|\widehat{\text{Rem}}_{nj}|$, we apply Lemma S.1 and note that $\tilde{\delta} \geq \delta/2$ (S.11) and $\lambda \leq 2$ (S.12),

$$|\text{Rem}^d(P)| < \delta^{-1}(s_n + t_n)CK_{nj} + (\delta^{-1} + 24\delta^{-2})s_n t_n C^2 K_{nj}^2.$$

S.5.8. Proof of Lemma S.6

From Lemma S.5,

$$|\hat{\sigma}_{nj}^2 - \text{var}_{P_0}\{D_{nj}(P_0)(O)\}| \lesssim \delta^{-4} \max(s_n^2, t_n^2) K_{nj} + \delta^{-2} \max(s_n^2, t_n^2) K_{nj}^2. \quad (\text{S.20})$$

It is easy to verify that, for a universal constant $C > 0$, the right hand side of the above inequality is bounded above by $\gamma/2$ for all j such that

$$j > C\gamma^{-1}\delta^{-2} \max\{\max(s_n^2, t_n^2), \gamma^{-2}\delta^{-6} \max(s_n^4, t_n^4)\} \log \max(n, p, q). \quad (\text{S.21})$$

Since s_n and t_n are both greater than 1, we can let $C(\delta, \gamma) = C\gamma^{-1}\delta^{-2} \max\{1, \gamma^{-2}\delta^{-6}\}$. Then, if $j > C(\delta, \gamma) \max(s_n^4, t_n^4) \log \max(n, p, q)$, it must satisfy (S.21) and hence

$$|\hat{\sigma}_{nj}^2 - \text{var}_{P_0}\{D_{nj}(P_0)(O)\}| \leq \gamma/2.$$

The result then follows from (14) and application of the triangle inequality.

S.5.9. Proof of Lemma S.7

Because we have fixed δ, γ and s as n grows, (S.6) implies that, for n large enough, ℓ_n is at least $J(n, \delta)$ and is at least $C(\delta, \gamma, s) \log \max(n, p, q)$, where these quantities are defined in Lemmas S.3 and S.6. We next verify properties (C1), (C2) and (C3) in Lemma S.7. 260

Property (C1): Because ℓ_n is at least $C(\delta, \gamma, s) \log \max(n, p, q)$, we have $\hat{\sigma}_{nj}^2 \geq \gamma/2$ from Lemma S.6. Because ℓ_n is at least $J(n, \delta)$, we have $|\hat{D}_{nj}(o)| \leq 8\delta^{-1}s$ from Lemma S.4. For $j \geq \ell_n$ and under the event $\mathcal{A}_n$, we have $|\hat{D}_{nj}(o)|/\hat{\sigma}_{nj} \lesssim s\delta^{-1}\gamma^{-1/2}$. By Lemma S.2, $\mathcal{A}_n$ occurs with probability at least $1 - 1/n$. Properties (C2) and (C3) can be verified in the same way as in Theorem A.10 of LL2018 and noting that $|\widehat{\text{Rem}}_{nj}| \lesssim K_{nj}^2$ and $|\hat{\sigma}_{nj}^2 - \sigma_{nj}^2| \lesssim K_{nj}$. 265

REFERENCES

- DiCiccio, C. J. & Romano, J. P. (2017). Robust permutation tests for correlation and regression coefficients. *Journal of the American Statistical Association* **112**, 1211–1220. 270
- Donoho, D. & Jin, J. (2004). Higher criticism for detecting sparse heterogeneous mixtures. *The Annals of Statistics* **32**, 962–994.
- Donoho, D. & Jin, J. (2015). Higher criticism for large-scale inference, especially for rare and weak effects. *Statistical Science* **30**, 1–25. 275
- Khim, J. T., Jog, V. & Loh, P.-L. (2016). Computing and maximizing influence in linear threshold and triggering models. In *Advances in Neural Information Processing Systems*.

- KRAUSE, A. & GOLOVIN, D. (2014). Submodular function maximization. In *Tractability: Practical Approaches to Hard Problems*. Cambridge University Press.
- 280 LUEDTKE, A. R. & VAN DER LAAN, M. (2018). Parametric-rate inference for one-sided differentiable parameters. *Journal of the American Statistical Association* **113**, 780–788.
- NEMHAUSER, G. L. & WOLSEY, L. A. (1978). Best algorithms for approximating the maximum of a submodular set function. *Mathematics of Operations Research* **3**, 177–188.
- 285 VAN DER VAART, A. W. & WELLNER, J. A. (1996). *Weak Convergence and Empirical Processes with Applications to Statistics*. New York, NY, USA: Springer.
- WANG, J. (2015). Joint estimation of sparse multivariate regression and conditional graphical models. *Statistica Sinica* **25**, 831–851.