

Simultaneous envelopes for multivariate linear regression

R. Dennis Cook* and Xin Zhang†

December 2, 2013

Abstract

We introduce envelopes for simultaneously reducing the predictors and the responses in multivariate linear regression, so the regression then depends only on estimated linear combinations of $\mathbf{X}$ and $\mathbf{Y}$. We use a likelihood-based objective function for estimating envelopes and then propose algorithms for estimation of a simultaneous envelope as well as for basic Grassmann manifold optimization. The asymptotic properties of the resulting estimator are studied under normality and extended to general distributions. We also investigate likelihood ratio tests and information criteria for determining the simultaneous envelope dimensions. Simulation studies and real data examples show substantial gain over the classical methods like, partial least squares, canonical correlation analysis and reduced-rank regression. This article has supplementary material available online.

Key Words: Canonical correlations; envelope model; Grassmann manifold; partial least squares; principal component analysis; reduced-rank regression; sufficient dimension reduction.

1 Introduction

Multivariate linear regression with predictors $\mathbf{X} \in \mathbb{R}^p$ and responses $\mathbf{Y} \in \mathbb{R}^r$ is a cornerstone of multivariate statistics. When p and r are not small, it is widely recognized that reducing the dimensionalities of $\mathbf{X}$ and $\mathbf{Y}$ may often result in improved performance. Perhaps the most popular methods for reducing the number of predictors and responses are principal component analysis (PCA), partial least squares (PLS) regression, canonical correlation analysis (CCA) and reduced-rank regression (RRR).

*R. Dennis Cook is Professor, School of Statistics, University of Minnesota, Minneapolis, MN 55455 (E-mail: dennis@stat.umn.edu).

†Xin Zhang is Ph.D student, School of Statistics, University of Minnesota, Minneapolis, MN, 55455 (Email: zhan0648@umn.edu).

Principal component analysis (Jolliffe, 1986, 2005) is an un-supervised dimension reduction method designed to select orthogonal linear combinations of $\mathbf{X}$ and $\mathbf{Y}$ with maximal variation. However, it does not use any information about the relationship between $\mathbf{X}$ and $\mathbf{Y}$, and thus separate PCA reductions of $\mathbf{X}$ and $\mathbf{Y}$ could be ineffective for regression. Partial least squares was proposed as iterative algorithms, NIPALS (Wold, 1966) and SIMPLS (de Jong, 1993), for predictor reduction. These methods, which are used extensively in chemometrics, reduce the predictors by iteratively estimating linear combinations of them that have maximal covariance with the response vector. Canonical correlation analysis (Hotelling, 1936; Anderson, 1984) is used to investigate the overall correlations between the two sets of variables $\mathbf{X}$ and $\mathbf{Y}$. It can simultaneously reduce $\mathbf{X}$ and $\mathbf{Y}$ by finding pairs of linear combinations such that the correlations of these pairs are in descending order and components are uncorrelated across different pairs. A probabilistic interpretation of CCA was given by Bach and Jordan (2005) as a latent variable model for two normal random vectors. Reduced-rank regression (Izenman 1975; Reinsel and Velu 1998) restricts on the rank of the regression coefficient matrix and therefore improves prediction by reducing the number of parameters in the model.

Sufficient dimension reduction methods for multivariate responses are also available for reducing dimensionality. For example, sliced inverse regression (Li, 1991) was extended to multivariate response data by Li, et al. (2003). Li, Wen and Zhu (2008) proposed “projective resampling” to deal with a multivariate regression by using univariate reduction methods. However, such methods are beyond the scope of this work since they are designed to estimate only a subspace as a preliminary step in an analysis. In contrast, we work in the context of the multivariate linear model with a view toward prediction and coefficient estimation.

Informally, multivariate linear regression can involve both material and immaterial variation in the responses and in the predictors. Material variation provides information that is directly relevant to the regression, while the immaterial variation is essentially irrelevant to the regression and serves to increase estimative variation. Envelopes, which were introduced by Cook, Li and Chiaromonte (2010) for response reduction, use a subspace to envelop the material information and thereby exclude the immaterial variation.

58 Essentially a form of targeted dimension reduction, this process can lead to substantial
59 efficiency gains when the immaterial variation is large relative to the material variation.
60 Cook, Helland and Su (2013) adapted envelopes to the predictors, and showed that the
61 SIMPLS algorithm for partial least squares regression converges to an envelope in the
62 predictor space. Following Cook, et al. (2010), they demonstrated that using a likelihood-
63 based objective function to separate the material and immaterial variation and to provide
64 an estimator of the coefficient matrix produces clear and often substantial estimative and
65 predictive advantages over SIMPLS.

66 However, little is known about using envelopes for joint reduction of the responses
67 and predictors. The previous developments kindle a hope that we can combine their
68 advantages to produce efficiency gains than are greater than those possible by reducing
69 either the responses or the predictors alone. In this article, we develop likelihood-based
70 envelope methods for simultaneously separating the material and immaterial variation in
71 the responses and in the predictors. We show a potential for synergy in a synchronized
72 reduction, producing an overall reduction in estimative variation surpassing that indicated
73 by the marginal reductions. Finding a likelihood-based envelope can be computationally
74 challenging, and so we propose a novel and fast optimization algorithm.

75 The rest of this paper is organized as follows. In Section 2, we briefly review the alge-
76 braic definition of an envelope as a prelude to our development of simultaneous reduction
77 that starts in Section 3. We depart a bit from convention and integrate our review of indi-
78 vidual response and predictor envelopes into our development of simultaneous envelopes.
79 We also link the simultaneous envelope method with partial least squares and canonical
80 correlation analysis. In Section 4 we introduce a likelihood-based objective function that
81 includes the objective functions used by Cook, et al. (2010) and Cook, et al. (2013) as
82 special cases. Novel algorithms for estimating a simultaneous envelope are also given in
83 this section. In Section 5, asymptotic properties of the simultaneous envelope estimators
84 are studied under normality and under general distributional assumptions. Encouraging
85 simulation results on prediction and on determine the dimensions of envelopes are given
86 in Section 7. In Section 8, we demonstrate the superiority of the simultaneous envelope
87 estimator compared to classical methods by simulations and by predicting the contents of

88 biscuit dough samples. Proofs and other technical details are included in the Supplement
89 to this article.

90 The following notations and definitions will be used in our exposition. Let $\mathbf{A} \sim \mathbf{B}$
91 denote that $\mathbf{A}$ has the same distribution as $\mathbf{B}$, let $\mathbf{A} \perp \mathbf{B}$ denote that $\mathbf{A}$ is independent
92 of $\mathbf{B}$ and let $\mathbf{A} \perp \mathbf{B} | \mathbf{C}$ indicate that $\mathbf{A}$ is conditionally independent of $\mathbf{B}$ given $\mathbf{C}$. Let
93 $\mathbb{R}^{m \times n}$ be the set of all real $m \times n$ matrices and let $\mathbb{S}^{k \times k}$ be the set of all real symmetric
94 $k \times k$ matrices. If $\mathbf{M} \in \mathbb{R}^{m \times n}$ then $\text{span}(\mathbf{M}) \subseteq \mathbb{R}^m$ is the subspace spanned by columns
95 of $\mathbf{M}$. We use $\hat{\boldsymbol{\theta}}_u$ to denote sample estimator for $\boldsymbol{\theta}$ with known parameters u . We use
96 $\mathbf{P}_{\mathbf{A}(\mathbf{V})} = \mathbf{A}(\mathbf{A}^T \mathbf{V} \mathbf{A})^{-1} \mathbf{A}^T \mathbf{V}$ to denote the projection onto $\text{span}(\mathbf{A})$ with the $\mathbf{V}$ inner
97 product and use $\mathbf{P}_{\mathbf{A}}$ to denote projection onto $\text{span}(\mathbf{A})$ with the identity inner product.
98 Let $\mathbf{Q}_{\mathbf{A}(\mathbf{V})} = \mathbf{I} - \mathbf{P}_{\mathbf{A}(\mathbf{V})}$. For an $m \times n$ matrix $\mathbf{A}$ and a $p \times q$ matrix $\mathbf{B}$, their direct sum is
99 defined as the $(m+p) \times (n+q)$ block diagonal matrix $\mathbf{A} \oplus \mathbf{B} = \text{diag}(\mathbf{A}, \mathbf{B})$. We will also
100 use the $\oplus$ operator for two subspaces. If $\mathcal{S} \subseteq \mathbb{R}^p$ and $\mathcal{R} \subseteq \mathbb{R}^q$ then $\mathcal{S} \oplus \mathcal{R} = \text{span}(\mathbf{S} \oplus \mathbf{R})$
101 where $\mathbf{S}$ and $\mathbf{R}$ are basis matrices for $\mathcal{S}$ and $\mathcal{R}$. The sum of two subspaces $\mathcal{S}_1$ and $\mathcal{S}_2$
102 of $\mathbb{R}^m$ is defined as $\mathcal{S}_1 + \mathcal{S}_2 = \{\mathbf{v}_1 + \mathbf{v}_2 | \mathbf{v}_1 \in \mathcal{S}_1, \mathbf{v}_2 \in \mathcal{S}_2\}$. We will use operators
103 $\text{vec} : \mathbb{R}^{a \times b} \rightarrow \mathbb{R}^{ab}$, which vectorizes an arbitrary matrix by stacking its columns, and
104 $\text{vech} : \mathbb{R}^{a \times a} \rightarrow \mathbb{R}^{a(a+1)/2}$, which vectorizes a symmetric matrix by extracting its columns
105 of elements below or on the diagonal.

106 The estimation of envelopes will eventually be done by optimizing over Grassmann
107 manifolds. A Grassmann manifold is the collection of all linear subspaces of a vector
108 space of a given dimension. We use $\mathcal{G}_{(k,n)}$ to denote all the k -dimensional subspaces in
109 $\mathbb{R}^n$, which is a Grassmann manifold with dimension $k(n-k)$. For more background
110 in Grassmann manifolds and Grassmann optimizations, see Edelman, Tomas and Smith
111 (1998). There are two currently available packages for Grassmann manifold optimiza-
112 tion: R package *GrassmannOptim* by Adraghi, Cook and Wu (2012) and the MATLAB
113 package *sg_min* by Ross A. Lippert (<http://web.mit.edu/~ripper/www/software/>).

2 Definition of an envelope

In this section we briefly review definitions and a property of reducing subspaces and envelopes.

Definition 1. A subspace $\mathcal{R} \subseteq \mathbb{R}^p$ is said to be a reducing subspace of $\mathbf{M} \in \mathbb{R}^{p \times p}$ if and only if $\mathcal{R}$ decomposes $\mathbf{M}$ as $\mathbf{M} = \mathbf{P}_{\mathcal{R}}\mathbf{M}\mathbf{P}_{\mathcal{R}} + \mathbf{Q}_{\mathcal{R}}\mathbf{M}\mathbf{Q}_{\mathcal{R}}$. If $\mathcal{R}$ is a reducing subspace of $\mathbf{M}$, we say that $\mathcal{R}$ reduces $\mathbf{M}$.

This definition of a reducing subspace is equivalent to the usual definition found in functional analysis (Conway, 1990) and in the literature on invariant subspaces, but the underlying notion of reduction is incompatible with how it is usually understood in statistics. Nevertheless, it is common terminology in those areas and is the basis for the definition of an envelope (Cook, et al., 2010) which is central to our developments.

Definition 2. Let $\mathbf{M} \in \mathbb{S}^{p \times p}$ and let $\mathcal{B} \subseteq \text{span}(\mathbf{M})$. Then the $\mathbf{M}$ -envelope of $\mathcal{B}$, denoted by $\mathcal{E}_{\mathbf{M}}(\mathcal{B})$, is the intersection of all reducing subspaces of $\mathbf{M}$ that contain $\mathcal{B}$.

The intersection of two reducing subspaces of $\mathbf{M}$ is still a reducing subspace of $\mathbf{M}$. This means that $\mathcal{E}_{\mathbf{M}}(\mathcal{B})$, which is unique by its definition, is the smallest reducing subspace containing $\mathcal{B}$. Also, the $\mathbf{M}$ -envelope of $\mathcal{B}$ always exist because of the requirement $\mathcal{B} \subseteq \text{span}(\mathbf{M})$.

The following proposition from Cook, et al. (2010) gives a characterization of envelopes.

Proposition 1. If $\mathbf{M} \in \mathbb{S}^{p \times p}$ has $q \leq p$ eigenspaces, then the $\mathbf{M}$ -envelope of $\mathcal{B} \subseteq \text{span}(\mathbf{M})$ can be constructed as $\mathcal{E}_{\mathbf{M}}(\mathcal{B}) = \sum_{i=1}^q \mathbf{P}_i \mathcal{B}$, where $\mathbf{P}_i$ is the projection onto the i -th eigenspace of $\mathbf{M}$.

From this proposition, we see that the $\mathbf{M}$ -envelope of $\mathcal{B}$ is the sum of the eigenspaces of $\mathbf{M}$ that are not orthogonal to $\mathcal{B}$; that is, the eigenspaces of $\mathbf{M}$ onto which $\mathcal{B}$ projects non-trivially. This implies that the envelope is the span of some subset of the eigenvectors of $\mathbf{M}$. In the regression context, $\mathcal{B}$ is typically the span of a regression coefficient matrix or a matrix of covariances, and $\mathbf{M}$ is chosen as a covariance matrix which is usually positive definite.

3 Simultaneous envelopes

3.1 Definition and structure

The standard multivariate linear model can be written as

$$\mathbf{Y} = \boldsymbol{\mu}_Y + \boldsymbol{\beta}^T(\mathbf{X} - \boldsymbol{\mu}_X) + \boldsymbol{\epsilon}, \quad (3.1)$$

where $\boldsymbol{\mu}_Y$ is the mean for $\mathbf{Y}$, $\boldsymbol{\mu}_X$ is the mean for $\mathbf{X}$, $\boldsymbol{\epsilon}$ is the error vector that has mean 0, variance $\boldsymbol{\Sigma}_{Y|X} > 0$ and is independent of $\mathbf{X}$, and $\boldsymbol{\beta} \in \mathbb{R}^{p \times r}$ is the regression coefficient matrix in which we are primarily interested. We use $\boldsymbol{\Sigma}_X > 0$ and $\boldsymbol{\Sigma}_Y > 0$ to denote the population covariance matrices of $\mathbf{X}$ and $\mathbf{Y}$, and use $\boldsymbol{\Sigma}_{XY}$ to denote their cross-covariance matrix. The covariance matrix of the population residual vector $\boldsymbol{\epsilon}$ is then $\boldsymbol{\Sigma}_{Y|X} = \boldsymbol{\Sigma}_Y - \boldsymbol{\Sigma}_{XY}^T \boldsymbol{\Sigma}_X^{-1} \boldsymbol{\Sigma}_{XY}$. Similarly, let $\boldsymbol{\Sigma}_{X|Y} = \boldsymbol{\Sigma}_X - \boldsymbol{\Sigma}_{XY} \boldsymbol{\Sigma}_Y^{-1} \boldsymbol{\Sigma}_{XY}^T$. The sample counterparts of these population covariance matrices are denoted by $\mathbf{S}_X$, $\mathbf{S}_Y$, $\mathbf{S}_{XY}$, $\mathbf{S}_{Y|X}$ and $\mathbf{S}_{X|Y}$. We also define $\boldsymbol{\Sigma}_{A|B}$ and $\mathbf{S}_{A|B}$ in the same way for two arbitrary random vectors $\mathbf{A}$ and $\mathbf{B}$.

Envelope methods have the potential to increase efficiency in estimation of $\boldsymbol{\beta}$ after reducing $\mathbf{Y}$ (Cook et al., 2010) and to improve prediction of $\mathbf{Y}$ after reducing $\mathbf{X}$ (Cook et al., 2013). Our goal is to combine their advantages by simultaneously reducing $\mathbf{X}$ and $\mathbf{Y}$ to decrease both predictive and estimative variation. We next give a coordinate representation of simultaneous envelopes.

Let $d \leq \min(r, p)$ denote the rank of $\boldsymbol{\beta}$ and consider the singular value decomposition of $\boldsymbol{\beta} = \mathbf{U}\mathbf{D}\mathbf{V}^T$, where $\mathbf{U} \in \mathbb{R}^{p \times d}$ and $\mathbf{V} \in \mathbb{R}^{r \times d}$ are orthogonal matrices, $\mathbf{U}^T \mathbf{U} = \mathbf{I}_d = \mathbf{V}^T \mathbf{V}$, and $\mathbf{D} = \text{diag}(\lambda_1, \dots, \lambda_d)$ is a diagonal matrix with elements $\lambda_1 \geq \dots \geq \lambda_d > 0$ being the d singular values of $\boldsymbol{\beta}$. We define the column space (the left eigenspace) and the row space (the right eigenspace) of $\boldsymbol{\beta}$ as $\text{span}(\boldsymbol{\beta}) = \text{span}(\mathbf{U}) \equiv \mathcal{L}$ and $\text{span}(\boldsymbol{\beta}^T) = \text{span}(\mathbf{V}) \equiv \mathcal{R}$. Then two envelopes can be constructed for simultaneously reducing the predictor space and the response space:

1. $\mathbf{X}$ -envelope (left envelope): $\mathcal{E}_{\mathbf{S}_X}(\mathcal{L})$ with $\dim(\mathcal{E}_{\mathbf{S}_X}(\mathcal{L})) = d_X$, $d \leq d_X \leq p$.
2. $\mathbf{Y}$ -envelope (right envelope): $\mathcal{E}_{\mathbf{S}_{Y|X}}(\mathcal{R})$ with $\dim(\mathcal{E}_{\mathbf{S}_{Y|X}}(\mathcal{R})) = d_Y$, $d \leq d_Y \leq r$.

168 We know from Cook, et al. (2010; Proposition 3.1) that $\mathcal{E}_{\Sigma_Y}(\mathcal{R}) = \mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R})$. Con-
 169 sequently, an alternative definition of the Y -envelope is $\mathcal{E}_{\Sigma_Y}(\mathcal{R})$, which then has the
 170 same form as X -envelope, by replacing X with Y and $\mathcal{L}$ with $\mathcal{R}$. We use $\mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R})$ as
 171 the definition of the Y -envelope to facilitate later parameterizations. If we imagine that
 172 the elements of β are generated with respect to Lebesgue measure then it follows from
 173 Proposition 1 that no reduction is possible: $\mathcal{E}_{\Sigma_X}(\mathcal{L}) = \mathbb{R}^p$ and $\mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R}) = \mathbb{R}^r$ with prob-
 174 ability one. Later in this section we show that proper envelopes imply certain relations
 175 between X and Y that may reasonably hold in practice.

176 From the definitions of the X - and Y -envelopes, if $d = r < p$ then we can reduce
 177 X only and $\mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R}) = \mathbb{R}^r$. Similarly, if $d = p < r$ then $\mathcal{E}_{\Sigma_X}(\mathcal{L}) = \mathbb{R}^p$ and reduction
 178 is possible only in the response space. Hence, we will assume $d < \min(r, p)$ from
 179 now on and discuss the general situation where simultaneous reduction is possible. Let
 180 $\mathbf{L} \in \mathbb{R}^{p \times d_X}$ be an orthogonal basis for $\mathcal{E}_{\Sigma_X}(\mathcal{L})$ and let $\mathbf{R} \in \mathbb{R}^{r \times d_Y}$ be an orthogonal basis
 181 for $\mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R})$. Also, $(\mathbf{L}, \mathbf{L}_0)$ is an orthogonal basis for $\mathbb{R}^p$ and $(\mathbf{R}, \mathbf{R}_0)$ is an orthogonal
 182 basis for $\mathbb{R}^r$. Then from Definition 1 we can write the covariance matrices as

$$\Sigma_X = \mathbf{L}\mathbf{\Omega}\mathbf{L}^T + \mathbf{L}_0\mathbf{\Omega}_0\mathbf{L}_0^T, \quad (3.2)$$

$$\Sigma_{Y|X} = \mathbf{R}\mathbf{\Phi}\mathbf{R}^T + \mathbf{R}_0\mathbf{\Phi}_0\mathbf{R}_0^T. \quad (3.3)$$

183 The covariance matrix decomposition in (3.2) indicates that the eigenvectors of Σ_X fall in
 184 either $\mathcal{E}_{\Sigma_X}(\mathcal{L})$ or $\mathcal{E}_{\Sigma_X}^\perp(\mathcal{L})$ with corresponding eigenvalues being the eigenvalues of $\mathbf{\Omega}$ and
 185 $\mathbf{\Omega}_0$. No relationship is assumed between the eigenvalues of $\mathbf{\Omega}$ and $\mathbf{\Omega}_0$: The eigenvalues of
 186 $\mathbf{\Omega}$ could be any subset of the eigenvalues of Σ_X . Similar results hold for the eigenvalues
 187 and eigenvectors of $\Sigma_{Y|X}$, as seen in (3.3).

188 All of the information that is available about β is carried by the reduced variables
 189 $\mathbf{R}^T\mathbf{Y}$ and $\mathbf{L}^T\mathbf{X}$, which can be seen as follows. Recall the singular value decomposition
 190 of β : $\beta = \mathbf{U}\mathbf{D}\mathbf{V}^T$, where $\text{span}(\mathbf{U}) = \mathcal{L} \subseteq \text{span}(\mathbf{L})$ and $\text{span}(\mathbf{V}) = \mathcal{R} \subseteq \text{span}(\mathbf{R})$
 191 by the definition of the X - and Y -envelopes. Hence, $\beta = \mathbf{U}\mathbf{D}\mathbf{V}^T = (\mathbf{L}\mathbf{A})\mathbf{D}(\mathbf{B}^T\mathbf{R}^T)$
 192 for some semi-orthogonal matrices $\mathbf{A} \in \mathbb{R}^{d_X \times d}$ and $\mathbf{B} \in \mathbb{R}^{d_Y \times d}$. Then $\mathbf{R}_0^T\beta^T = 0$,
 193 $\beta^T\mathbf{L}_0 = 0$ and model (3.1) can be reduced to

$$\begin{cases} \mathbf{R}^T\mathbf{Y} = \mathbf{R}^T\mu_Y + \eta^T\{\mathbf{L}^T(\mathbf{X} - \mu_X)\} + \mathbf{R}^T\epsilon, \\ \mathbf{R}_0^T\mathbf{Y} = \mathbf{R}_0^T\mu_Y + \mathbf{R}_0^T\epsilon, \end{cases} \quad (3.4)$$

194 where $\boldsymbol{\eta} = \mathbf{BDA}^T \in \mathbb{R}^{d_X \times d_Y}$ has rank d . The *simultaneous envelope model* then be-
 195 comes

$$\begin{aligned} \mathbf{Y} &= \mathbf{R}\mathbf{R}^T\mathbf{Y} + \mathbf{R}_0\mathbf{R}_0^T\mathbf{Y} \\ &= \boldsymbol{\mu}_Y + \mathbf{R}\boldsymbol{\eta}^T\mathbf{L}^T(\mathbf{X} - \boldsymbol{\mu}_X) + \boldsymbol{\epsilon}. \end{aligned} \quad (3.5)$$

196 with $\boldsymbol{\Sigma}_X$ and $\boldsymbol{\Sigma}_{Y|X}$ given by (3.2) and (3.3).

197 Comparing to (3.1), we see that the regression coefficient matrix is now $\boldsymbol{\beta} = \mathbf{L}\boldsymbol{\eta}\mathbf{R}^T$,
 198 where $\boldsymbol{\eta}$ contains the coordinates of $\boldsymbol{\beta}$ relative to $\mathbf{L}$ and $\mathbf{R}$. This implies that the columns
 199 and rows of $\boldsymbol{\beta}$ vary only within the $\mathbf{Y}$ -envelope and the $\mathbf{X}$ -envelope. By letting $\mathbf{L} = \mathbf{I}_p$ or
 200 $\mathbf{R} = \mathbf{I}_r$, there will be reductions only in the column space or the row space of $\boldsymbol{\beta}$. These
 201 are the two special situations studied by Cook et al. (2010) and Cook et al. (2013).

202 If $\mathbf{L} = \mathbf{I}_p$, it follows from (3.3) and (3.4) that $\text{cov}(\mathbf{R}^T\mathbf{Y}, \mathbf{R}_0^T\mathbf{Y}|\mathbf{X}) = 0$ and $\mathbf{R}_0^T\mathbf{Y}|\mathbf{X} \sim$
 203 $\mathbf{R}_0^T\mathbf{Y}$, which motivated the construction of response envelopes (Cook, et al., 2010). If $\boldsymbol{\epsilon}$
 204 is normally distributed, then this pair of conditions is equivalent to $\mathbf{R}_0^T\mathbf{Y} \perp\!\!\!\perp \mathbf{R}^T\mathbf{Y}|\mathbf{X}$ and
 205 $\mathbf{R}_0^T\mathbf{Y}|\mathbf{X} \sim \mathbf{R}_0^T\mathbf{Y}$. If $(\mathbf{X}, \mathbf{Y})$ is jointly normal then this pair of conditions is equivalent to
 206 $\mathbf{R}_0^T\mathbf{Y} \perp\!\!\!\perp (\mathbf{R}^T\mathbf{Y}, \mathbf{X})$. We refer to $\mathbf{R}^T\mathbf{Y}$ and $\mathbf{R}_0^T\mathbf{Y}$ as the material and immaterial parts of
 207 the responses. This is because $\mathbf{R}_0^T\mathbf{Y}$ is neither affected by the predictors nor correlated
 208 with the complementary part of the responses $\mathbf{R}^T\mathbf{Y}$ and in this sense has no contribution
 209 to linear regression.

210 If $\mathbf{R} = \mathbf{I}_r$ then, from (3.2) and (3.5), $\text{cov}(\mathbf{Y}, \mathbf{L}_0^T\mathbf{X}|\mathbf{L}^T\mathbf{X}) = 0$ and $\text{cov}(\mathbf{L}^T\mathbf{X}, \mathbf{L}_0^T\mathbf{X}) =$
 211 0 , which are the two conditions used by Cook, et al. (2013, Proposition 2.1) for predictor
 212 reduction. If $(\mathbf{X}, \mathbf{Y})$ has a joint normal distribution, then this pair of conditions is equiv-
 213 alent to $\mathbf{L}_0^T\mathbf{X} \perp\!\!\!\perp (\mathbf{L}^T\mathbf{X}, \mathbf{Y})$. We refer to $\mathbf{L}^T\mathbf{Y}$ and $\mathbf{L}_0^T\mathbf{Y}$ as the material and immaterial
 214 parts of the predictors. Similar to the response, this is because $\mathbf{L}_0^T\mathbf{Y}$ is neither affected
 215 by the response nor correlated with the rest of the predictors.

216 The above conditions for $\mathbf{L}$ and $\mathbf{R}$ are not stated symmetrically because of the natural
 217 of regression. Cook et al. (2010) treated $\mathbf{X}$ as fixed since $\mathbf{S}_X$ is an ancillary statistic
 218 for the $\mathbf{Y}$ -envelope, while Cook et al. (2013) treated $\mathbf{X}$ as random because $\mathbf{S}_X$ is not
 219 ancillary for $\mathbf{X}$ reduction. For simultaneous reductions, we assume that $\mathbf{X}$ and $\mathbf{Y}$ have a
 220 joint distribution throughout this article. The covariance decompositions (3.2) and (3.3)

221 play a critical role in obtaining the above relationships, and they distinguish the envelope
222 reductions from other methods for reducing the column and row dimensions of β .

223 The previous relationships follow from the marginal response and predictor envelopes.
224 The following lemma describes additional relationships between the material part in $\mathbf{X}$
225 (or $\mathbf{Y}$) and the immaterial part in $\mathbf{Y}$ (or $\mathbf{X}$) that come with the simultaneous envelopes.

226 **Lemma 1.** *Assume the simultaneous envelope model (3.5). Then $\text{cov}(\mathbf{R}^T \mathbf{Y}, \mathbf{L}_0^T \mathbf{X}) = 0$
227 and $\text{cov}(\mathbf{L}^T \mathbf{X}, \mathbf{R}_0^T \mathbf{Y}) = 0$.*

228 This lemma, which does not require normality of $\mathbf{X}$ or $\mathbf{Y}$, is implied by the previous
229 discussion if $\mathbf{L} = \mathbf{I}_p$ or $\mathbf{R} = \mathbf{I}_r$. It shows a similarity between simultaneous envelope
230 reduction and canonical correlation analysis: the selected components are uncorrelated
231 with the rest of the components. Additional discussion of the connection between en-
232 velopes and canonical correlation is given in Section 3.3.

233 3.2 A visualized example of simultaneous envelope

234 Figure 3.1 (a) and (b) illustrate the working mechanism of the simultaneous envelope
235 reduction for a multivariate regression with two response $\mathbf{Y} = (Y_1, Y_2)^T$ and two predic-
236 tors $\mathbf{X} = (X_1, X_2)^T$. For ease of illustration, we assume that $\mathbf{Y} = \beta^T \mathbf{X} + \epsilon$ with rank
237 one regression coefficient matrix $\beta = \mathbf{L}\mathbf{R}^T$ for some 2×1 matrices $\mathbf{L}$ and $\mathbf{R}$ such that
238 $(\mathbf{L}, \mathbf{L}_0) \in \mathbb{R}^{2 \times 2}$ and $(\mathbf{R}, \mathbf{R}_0) \in \mathbb{R}^{2 \times 2}$ are orthogonal matrices. Then the plots demonstrate
239 the set-up where $\mathbf{L}$ and $\mathbf{R}$ span the $\mathbf{X}$ - and $\mathbf{Y}$ -envelope.

240 In the first plot, the conditional distribution of $\mathbf{Y}|\mathbf{X}$ is represented by the ellipses,
241 whose axes are the directions of the eigenvectors of $\Sigma_{\mathbf{Y}|\mathbf{X}}$. The shift from one contour to
242 another is captured by $\beta^T \mathbf{X} = \mathbf{R}\mathbf{L}^T \mathbf{X}$, which is in the direction of $\mathbf{R}$ and has magnitude
243 proportional to $\mathbf{L}^T \mathbf{X}$. From Proposition 1, the $\mathbf{Y}$ -envelope is the sum of eigenspaces of
244 $\Sigma_{\mathbf{Y}|\mathbf{X}}$ that are not orthogonal to $\mathbf{R}$. In this plot, the eigenvector corresponds to the larger
245 eigenvalue of $\Sigma_{\mathbf{Y}|\mathbf{X}}$ is orthogonal to $\mathbf{R}$ and hence represents the immaterial information
246 of the regression. By projecting the data onto $\mathbf{R}^T \mathbf{Y}$, we will eliminate the immaterial
247 variation in the response. The response envelope reduction in this case is very efficient
248 because $\mathbf{R}$ lies in the eigenspace corresponds to the small eigenvalue of $\Sigma_{\mathbf{Y}|\mathbf{X}}$.

249 The second plot represents the marginal distribution of $\mathbf{X}$. From Proposition 1, the
 250 reduction by $\mathbf{X}$ -envelope is available when $\mathbf{L}$ is contained in a subset of the eigenspaces
 251 of $\Sigma_{\mathbf{X}}$. In this plot, $\mathbf{L}$ happens to be the first eigenvector of $\Sigma_{\mathbf{X}}$, which means it spans
 252 the $\mathbf{X}$ -envelope and $\mathbf{L}_0^T \mathbf{X}$ will be immaterial information of the regression. It should
 253 be pointed out that if we assign the Lebesgue measure to this 2-dimensional space, the
 254 probability of $\mathbf{L}$ being one of the two eigenvector of $\Sigma_{\mathbf{X}}$ will be zero. But statistically,
 255 this event could happen because it is equivalent to requiring (a) $\text{cov}(\mathbf{Y}, \mathbf{L}_0^T \mathbf{X} | \mathbf{L}^T \mathbf{X}) = 0$
 256 and (b) $\text{cov}(\mathbf{L}^T \mathbf{X}, \mathbf{L}_0^T \mathbf{X}) = 0$, see discussion in Section 3.1. As represented by this plot,
 257 the predictor envelope reduction has great advantage over OLS when $\mathbf{L}$ corresponds to
 258 the larger eigenvalue of $\Sigma_{\mathbf{X}}$. This can be seen from the first plot, where the magnitude of
 259 $\mathbf{L}^T \mathbf{X}$ is proportional to the strength of the linear relationship.

260 For a toy data example, we use the meat data analyzed by Cook et al. (2013) for enve-
 261 lope predictor reduction in multivariate linear regression. This dataset consists of spectral
 262 measurements from infrared transmittance for 103 meat samples. There are three re-
 263 sponse variables: percentages of protein, fat and water. The sum of the three percentages
 264 is not one because of the other chemical content in the sample. We take Y_1 to be protein
 265 and Y_2 to be the sum of water and fat. We take two spectral measurements at 910nm
 266 and 960nm for illustration. The estimated $\mathbf{X}$ -envelope and the $\mathbf{Y}$ -envelope are both one-
 267 dimensional. The left plot was constructed by conditioning on the high, median and low
 268 values of $\hat{\mathbf{L}}^T \mathbf{X}$, we can clearly see that the estimated envelope direction $\hat{\mathbf{R}}$ matches the
 269 major axes of the contour of each sub-sample. So in this example the $\mathbf{Y}$ -envelope reduces
 270 the dimension but not very much immaterial information. On the other hand, the right
 271 plot closely resembles the schematic representation shown in top-right. Consequently,
 272 the simultaneous envelope method offers a more precise estimate of β than the standard
 273 method. The standard errors of the OLS estimated coefficients in $\hat{\beta}_{\text{OLS}}$ are 1.2, 12.7, 12.8
 274 and 49.5 times of that of the simultaneous envelope estimator.

275 **3.3 Links to PCA, PLS, CCA and RRR**

276 As stated previously, the eigenvalues of Ω_0 could be any subset of the eigenvalues of
 277 $\Sigma_{\mathbf{X}}$. When some or all of the largest few eigenvalues of $\Sigma_{\mathbf{X}}$ come from Ω_0 , the first few

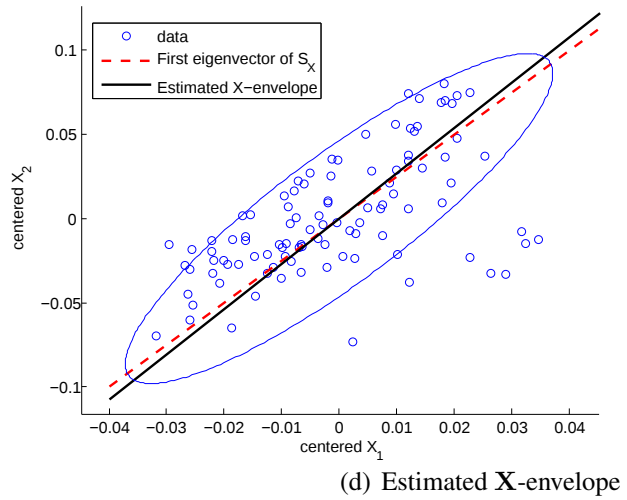
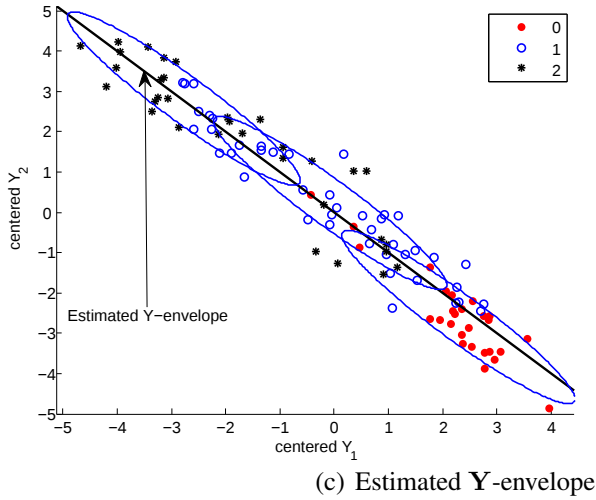
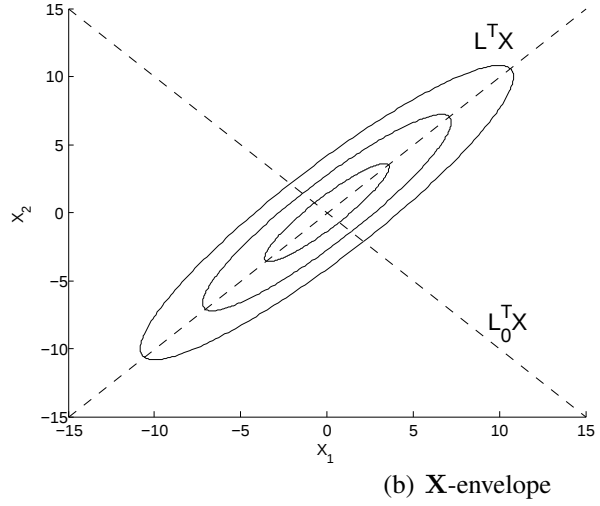
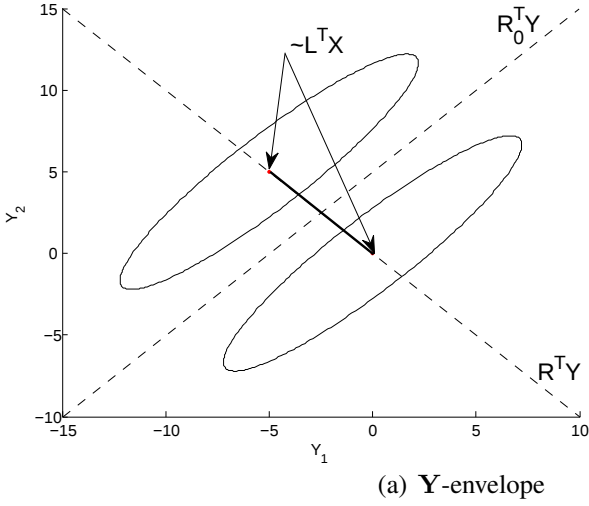


Figure 3.1: Working mechanism of simultaneous envelope reduction. (a) Schematic representation of response envelope reduction; (b) schematic representation of envelope in the predictor space; (c) the meat data with estimated Y-envelope, where the data points were marked differently according to their values in the predictor envelope $\hat{\mathbf{L}}^T \mathbf{X}_i$. ; (d) the meat data with estimated X-envelope and the first eigenvector of $\mathbf{S}_X$. To help visualization, elliptical contours that cover 90% of the data points were added in (c) and (d).

principal components of $\mathbf{X}$ will be from the immaterial part of $\mathbf{X}$, which is ineffective for the regression as we can see from the above relationships. Principal components may be effective only if the larger eigenvalues of $\Sigma_{\mathbf{X}}$ all come from Ω ; that is, from the material variation. Similar problems of principal component analysis arise for reducing $\mathbf{Y}$.

As mentioned in the Introduction, the partial least squares method reduces only the predictors and so it is comparable to simultaneous envelopes method when setting $\mathbf{R} = \mathbf{I}_r$. Cook, et al. (2013) showed that in this case the SIMPLS algorithm (de Jong, 1993) produces a $\sqrt{n}$ -consistent estimator of $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$, and that a likelihood-based estimator can do much better for prediction than the SIMPLS estimator.

Canonical correlation analysis is widely used for the purpose of simultaneously reducing the predictors and the responses. In the population, it finds canonical pairs of directions $\{\mathbf{a}_i, \mathbf{b}_i\}$, $i = 1, \dots, d$, such that the correlations between $\mathbf{a}_i^T \mathbf{X}$ and $\mathbf{b}_i^T \mathbf{Y}$ are maximized. The maximization is over the constraints $\mathbf{a}_j^T \Sigma_{\mathbf{X}} \mathbf{a}_k = 0$, $\mathbf{a}_j^T \Sigma_{\mathbf{X}\mathbf{Y}} \mathbf{b}_k = 0$ and $\mathbf{b}_j^T \Sigma_{\mathbf{Y}} \mathbf{b}_k = 0$ for all $j \neq k$, and $\mathbf{a}_j^T \Sigma_{\mathbf{X}} \mathbf{a}_j = 1$ and $\mathbf{b}_j^T \Sigma_{\mathbf{Y}} \mathbf{b}_j = 1$ for all j . The solution is then $\{\mathbf{a}_i, \mathbf{b}_i\} = \{\Sigma_{\mathbf{X}}^{-1/2} \mathbf{e}_i, \Sigma_{\mathbf{Y}}^{-1/2} \mathbf{f}_i\}$, where $\{\mathbf{e}_i, \mathbf{f}_i\}$ is the i -th left-right eigenvector pair of the correlation matrix $\boldsymbol{\rho} = \Sigma_{\mathbf{X}}^{-1/2} \Sigma_{\mathbf{X}\mathbf{Y}} \Sigma_{\mathbf{Y}}^{-1/2}$.

Lemma 2. *Under the simultaneous envelope model (3.5), canonical correlation analysis can find at most d directions in the population, where $d = \text{rank}(\boldsymbol{\beta}) \leq \min(d_X, d_Y)$. Moreover, the directions are contained in the simultaneous envelope as*

$$\text{span}(\mathbf{a}_1, \dots, \mathbf{a}_d) \subseteq \mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L}), \text{span}(\mathbf{b}_1, \dots, \mathbf{b}_d) \subseteq \mathcal{E}_{\Sigma_{\mathbf{Y}|\mathbf{X}}}(\mathcal{R}). \quad (3.6)$$

Hence, canonical correlation analysis may miss some information about the regression by ignoring some material parts of $\mathbf{X}$ and/or $\mathbf{Y}$. For example, when r is small, it can find at most r linear combinations of $\mathbf{X}$, which can be insufficient for regression. Moreover, the most correlated pairs are not, in general, the most predictable pairs for regression. Canonical correlations are often used for data visualization instead of regression, so it may be expected that they can fail in prediction. In the simulation studies of Section 7.1 and Section J in the Supplement, we found that the prediction performances based on canonical correlations varied widely for different covariance structures.

305 The maximum likelihood estimator of RRR is obtained as

$$\hat{\beta}_{\text{RRR}} = \arg \min_{\text{rank}(\beta)=d} \left\{ \sum_{i=1}^n (\mathbf{Y}_i - \beta^T \mathbf{X}_i)^T \mathbf{S}_{\mathbf{Y}|\mathbf{X}}^{-1} (\mathbf{Y}_i - \beta^T \mathbf{X}_i) \right\}, \quad (3.7)$$

306 which is the same as the regression coefficient matrix obtained by using the canonical
 307 variables (see Reinsel and Velu 1998; Section 2.4.2). Therefore, in the purpose of estima-
 308 tion and prediction, the maximum likelihood RRR estimator is equivalent to the estimator
 309 based on CCA estimator. RRR can also be applied with identity inner product in (3.7)
 310 instead of $\mathbf{S}_{\mathbf{Y}|\mathbf{X}}$. And for any minimization criteria of RRR, the estimators always have
 311 the form $\hat{\beta} = \hat{\mathbf{A}}\hat{\mathbf{B}}$ where $\hat{\mathbf{A}} \in \mathbb{R}^{p \times d}$ and $\hat{\mathbf{B}} \in \mathbb{R}^{d \times r}$ are both $\sqrt{n}$ -consistent estimators
 312 for their population counterparts $\mathbf{A}$ and $\mathbf{B}$. Similar to Lemma 2, we have the following
 313 relations by definition,

$$\text{span}(\mathbf{A}) = \text{span}(\beta) \subseteq \mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L}), \quad \text{span}(\mathbf{B}^T) = \text{span}(\beta^T) \subseteq \mathcal{E}_{\Sigma_{\mathbf{Y}}}(\mathcal{R}). \quad (3.8)$$

314 Therefore, RRR has the same drawback as CCA that it may loss some material infor-
 315 mation in regression. Simulation studies in Section 7 includes RRR estimator based on
 316 identity inner product.

317 3.4 Potential gain

318 To gain intuition about the potential advantages of simultaneous envelopes, we next con-
 319 sider the case where $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$ and $\mathcal{E}_{\Sigma_{\mathbf{Y}}}(\mathcal{R})$ are known. Estimation of these envelopes in
 320 practice will mitigate the findings in this section, but we have nevertheless found them to
 321 be useful qualitative indicators of the benefits of simultaneous reduction.

322 The envelope estimator for β with known semi-orthogonal basis matrices $\mathbf{L}$ and $\mathbf{R}$,
 323 denoted by $\hat{\beta}_{\mathbf{L},\mathbf{R}}$, can be written as

$$\hat{\beta}_{\mathbf{L},\mathbf{R}} = \mathbf{L}\hat{\eta}_{\mathbf{L},\mathbf{R}}\mathbf{R}^T = \mathbf{L}(\mathbf{L}^T\mathbf{S}_{\mathbf{X}}\mathbf{L})^{-1}\mathbf{L}^T\mathbf{S}_{\mathbf{X}\mathbf{Y}}\mathbf{R}\mathbf{R}^T = \mathbf{P}_{\mathbf{L}(\mathbf{S}_{\mathbf{X}})}\hat{\beta}_{\text{OLS}}\mathbf{P}_{\mathbf{R}}, \quad (3.9)$$

324 where $\hat{\beta}_{\text{OLS}} = \mathbf{S}_{\mathbf{X}}^{-1}\mathbf{S}_{\mathbf{X}\mathbf{Y}}$ is the ordinary least squares estimator. Clearly, the estimator
 325 $\hat{\beta}_{\mathbf{L},\mathbf{R}}$ uses only the material variation in $\mathbf{Y}$ and $\mathbf{X}$. It can be obtained by projecting $\hat{\beta}_{\text{OLS}}$
 326 to the reduced predictor space and the reduced response space, and so does not depend
 327 on the particular bases $\mathbf{L}$ and $\mathbf{R}$ selected. The estimator $\hat{\eta}_{\mathbf{L},\mathbf{R}}$ is the ordinary least squares
 328 estimator of the coefficient matrix for the regression of $\mathbf{R}^T\mathbf{Y}$ on $\mathbf{L}^T\mathbf{X}$.

329 The next proposition shows how the variance of the ordinary least squares estimator
 330 with normal predictors can be potentially reduced by using simultaneous envelopes. Let
 331 $f_p = n - p - 2$ and $f_x = n - d_X - 2$.

332 **Proposition 2.** Assume that $\mathbf{X} \sim N_p(\boldsymbol{\mu}_X, \boldsymbol{\Sigma}_X)$, $n > p + 2$ and that semi-orthogonal
 333 basis matrices $\mathbf{L}$, $\mathbf{R}$ for the left and right envelopes are known. Then $\text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\text{OLS}})) =$
 334 $f_p^{-1} \boldsymbol{\Sigma}_{Y|X} \otimes \boldsymbol{\Sigma}_X^{-1}$ and

$$\begin{aligned} \text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}})) &= f_x^{-1} (\mathbf{R}\boldsymbol{\Phi}\mathbf{R}^T) \otimes (\mathbf{L}\boldsymbol{\Omega}^{-1}\mathbf{L}^T) \\ &= f_p f_x^{-1} \text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\text{OLS}})) - f_x^{-1} \mathbf{R}_0 \boldsymbol{\Phi}_0 \mathbf{R}_0^T \otimes \mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T \\ &\quad - f_x^{-1} \mathbf{R} \boldsymbol{\Phi} \mathbf{R}^T \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T - f_x^{-1} \mathbf{R}_0 \boldsymbol{\Phi}_0 \mathbf{R}_0^T \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T, \end{aligned}$$

335 where $\boldsymbol{\Omega} = \mathbf{L}^T \boldsymbol{\Sigma}_X \mathbf{L}$, $\boldsymbol{\Omega}_0 = \mathbf{L}_0^T \boldsymbol{\Sigma}_X \mathbf{L}_0$, $\boldsymbol{\Phi}_0 = \mathbf{R}_0^T \boldsymbol{\Sigma}_{Y|X} \mathbf{R}_0$ and $\boldsymbol{\Phi} = \mathbf{R}^T \boldsymbol{\Sigma}_{Y|X} \mathbf{R}$.

336 This proposition shows that the variation in $\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}}$ can be seen in two parts: the first
 337 part is the variation in $\hat{\boldsymbol{\beta}}_{\text{OLS}}$ times a constant $f_p f_x^{-1} \leq 1$, and the second consists of
 338 terms that reduce this value depending on the variances associated with the immaterial
 339 information $\mathbf{L}_0^T \mathbf{X}$ and $\mathbf{R}_0^T \mathbf{Y}$.

340 If $\mathbf{R} = \mathbf{I}_r$, then $\boldsymbol{\Phi} = \boldsymbol{\Sigma}_{Y|X}$ and we get the multivariate version of Proposition 2.3 by
 341 Cook, et al. (2013) for univariate response regression:

$$\text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{L}})) = f_p f_x^{-1} \text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\text{OLS}})) - f_x^{-1} \boldsymbol{\Sigma}_{Y|X} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T. \quad (3.10)$$

342 When p is close to n and the $\mathbf{X}$ -envelope dimension d_X is small, the constant $f_p f_x^{-1}$ could
 343 be small and the gain by $\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}}$ over $\hat{\boldsymbol{\beta}}_{\text{OLS}}$ could be substantial. If there is substantial co-
 344 linearity in the predictors, so $\boldsymbol{\Sigma}_X$ has some small eigenvalues, and if the corresponding
 345 eigenvectors of $\boldsymbol{\Sigma}_X$ fall in $\mathcal{E}_{\boldsymbol{\Sigma}_X}^\perp(\mathcal{L})$, then the variance of $\hat{\boldsymbol{\beta}}_{\mathbf{L}}$ could be reduced consider-
 346 ably since $\boldsymbol{\Omega}_0^{-1}$ will be large. It is widely known that colinearity in $\mathbf{X}$ can increase the
 347 variance of $\hat{\boldsymbol{\beta}}_{\text{OLS}}$. However, when the eigenvectors of $\boldsymbol{\Sigma}_X$ corresponding to these small
 348 eigenvalues lie in $\mathcal{E}_{\boldsymbol{\Sigma}_X}^\perp(\mathcal{L})$, the variance of $\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}}$ is not affected by colinearity.

349 Similarly, if $\mathbf{L} = \mathbf{I}_p$ then $\boldsymbol{\Omega} = \boldsymbol{\Sigma}_X$ and we get the following new expression for $\mathbf{Y}$
 350 reduction:

$$\text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{R}})) = f_p f_x^{-1} \text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\text{OLS}})) - f_x^{-1} \mathbf{R}_0 \boldsymbol{\Phi}_0 \mathbf{R}_0^T \otimes \boldsymbol{\Sigma}_X. \quad (3.11)$$

351 If the eigenvectors with larger eigenvalues of $\Sigma_{Y|X}$ lie in $\mathcal{E}_{\Sigma_{Y|X}}^\perp(\mathcal{R})$, then the variance of
 352 $\hat{\beta}_R$ may be reduced considerably since then Φ_0 will be large.

353 More importantly, the last term of the expansion in Proposition 2 represents a syn-
 354 ergy between the X and Y reductions that is not present in the marginal reductions. If
 355 the eigenvectors of $\Sigma_{Y|X}$ with large eigenvalues lie in $\mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R})$ or if the eigenvectors
 356 of Σ_X with small eigenvalues lie in $\mathcal{E}_{\Sigma_X}(\mathcal{L})$, then the variance reductions in either (3.11)
 357 or (3.10) could be insignificant. However, the synergy in simultaneous X and Y reduc-
 358 tions may still reduce the variance substantially because one of factors in the Kronecker
 359 product $f_x^{-1}R_0\Phi_0R_0^T \otimes L_0\Omega_0^{-1}L_0^T$ could be still be large.

360 Let $\mathbf{x}_N$ denote a new observation from $N(\mu_X, \Sigma_X)$, let $\mathbf{z}_N = \mathbf{x}_N - \mu_X$ be held fixed
 361 and consider $\text{var}(\hat{\beta}^T \mathbf{z}_N) = \text{var}\left((\mathbf{z}_N^T \otimes \mathbf{I})\text{vec}(\hat{\beta}^T)\right)$, which is the variance of a fitted
 362 vector for $\hat{\beta} = \hat{\beta}_{OLS}$ and $\hat{\beta} = \hat{\beta}_{L,R}$. It is straightforward from Proposition 2 that,

$$\begin{aligned} f_p \text{var}(\hat{\beta}_{OLS}^T \mathbf{z}_N) &= f_x \text{var}(\hat{\beta}_{L,R}^T \mathbf{z}_N) + R_0 \Phi_0 R_0^T (\mathbf{z}_N^T L \Omega^{-1} L^T \mathbf{z}_N) \\ &\quad + R \Phi R^T (\mathbf{z}_N^T L_0 \Omega_0^{-1} L_0^T \mathbf{z}_N) + R_0 \Phi_0 R_0^T (\mathbf{z}_N^T L_0 \Omega_0^{-1} L_0^T \mathbf{z}_N). \end{aligned}$$

363 We see from the above equality that only the part of $\mathbf{z}_N$ that lies in $\mathcal{E}_{\Sigma_X}(\mathcal{L})$ will contribute
 364 to the variance in prediction from $\hat{\beta}_{L,R}$, while the prediction variance from $\hat{\beta}_{OLS}$ depends
 365 on the whole of $\mathbf{z}_N$. If, in an extreme case, $\mathbf{z}_N \in \mathcal{E}_{\Sigma_X}^\perp(\mathcal{L})$, then $\text{var}(\hat{\beta}_{L,R}^T \mathbf{z}_N) = 0$.

366 4 Estimating envelopes

367 Let $\mathbf{C} = (\mathbf{X}^T, \mathbf{Y}^T)^T$ denote the concatenated random vectors, which has mean μ_C
 368 and covariance Σ_C , and let $\mathbf{S}_C$ be the sample covariance matrix of $\mathbf{C}$. In order to es-
 369 timate the parameters of the simultaneous envelope model (3.5), we introduce and probe
 370 a likelihood-based objective function that includes the objective functions in Cook, et al.
 371 (2010) and Cook, et al. (2013) as special cases. Variations on the objective function and
 372 their corresponding algorithms are also studied. We first give coordinate dependent and
 373 coordinate independent representations of Σ_C in Section 4.1 to facilitate estimation.

374 4.1 Structure of Σ_C

375 Since we already have the structure of Σ_X and $\Sigma_{Y|X}$ in (3.2) and (3.3), and $\Sigma_{XY} =$
 376 $\Sigma_X \beta = L \Omega \eta R^T$, we need only the following expression for Σ_Y to complete the neces-
 377 sary ingredients of Σ_C ,

$$\begin{aligned} \Sigma_Y &= \Sigma_{Y|X} + \Sigma_{XY}^T \Sigma_X^{-1} \Sigma_{XY} \\ &= R \Phi R^T + R_0 \Phi_0 R_0^T + R \eta^T \Omega L^T (L \Omega L^T + L_0 \Omega_0 L_0^T)^{-1} L \Omega \eta R^T \\ &= R(\Phi + \eta^T \Omega \eta) R^T + R_0 \Phi_0 R_0^T. \end{aligned}$$

378 Then we get the coordinate representation of the covariance Σ_C as

$$\begin{aligned} \Sigma_C &= \begin{pmatrix} \Sigma_X & \Sigma_{XY} \\ \Sigma_{XY}^T & \Sigma_Y \end{pmatrix} \\ &= \begin{pmatrix} L \Omega L^T + L_0 \Omega_0 L_0^T & L \Omega \eta R^T \\ R \eta^T \Omega L^T & R(\Phi + \eta^T \Omega \eta) R^T + R_0 \Phi_0 R_0^T \end{pmatrix}. \end{aligned} \quad (4.1)$$

379 By noticing that $\Sigma_X = P_L \Sigma_X P_L + Q_L \Sigma_X Q_L$, $\Sigma_{XY} = P_L \Sigma_{XY} P_R$ and $\Sigma_Y =$
 380 $P_R \Sigma_Y P_R + Q_R \Sigma_Y Q_R$ from the above expression, we can further obtain the coordi-
 381 nate independent representation

$$\begin{aligned} \Sigma_C &= P_{L \oplus R} \Sigma_C P_{L \oplus R} + Q_{L \oplus R} \Sigma_C Q_{L \oplus R}, \\ &= P_{L \oplus R} \Sigma_C P_{L \oplus R} + Q_{L \oplus R} \Sigma_D Q_{L \oplus R}, \end{aligned} \quad (4.2)$$

382 where $\Sigma_D \equiv \Sigma_X \oplus \Sigma_Y$ and $P_{L \oplus R} = P_L \oplus P_R$ is the projection onto the direct sum of
 383 two envelopes, $\mathcal{E}_{\Sigma_X}(\mathcal{L}) \oplus \mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R})$.

384 So far we have considered $\mathcal{E}_{\Sigma_X}(\mathcal{L})$ and $\mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R})$ as separate subspaces. Motivated
 385 by (4.2), the next lemma states that the direct sum of two arbitrary envelopes is itself an
 386 envelope. Let $M_1 \in \mathbb{S}^{p_1 \times p_1}$, $M_2 \in \mathbb{S}^{p_2 \times p_2}$, and let $\mathcal{S}_1$ and $\mathcal{S}_2$ be subspaces of $\text{span}(M_1)$
 387 and $\text{span}(M_2)$, which is required by Definition 2. Then

388 **Lemma 3.** $\mathcal{E}_{M_1}(\mathcal{S}_1) \oplus \mathcal{E}_{M_2}(\mathcal{S}_2) = \mathcal{E}_{M_1 \oplus M_2}(\mathcal{S}_1 \oplus \mathcal{S}_2)$.

389 From this lemma, we have $\mathcal{E}_{\Sigma_X}(\mathcal{L}) \oplus \mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R}) = \mathcal{E}_{\Sigma_X}(\mathcal{L}) \oplus \mathcal{E}_{\Sigma_Y}(\mathcal{R}) = \mathcal{E}_{\Sigma_X \oplus \Sigma_Y}(\mathcal{L} \oplus$
 390 $\mathcal{R}) = \mathcal{E}_{\Sigma_D}(\mathcal{L} \oplus \mathcal{R})$. We call this the simultaneous envelope for β .

4.2 The estimation criterion and resulting estimators

Assuming multivariate normality for $\mathbf{C}$, the negative log-likelihood minimized over $\mu_{\mathbf{C}}$ leads to the following objective function for estimation of $\Sigma_{\mathbf{C}}$,

$$F(\Sigma_{\mathbf{C}}) = \log |\Sigma_{\mathbf{C}}| + \text{trace}(\mathbf{S}_{\mathbf{C}}\Sigma_{\mathbf{C}}^{-1}). \quad (4.3)$$

We use this as a multi-purpose objective function, which gives the maximum likelihood estimator of β under normality of $\mathbf{C}$ and a $\sqrt{n}$ -consistent simultaneous envelope estimator of β when $\mathbf{C}$ has finite fourth moments. We also use $F(\cdot)$ as a generic objective function whose definition changes and is implied by its own argument. This should cause no confusion since $F(\cdot)$ will always be written with its arguments.

Substituting the coordinate form for $\Sigma_{\mathbf{C}}$ from (4.1) into $F(\Sigma_{\mathbf{C}})$ leads to an objective function that can be minimized explicitly over Ω , Ω_0 , Φ , Φ_0 and η with $\mathbf{R}$ and $\mathbf{L}$ held fixed. The resulting partially minimized objective function for simultaneous envelopes can be expressed as follows.

$$F(\mathbf{L} \oplus \mathbf{R}) = \log |(\mathbf{L}^T \oplus \mathbf{R}^T)\mathbf{S}_{\mathbf{C}}(\mathbf{L} \oplus \mathbf{R})| + \log |(\mathbf{L}^T \oplus \mathbf{R}^T)\mathbf{S}_{\mathbf{D}}^{-1}(\mathbf{L} \oplus \mathbf{R})|, \quad (4.4)$$

where $\mathbf{S}_{\mathbf{D}} = \mathbf{S}_{\mathbf{X}} \oplus \mathbf{S}_{\mathbf{Y}}$ is the sample version of $\Sigma_{\mathbf{D}}$. Let $|\mathbf{A}|_0$ denote the product of all nonzero eigenvalues of a matrix $\mathbf{A}$. Then we have the following coordinate independent representation of (4.4),

$$F(\mathbf{P}_{\mathbf{L} \oplus \mathbf{R}}) = \log |\mathbf{P}_{\mathbf{L} \oplus \mathbf{R}}\mathbf{S}_{\mathbf{C}}\mathbf{P}_{\mathbf{L} \oplus \mathbf{R}}|_0 + \log |\mathbf{P}_{\mathbf{L} \oplus \mathbf{R}}\mathbf{S}_{\mathbf{D}}^{-1}\mathbf{P}_{\mathbf{L} \oplus \mathbf{R}}|_0. \quad (4.5)$$

Moreover, if we substitute the coordinate free representation of $\Sigma_{\mathbf{C}}$ from (4.2) into $F(\Sigma_{\mathbf{C}})$, we immediately get (4.5).

Objective function (4.4) is invariant under orthogonal transformations: for any orthogonal $(d_X + d_Y) \times (d_X + d_Y)$ matrix $\mathbf{O}$, $F(\mathbf{L} \oplus \mathbf{R}) = F((\mathbf{L} \oplus \mathbf{R})\mathbf{O})$. The same result holds if we replace $\mathbf{S}_{\mathbf{C}}$ and $\mathbf{S}_{\mathbf{D}}$ with their population counterparts. Minimization of $F(\mathbf{L} \oplus \mathbf{R})$ is thus a Grassmann optimization problem with special direct sum structure. Consequently, neither $\mathbf{L}$ nor $\mathbf{R}$ is identified. However, $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$ and $\mathcal{E}_{\Sigma_{\mathbf{Y}|\mathbf{X}}}(\mathcal{R})$ are identified, and these are all that is needed to estimate β , $\Sigma_{\mathbf{X}}$ and $\Sigma_{\mathbf{Y}|\mathbf{X}}$. While the estimators of η , Ω , Ω_0 , Φ and Φ_0 depend on the particular bases chosen to minimize $F(\mathbf{L} \oplus \mathbf{R})$,

the estimators of β , $\Sigma_{\mathbf{X}}$ and $\Sigma_{\mathbf{Y}|\mathbf{X}}$ are independent of the choice since the estimated projection matrices $\mathbf{P}_{\hat{\mathbf{L}}}$ and $\mathbf{P}_{\hat{\mathbf{R}}}$ do not depend on the basis. In short, any values $\hat{\mathbf{L}}$ of $\mathbf{L}$ and $\hat{\mathbf{R}}$ of $\mathbf{R}$ that minimize $F(\mathbf{L} \oplus \mathbf{R})$ are allowed. Estimators of $\mathbf{L}_0$ and $\mathbf{R}_0$ are then any semi-orthogonal matrices $\hat{\mathbf{L}}_0$ and $\hat{\mathbf{R}}_0$ so that $(\hat{\mathbf{L}}, \hat{\mathbf{L}}_0)$ and $(\hat{\mathbf{R}}, \hat{\mathbf{R}}_0)$ are orthogonal matrices.

The next lemma summarizes the estimators that result from this process.

Lemma 4. *Let $\hat{\mathbf{L}} \oplus \hat{\mathbf{R}} = \arg \min F(\mathbf{L} \oplus \mathbf{R})$, where $\mathbf{L} \in \mathbb{R}^{p \times d_X}$ and $\mathbf{R} \in \mathbb{R}^{r \times d_Y}$ are semi-orthogonal basis matrices. Then the estimators of the remaining parameters are $\hat{\Omega} = \hat{\mathbf{L}}^T \mathbf{S}_{\mathbf{X}} \hat{\mathbf{L}}$, $\hat{\Omega}_0 = \hat{\mathbf{L}}_0^T \mathbf{S}_{\mathbf{X}} \hat{\mathbf{L}}_0$, $\hat{\eta} = (\hat{\mathbf{L}}^T \mathbf{S}_{\mathbf{X}} \hat{\mathbf{L}})^{-1} (\hat{\mathbf{L}}^T \mathbf{S}_{\mathbf{X}\mathbf{Y}} \hat{\mathbf{R}})$, $\hat{\Phi}_0 = \hat{\mathbf{R}}_0^T \mathbf{S}_{\mathbf{Y}} \hat{\mathbf{R}}_0$ and*

$$\hat{\Phi} = \hat{\mathbf{R}}^T \left(\mathbf{S}_{\mathbf{Y}} - \mathbf{S}_{\mathbf{X}\mathbf{Y}}^T \hat{\mathbf{L}} \left(\hat{\mathbf{L}}^T \mathbf{S}_{\mathbf{X}} \hat{\mathbf{L}} \right)^{-1} \hat{\mathbf{L}}^T \mathbf{S}_{\mathbf{X}\mathbf{Y}} \right) \hat{\mathbf{R}}.$$

The simultaneous envelope estimator for β is then

$$\hat{\beta} = \hat{\mathbf{L}} \hat{\eta} \hat{\mathbf{R}}^T = \mathbf{P}_{\hat{\mathbf{L}}(\mathbf{S}_{\mathbf{X}})} \hat{\beta}_{\text{OLS}} \mathbf{P}_{\hat{\mathbf{R}}}. \quad (4.6)$$

The simultaneous envelope estimator for β in (4.6) coincides with the plug-in envelope estimator $\hat{\beta}_{\hat{\mathbf{L}}, \hat{\mathbf{R}}}$ obtained by regarding estimated $\hat{\mathbf{L}}$ and $\hat{\mathbf{R}}$ as the known values of $\mathbf{L}$ and $\mathbf{R}$. We next turn to methods for minimizing (4.4).

4.3 Alternating algorithm

If we fix $\mathbf{L}$ as an arbitrary orthogonal basis, then the objective function $F(\mathbf{L} \oplus \mathbf{R})$ in (4.4) can be re-expressed as an objective function in $\mathbf{R}$:

$$F(\mathbf{R}|\mathbf{L}) = \log |\mathbf{R}^T \mathbf{S}_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}} \mathbf{R}| + \log |\mathbf{R}^T \mathbf{S}_{\mathbf{Y}}^{-1} \mathbf{R}|. \quad (4.7)$$

Similarly, if we fix $\mathbf{R}$, the objective function $F(\mathbf{L} \oplus \mathbf{R})$ reduces to the conditional objective function

$$F(\mathbf{L}|\mathbf{R}) = \log |\mathbf{L}^T \mathbf{S}_{\mathbf{X}|\mathbf{R}^T \mathbf{Y}} \mathbf{L}| + \log |\mathbf{L}^T \mathbf{S}_{\mathbf{X}}^{-1} \mathbf{L}|. \quad (4.8)$$

We use the following alternating algorithm based on (4.7) and (4.8) to obtain a minimizer of the objective function $F(\mathbf{L} \oplus \mathbf{R})$ in (4.4).

1. Initialization. Set the starting value $\mathbf{L}^{(0)}$ and get $\mathbf{R}^{(0)} = \arg \min_{\mathbf{R}} F(\mathbf{R}|\mathbf{L}^{(0)})$.

- 435 2. Alternating. For the k -th stage, obtain $\mathbf{L}^{(k)} = \arg \min_{\mathbf{L}} F(\mathbf{L}|\mathbf{R} = \mathbf{R}^{(k-1)})$ and
 436 $\mathbf{R}^{(k)} = \arg \min_{\mathbf{R}} F(\mathbf{R}|\mathbf{L} = \mathbf{L}^{(k)})$.
- 437 3. Convergence criterion. Evaluate $\{F(\mathbf{L}^{(k-1)} \oplus \mathbf{R}^{(k-1)}) - F(\mathbf{L}^{(k)} \oplus \mathbf{R}^{(k)})\}$ and re-
 438 turn to the alternating step if it is bigger than a tolerance; otherwise, stop the itera-
 439 tion and use $\mathbf{L}^{(k)} \oplus \mathbf{R}^{(k)}$ as the final estimator.

440 In the initialization step of the algorithm, we could also set $\mathbf{R}^{(0)}$ to be some initial value
 441 and get $\mathbf{L}^{(0)} = \arg \min_{\mathbf{L}} F(\mathbf{L}|\mathbf{R}^{(0)})$. Then interchanging the roles of $\mathbf{L}$ and $\mathbf{R}$ in the
 442 alternating step will give us another algorithm, which has the same performance as the
 443 alternating algorithm outlined above. Comparing to the objective function used by Cook
 444 et al. (2010), we see that $F(\mathbf{R}|\mathbf{L})$ is the objective function for estimating the $\mathbf{Y}$ -envelope
 445 in the regression of $\mathbf{Y}$ on the reduced predictors $\mathbf{L}^T \mathbf{X}$ for fixed $\mathbf{L}$. Similarly from the
 446 objective function used by Cook et al. (2013), we notice that $F(\mathbf{L}|\mathbf{R})$ is the objective
 447 function for estimating the $\mathbf{X}$ -envelope in the regression of the reduced responses $\mathbf{R}^T \mathbf{Y}$
 448 on the predictors $\mathbf{X}$ for fixed $\mathbf{R}$.

449 In our experience, as long as we use good initial values, the alternating algorithm,
 450 which monotonically decreases $F(\mathbf{L} \oplus \mathbf{R})$, will converge after only a few cycles, typically
 451 less than four. Root- n consistent starting values are particularly important to mitigate
 452 potential problems caused by multiple local minima and to ensure efficient estimation.
 453 For instance, under joint normality of $\mathbf{X}$ and $\mathbf{Y}$, one Newton-Raphson iteration from any
 454 $\sqrt{n}$ -consistent estimator will be asymptotically as efficient as the maximum likelihood
 455 estimator, even if local minima are present (Small et al., 2000; Lehmann and Casella,
 456 1998, Theorem 4.3).

457 4.4 One-dimensional optimization

458 In this section we propose a fast 1D algorithm for $\sqrt{n}$ -consistent envelope estimation.
 459 The algorithm, which does not itself require starting values, can be used for standalone
 460 estimation or to obtain starting values for the alternating algorithm. We state our algo-
 461 rithm in terms of estimating $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$ from $f(\mathbf{L}) \equiv F(\mathbf{L}|\mathbf{R} = \mathbf{I}_r)$. It works similarly for
 462 estimating $\mathcal{E}_{\Sigma_{\mathbf{Y}|\mathbf{X}}}(\mathcal{R})$ and for estimating a general envelope $\mathcal{E}_{\mathbf{M}}(\mathcal{S})$.

Optimization over $\mathcal{G}_{(p,d_X)}$, as required in each cycle of the alternating algorithm, can be computationally expensive. The next lemma offers a way to break down the d_X -dimensional manifold into fast 1D components.

Lemma 5. *Let $\mathbf{M} \in \mathbb{S}^{p \times p}$ be a positive definite matrix and let $\mathcal{S} \subset \mathbb{R}^p$. Suppose $(\mathbf{B}, \mathbf{B}_0)$ is an orthogonal basis of $\mathbb{R}^p$ where $\mathbf{B} \in \mathbb{R}^{p \times q}$, $\mathbf{B}_0 \in \mathbb{R}^{p \times (p-q)}$ and $\text{span}(\mathbf{B}) \subseteq \mathcal{E}_{\mathbf{M}}(\mathcal{S})$. Then $\mathbf{v} \in \mathcal{E}_{\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0}(\mathbf{B}_0^T \mathcal{S})$ implies $\mathbf{B}_0 \mathbf{v} \in \mathcal{E}_{\mathbf{M}}(\mathcal{S})$.*

Suppose we have an estimate $\mathbf{B}$ of q directions in the envelope $\mathcal{E}_{\mathbf{M}}(\mathcal{S})$. Then we can estimate the rest of $\mathcal{E}_{\mathbf{M}}(\mathcal{S})$ by focusing our attention on $\mathcal{E}_{\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0}(\mathbf{B}_0^T \mathcal{S})$, which is a lower dimensional envelope. This leads to the following 1D algorithm.

Let $\ell_k \in \mathbb{R}^p$, $k = 1, \dots, d_X$, be the stepwise directions obtained. Let $\mathbf{L}_k = (\ell_1, \dots, \ell_k)$, let $(\mathbf{L}_k, \mathbf{L}_{0k})$ be an orthogonal basis for $\mathbb{R}^p$ and set initial value $\ell_0 = \mathbf{L}_0 = 0$, then for $k = 0, \dots, d_X - 1$, get

$$\begin{aligned} \mathbf{v}_{k+1} &= \arg \min_{\mathbf{v} \in \mathbb{R}^{p-k}} f_k(\mathbf{v}), \text{ subject to } \mathbf{v}^T \mathbf{v} = 1, \\ \ell_{k+1} &= \mathbf{L}_{0k} \mathbf{v}_{k+1}, \end{aligned} \quad (4.9)$$

where $f_k(\mathbf{v}) = \log |\mathbf{v}^T \mathbf{L}_{0k}^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \mathbf{L}_{0k} \mathbf{v}| + \log |\mathbf{v}^T (\mathbf{L}_{0k}^T \mathbf{S}_{\mathbf{X}} \mathbf{L}_{0k})^{-1} \mathbf{v}|$ for $\mathbf{v} \in \mathbb{R}^{p-k}$. This algorithm does not require an initial value search. It starts simply from $\ell_0 = \mathbf{L}_0 = 0$ and builds the estimator sequentially.

Suppose we know $\text{span}(\mathbf{L}_k) \subset \mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$, then we can estimate the remaining part of $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$ by substituting $\mathbf{L} = (\mathbf{L}_k, \mathbf{L}_{0k} \mathbf{V})$, where $\mathbf{V} \in \mathbb{R}^{(p-k) \times (d_X-k)}$, into $f(\mathbf{L})$, leading to the objective function $\log |\mathbf{V}^T \mathbf{L}_{0k}^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \mathbf{L}_{0k} \mathbf{V}| + \log |\mathbf{V}^T (\mathbf{L}_{0k}^T \mathbf{S}_{\mathbf{X}} \mathbf{L}_{0k})^{-1} \mathbf{V}|$, whose one dimensional version is exactly $f_k(\mathbf{v})$. This describes how the 1D algorithm is minimizing the full objective function $f(\mathbf{L})$ sequentially. Let $\hat{\mathbf{L}}_{1D}$ denote the estimator from this algorithm.

The following proposition formally states the $\sqrt{n}$ -consistency of the 1D algorithm for estimating $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$. Similar results hold for estimating $\mathcal{E}_{\Sigma_{\mathbf{Y}|\mathbf{X}}}(\mathcal{R})$ and for estimating an arbitrary envelope using the 1D algorithm.

Proposition 3. *Assume that $\mathbf{S}_{\mathbf{X}}$ and $\mathbf{S}_{\mathbf{X}|\mathbf{Y}}$ are $\sqrt{n}$ -consistent estimators of $\Sigma_{\mathbf{X}}$ and $\Sigma_{\mathbf{X}|\mathbf{Y}}$. Let $\hat{\mathbf{L}}_{1D}$ denote the estimator obtained from the 1D algorithm. Then $\mathbf{P}_{\hat{\mathbf{L}}_{1D}}$ is $\sqrt{n}$ -consistent for the projection onto $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$.*

490 We have also considered sequentially minimizing $f(\ell) = \log |\ell^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell| + \log |\ell^T \mathbf{S}_{\mathbf{X}}^{-1} \ell|$,
 491 so that $\ell_i^T \ell_j = 0$ if $i \neq j$ and equals one if $i = j$. Although this seems to be a reasonable
 492 way of sequentially constructing the envelope, we do not know if this procedure results in
 493 a consistent estimator. Our numerical experience is that the 1D algorithm (4.9) performs
 494 much better than such sequential optimization.

495 The 1D algorithm can be used to get fast $\sqrt{n}$ -consistent starting values for the al-
 496 ternating algorithm in Section 4.3 by separately estimating the $\mathbf{X}$ -envelope and the $\mathbf{Y}$ -
 497 envelope bases. Since the 1D algorithm turns the optimizations over d_X and d_Y -dimensional
 498 manifolds to d_X and d_Y sequential optimizations over 1D manifolds, the computation
 499 complexity is reduced drastically. For the simulation examples we studied, the computa-
 500 tion costs of the 1D algorithm are from tens to hundreds times less than the computation
 501 costs of using d_X and d_Y -dimensional Grassmann manifold optimizations.

502 Perhaps more importantly, we have found $\hat{\mathbf{L}}_{1D}$ and $\hat{\mathbf{R}}_{1D}$ to be practically as efficient
 503 as the final estimators obtained by alternating because the alternating algorithm nearly
 504 always converges after only a few iterations and produces only small changes. To em-
 505 phasize the utility of the 1D estimators, we use them in the simulations and real data
 506 examples that follow in later sections. In Section 7.3, we use a simulation example to
 507 demonstrate that estimators from the 1D manifold algorithm may have the same behavior
 508 as maximum likelihood estimators.

509 5 Asymptotic properties

510 The parameters involved in the coordinate representation of the simultaneous envelope
 511 model are $\boldsymbol{\eta}$, $\boldsymbol{\Omega}$, $\boldsymbol{\Omega}_0$, $\boldsymbol{\Phi}$, $\boldsymbol{\Phi}_0$, $\mathbf{L}$ and $\mathbf{R}$. In Sections 5.1 and 5.2, we focus on the asymptotic
 512 properties of the estimable functions $\boldsymbol{\beta} = \mathbf{L}\boldsymbol{\eta}\mathbf{R}^T$, $\boldsymbol{\Sigma}_{\mathbf{X}} = \mathbf{L}\boldsymbol{\Omega}\mathbf{L}^T + \mathbf{L}_0\boldsymbol{\Omega}_0\mathbf{L}_0^T$, and $\boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}} =$
 513 $\mathbf{R}\boldsymbol{\Phi}\mathbf{R}^T + \mathbf{R}_0\boldsymbol{\Phi}_0\mathbf{R}_0^T$. Specifically, we study the asymptotic covariances of the following

parameters ϕ and estimable functions $\mathbf{h}$.

$$\phi = \begin{pmatrix} \phi_1 \\ \phi_2 \\ \phi_3 \\ \phi_4 \\ \phi_5 \\ \phi_6 \\ \phi_7 \end{pmatrix} \equiv \begin{pmatrix} \text{vec}(\boldsymbol{\eta}) \\ \text{vec}(\mathbf{L}) \\ \text{vec}(\mathbf{R}) \\ \text{vech}(\boldsymbol{\Omega}) \\ \text{vech}(\boldsymbol{\Omega}_0) \\ \text{vech}(\boldsymbol{\Phi}) \\ \text{vech}(\boldsymbol{\Phi}_0) \end{pmatrix}, \quad \mathbf{h} = \begin{pmatrix} \mathbf{h}_1(\phi) \\ \mathbf{h}_2(\phi) \\ \mathbf{h}_3(\phi) \end{pmatrix} \equiv \begin{pmatrix} \text{vec}(\boldsymbol{\beta}) \\ \text{vech}(\boldsymbol{\Sigma}_{\mathbf{X}}) \\ \text{vech}(\boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}}) \end{pmatrix}.$$

Since $\mathbf{h} \in \mathbb{R}^{\frac{1}{2}(p+r)(p+r+1)}$ and $\phi \in \mathbb{R}^{\frac{1}{2}p(p+1) + \frac{1}{2}r(r+1) + d_X d_Y}$, the simultaneous envelope model reduces the total number of parameters by $\frac{1}{2}(p+r)(p+r+1) - \frac{1}{2}p(p+1) - \frac{1}{2}r(r+1) - d_X d_Y = pr - d_X d_Y$.

5.1 Without the normality assumption

Let $\hat{\mathbf{h}}_{\text{full}} = \left(\text{vec}^T(\hat{\boldsymbol{\beta}}_{\text{OLS}}), \text{vech}^T(\mathbf{S}_{\mathbf{X}}), \text{vech}^T(\mathbf{S}_{\mathbf{Y}|\mathbf{X}}) \right)^T$ denote the OLS estimator of $\mathbf{h}$ under the standard model (3.1), and let $\hat{\mathbf{h}}$ denote the estimator from Lemma 4 under the simultaneous envelope model (3.5). The true values of $\mathbf{h}$ and ϕ are denoted as $\mathbf{h}_0$ and ϕ_0 . Define $\boldsymbol{\Delta} \equiv \partial \mathbf{h}(\phi) / \partial \phi$ to be the gradient matrix, whose explicit form is given in the Supplement, Section I.1. We use “avar” to denote an asymptotic covariance matrix: $\text{avar}(\sqrt{n}\hat{\mathbf{h}}) = \mathbf{A}$ is equivalent to $\sqrt{n}(\hat{\mathbf{h}} - \mathbf{h}_0) \rightarrow N(0, \mathbf{A})$ in distribution. We use the expansion matrix (Henderson and Searle, 1979) $\mathbf{E}_p \in \mathbb{R}^{p^2 \times p(p+1)/2}$ to relate the vec operation and vech operation: for any symmetric matrix $\mathbf{M} \in \mathbb{S}^{p \times p}$, $\text{vec}(\mathbf{M}) = \mathbf{E}_p \text{vech}(\mathbf{M})$. Then $\sqrt{n}(\hat{\mathbf{h}}_{\text{full}} - \mathbf{h}_0) \rightarrow N(0, \boldsymbol{\Gamma})$, for some positive definite covariance matrix $\boldsymbol{\Gamma}$. Since there is a one-to-one relationship between $\mathbf{h}$ and $\boldsymbol{\Sigma}_{\mathbf{C}}$, we treat the objective function $F(\boldsymbol{\Sigma}_{\mathbf{C}})$ in (4.3) as a function of $\mathbf{h}$ and $\hat{\mathbf{h}}_{\text{full}}$ and write it as $F(\mathbf{h}, \hat{\mathbf{h}}_{\text{full}})$. Let $\mathbf{J}_{\mathbf{h}} = 1/2 \times \partial^2 F(\mathbf{h}, \hat{\mathbf{h}}_{\text{full}}) / \partial \mathbf{h} \partial \mathbf{h}^T$ evaluated at $\hat{\mathbf{h}}_{\text{full}} = \mathbf{h} = \mathbf{h}_0$, which is the Fisher information matrix for $\mathbf{h}$ when $\mathbf{C}$ is normal,

$$\mathbf{J}_{\mathbf{h}} = \begin{pmatrix} \boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}} & 0 & 0 \\ 0 & \frac{1}{2} \mathbf{E}_p^T (\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}) \mathbf{E}_p & 0 \\ 0 & 0 & \frac{1}{2} \mathbf{E}_r^T (\boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}}^{-1}) \mathbf{E}_r \end{pmatrix}. \quad (5.1)$$

The following proposition formally states the asymptotic distribution of $\hat{\mathbf{h}}$.

Proposition 4. *Assume that the data $(\mathbf{x}_i, \mathbf{y}_i)$ are i.i.d. from a joint distribution with finite fourth moments. Then $\sqrt{n}(\text{vec}(\hat{\mathbf{h}}) - \text{vec}(\mathbf{h}_0))$ converges in distribution to a normal*

random variable with mean $\mathbf{0}$ and covariance matrix

$$\mathbf{W} = \Delta (\Delta^T \mathbf{J}_h \Delta)^\dagger \Delta^T \mathbf{J}_h \Gamma \mathbf{J}_h \Delta (\Delta^T \mathbf{J}_h \Delta)^\dagger \Delta^T,$$

where $\dagger$ indicates the Moore-Penrose inverse. In particular, $\sqrt{n}(\text{vec}(\hat{\beta}) - \text{vec}(\beta))$ converges in distribution to a normal random variable with mean $\mathbf{0}$ and covariance $\mathbf{W}_{11}$, the upper-left $pr \times pr$ block of $\mathbf{W}$. Moreover, $\text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{full}}) \geq \text{avar}(\sqrt{n}\hat{\mathbf{h}})$ if $\text{span}(\mathbf{J}_h^{1/2} \Delta)$ is a reducing subspace of $\mathbf{J}_h^{1/2} \Gamma \mathbf{J}_h^{1/2}$.

The $\sqrt{n}$ -consistency of the estimator $\hat{\beta}$ is essentially because $\mathbf{S}_X$, $\mathbf{S}_Y$ and $\mathbf{S}_{XY}$ are $\sqrt{n}$ -consistent and because of the properties of $F(\mathbf{h}, \hat{\mathbf{h}}_{\text{full}})$. The asymptotic covariance matrix $\mathbf{W}_{11}$ can be computed straightforwardly, but its accuracy for any fixed sample size may depend on the distribution of $\mathbf{C}$. Fortunately, bootstrap methods can provide a good approximation of $\mathbf{W}_{11}$, as discussed in Section 5.3.

5.2 Under the normality assumption

The asymptotic covariance $\mathbf{W}$ simplifies when $\mathbf{C} \sim N(\mu_C, \Sigma_C)$ because then objective function $F(\mathbf{h}, \hat{\mathbf{h}}_{\text{full}})$ agrees with the negative log-likelihood function and $\Gamma = \mathbf{J}_h^{-1} = \text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{full}})$. And $\{\text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{full}}) - \text{avar}(\sqrt{n}\hat{\mathbf{h}})\}$ will be positive semi-definite because $\text{span}(\mathbf{J}_h^{1/2} \Delta)$ is always a reducing subspace of $\mathbf{J}_h^{1/2} \Gamma \mathbf{J}_h^{1/2} = \mathbf{I}$.

Proposition 5. Assume that $\mathbf{C} \sim N(\mu_C, \Sigma_C)$. Then $\text{avar}(\sqrt{n}\hat{\mathbf{h}}) = \Delta(\Delta^T \mathbf{J}_h \Delta)^\dagger \Delta^T$. Moreover, $\text{avar}(\sqrt{n}\hat{\mathbf{h}}) \leq \text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{full}})$,

$$\mathbf{J}_h^{-1} - \Delta(\Delta^T \mathbf{J}_h \Delta)^\dagger \Delta^T = \mathbf{J}_h^{-\frac{1}{2}} \mathbf{Q}_{\mathbf{J}_h^{\frac{1}{2}} \Delta} \mathbf{J}_h^{-\frac{1}{2}} \geq 0.$$

In particular, $\text{avar}(\sqrt{n}\text{vec}(\hat{\beta})) \leq \text{avar}(\sqrt{n}\text{vec}(\hat{\beta}_{\text{OLS}}))$.

This proposition is obtained by direct computation, and is consistent with Cook et al. (2010) and Cook et al. (2013). Also, because the estimator $\hat{\beta}$ is the MLE under normality, its asymptotic variance is no larger than that of $\mathbf{X}$ -envelope estimator or $\mathbf{Y}$ -envelope estimator. Explicit expressions can be found in the Supplement, Section I. To further interpret this result, we next consider the asymptotic variance of $\text{vec}(\hat{\beta})$, which can lead to the asymptotic variances for predictions and fitted values.

556 **Proposition 6.** Assume that $\mathbf{C} \sim N(\boldsymbol{\mu}_C, \boldsymbol{\Sigma}_C)$. Then

$$\begin{aligned} \text{avar}(\sqrt{n}\text{vec}(\hat{\boldsymbol{\beta}})) &= \text{avar}(\sqrt{n}\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}})) + \text{avar}\left(\sqrt{n}\text{vec}(\mathbf{Q}_{\mathbf{L}}\hat{\boldsymbol{\beta}}_{\boldsymbol{\eta},\mathbf{R}})\right) \\ &+ \text{avar}\left(\sqrt{n}\text{vec}(\hat{\boldsymbol{\beta}}_{\boldsymbol{\eta},\mathbf{L}}\mathbf{Q}_{\mathbf{R}})\right), \end{aligned}$$

557 where we use $\hat{\boldsymbol{\beta}}_{\boldsymbol{\eta},\mathbf{R}}$, $\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}}$ and $\hat{\boldsymbol{\beta}}_{\boldsymbol{\eta},\mathbf{L}}$ to denote the maximum likelihood estimators when
558 the parameters in their subscripts are known. Explicit expressions for the asymptotic
559 variances of $\hat{\boldsymbol{\beta}}_{\boldsymbol{\eta},\mathbf{R}}$, $\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}}$ and $\hat{\boldsymbol{\beta}}_{\boldsymbol{\eta},\mathbf{L}}$ are given in the Supplement, Section I.

560 The first term in the above decomposition, $\text{avar}(\sqrt{n}\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}}))$, is the same as that in
561 Proposition 2, which gave the asymptotic variance if we knew bases for the true envelope.
562 If we set $\mathbf{L} = \mathbf{I}_p$ in the above decomposition, so we are pursuing $\mathbf{Y}$ reduction only, then
563 the decomposition reduces to that given by Cook et al. (2010; Theorem 6.1). Setting
564 $\mathbf{R} = \mathbf{I}_r$, which indicates $\mathbf{X}$ reduction only gives the corresponding result of Cook et al.
565 (2013; Proposition 4.4). The projections $\mathbf{Q}_{\mathbf{L}}$ and $\mathbf{Q}_{\mathbf{R}}$ serve to orthogonalize the random
566 vectors so that the asymptotic variances are additive.

567 5.3 Residual bootstrap

568 To illustrate the application of the above bootstrap method, we consider a simple model
569 with $p = r = 3$ and $d_X = d_Y = 1$. We generated $\mathbf{L}$, $\mathbf{R}$ and $\boldsymbol{\eta}$ by filling in random
570 numbers from uniform(0, 1) distribution. Then we orthonormalized $\mathbf{L}$, $\mathbf{R}$ and obtained
571 corresponding $\mathbf{L}_0$, $\mathbf{R}_0$. The covariance matrices were $\boldsymbol{\Omega} = 5$, $\boldsymbol{\Omega}_0 = \mathbf{I}_2$, $\boldsymbol{\Phi} = 1$ and
572 $\boldsymbol{\Phi}_0 = 10\mathbf{I}_2$. The data vectors $\mathbf{C}_i = (\mathbf{X}_i^T, \mathbf{Y}_i^T)^T$, $i = 1, \dots, n$, were simulated from
573 $\boldsymbol{\Sigma}_{\mathbf{C}}^{1/2}\mathbf{U}_i$ where $\mathbf{U}_i$ is a vector of i.i.d uniform random variables with mean 0 standard
574 deviation 1. Therefore, $\mathbf{C}_i$ would follow a distribution with mean 0, covariance $\boldsymbol{\Sigma}_{\mathbf{C}}$ and
575 finite fourth moments. A dataset with $n = 100$ observations was generated and $B = 100$
576 bootstrap datasets were used throughout. We used the bootstrap method to estimate the
577 variances of two estimators: the OLS estimator $\hat{\boldsymbol{\beta}}_{\text{OLS}}$ and the simultaneous envelope
578 estimator $\hat{\boldsymbol{\beta}}_{\text{1D}}$ which was obtained by using the 1D algorithm without alternating.

579 Table 1 summarizes all the $p \times r = 9$ elements of $\text{vec}(\boldsymbol{\beta})$, $\text{vec}(\hat{\boldsymbol{\beta}}_{\text{OLS}})$ and $\text{vec}(\hat{\boldsymbol{\beta}}_{\text{1D}})$
580 and their asymptotic, bootstrap and actual standard errors. We included the asymptotic
581 standard errors of the elements in $\hat{\boldsymbol{\beta}}_{\text{OLS}}$, which were the square roots of diagonals in

True parameter									
$\text{vec}(\beta) \times 100$	30	30	32	36	36	38	14	14	15
OLS estimator									
$\text{vec}(\hat{\beta}_{\text{OLS}}) \times 100$	37	37	14	29	42	53	39	2	15
$E_N \left\{ \text{vec}(\hat{\beta}_{\text{OLS}}) \right\} \times 100$	29	29	32	37	36	38	13	17	13
Asymptotic se $\times 100$	22	22	22	20	20	19	26	26	26
Bootstrap se $\times 100$	21	20	17	17	17	16	25	28	25
Actual se $\times 100$	22	23	22	21	22	18	30	24	23
Simultaneous envelope estimator									
$\text{vec}(\hat{\beta}_{\text{1D}}) \times 100$	31	34	30	39	43	38	16	18	16
$E_N \left\{ \text{vec}(\hat{\beta}_{\text{1D}}) \right\} \times 100$	30	30	32	36	35	38	14	14	15
Asymptotic se $\times 100$	3	3	3	3	3	3	2	2	2
Bootstrap se $\times 100$	3	2	3	3	3	3	2	2	2
Actual se $\times 100$	3	2	3	3	3	3	2	2	2

Table 1: Bootstrap and estimated asymptotic standard errors of the 9 elements in $\hat{\beta}$ under the OLS estimator and the simultaneous envelope estimator with fixed $d_X = d_Y = 1$.

582 $\Sigma_{Y|X} \otimes \Sigma_X^{-1}/n$ and the asymptotic standard errors of the MLE under normality for the
583 simultaneous envelope model, which is given in the Supplement, in order to compare with
584 the estimator from the 1D manifold algorithm. We also repeatedly simulated $N = 100$
585 datasets under the same setting, and estimated the averaged coefficient estimates $E_N(\hat{\beta})$
586 and the actual standard errors of $\text{vec}(\hat{\beta})$ from the $N = 100$ replicates.

587 The asymptotic standard errors and the actual standard errors were well estimated
588 by the bootstrap method. Moreover, the estimated asymptotic standard errors for the
589 1D algorithm are really close to the asymptotic standard errors of the MLE although
590 the normality assumption was violated. As expected, the envelope estimator had much
591 smaller standard errors than that of the OLS estimator.

592 6 Selecting the dimensions

593 6.1 Rank of β

594 Since the simultaneous envelope contains the row and the column space of β , the di-
595 mensions d_X and d_Y are bounded below by $d = \text{rank}(\beta)$. Thus when determine the
596 dimensions d_X and d_Y , it is helpful to first have some guidance on d . Bura and Cook

(2003) developed a chi-squared test for the rank d that requires only that the response variables have finite second moments. The test statistic is $\Lambda_d = n \sum_{j=d+1}^{\min(p,r)} \varphi_j^2$, where $\varphi_1 \geq \dots \geq \varphi_{\min(p,r)}$ are eigenvalues of the $p \times r$ matrix

$$\hat{\beta}_{\text{std}} = \{(n - p - 1)\}/n\}^{1/2} \mathbf{S}_{\mathbf{X}}^{1/2} \hat{\beta}_{\text{OLS}} \mathbf{S}_{\mathbf{Y}|\mathbf{X}}^{-1/2}. \quad (6.1)$$

Then they derived that Λ_d is asymptotically distributed as a $\chi_{(p-d)(r-d)}^2$ random variable under the null hypothesis that $\text{rank}(\beta) = d$. The rank is then determined by comparing a sequence of test statistics Λ_k , $k = 0, \dots, \min(p, r) - 1$, to the percentiles of their null distribution $\chi_{(p-k)(r-k)}^2$. The first non-significant value of k will serve as our estimate of the rank of β .

6.2 Selecting d_X and d_Y

One way to determine d_X and d_Y is by using a sequence of likelihood ratio tests based on the statistic $\Lambda_{d_X, d_Y} = 2(\hat{L}_{\text{full}} - \hat{L}_{d_X, d_Y})$, where $\hat{L}_{\text{full}} = -(n/2) \log |\mathbf{S}_{\mathbf{C}}| - n(p + r)/2$ is the log likelihood for the standard model and $\hat{L}_{d_X, d_Y} = -(n/2) \log |\hat{\Sigma}_{\mathbf{C}}| - (n/2) \text{trace}(\mathbf{S}_{\mathbf{C}} \hat{\Sigma}_{\mathbf{C}}^{-1})$ is the maximum log likelihood for simultaneous envelope model, where $\hat{\Sigma}_{\mathbf{C}} = \mathbf{P}_{\hat{\mathbf{L}} \oplus \hat{\mathbf{R}}} \mathbf{S}_{\mathbf{C}} \mathbf{P}_{\hat{\mathbf{L}} \oplus \hat{\mathbf{R}}} + \mathbf{Q}_{\hat{\mathbf{L}} \oplus \hat{\mathbf{R}}} \mathbf{S}_{\mathbf{D}} \mathbf{Q}_{\hat{\mathbf{L}} \oplus \hat{\mathbf{R}}}$ is the MLE of the simultaneous envelope model. Under the null hypothesis, the test statistic Λ_{d_X, d_Y} is asymptotically distributed as a $\chi_{(p-d_X)(r-d_Y)}^2$ random variable. We start testing with $d_X = d_Y = d_0$ where the choice of d_0 can be guided by the Bura-Cook estimator. Then sequentially test $d_X = d_0, d_0 + 1, \dots$, until first non-significant value, and sequentially test $d_Y = d_0, d_0 + 1, \dots$, until first non-significant value.

Information criteria such as Akaike's information criterion (AIC) and Bayesian information criterion (BIC) could also be used to select the dimension d_X and d_Y , similar to Cook et al. (2010). For dimensions d_X and d_Y , AIC is $A(d_X, d_Y) = 2k_{d_X, d_Y} - 2\hat{L}_{d_X, d_Y}$, where $k_{d_X, d_Y} = p(p + 3)/2 + r(r + 3)/2 + d_X d_Y$ is the number of parameters in the model; and BIC is $B(d_X, d_Y) = k_{d_X, d_Y} \log(n) - 2\hat{L}_{d_X, d_Y}$. We search (d_X, d_Y) from (d_0, d_0) to (p, r) and choose the pair that has the smallest AIC or BIC. See Section 7.3 for a simulation to illustrate determining the dimensions by likelihood ratio testing, AIC and BIC.

Our experience suggests that BIC is the most favorable criterion for determining dimensions when sample size is not too small. Since the true dimension always exist, BIC is consistent in the sense that the probability of selecting the correct dimension approaches 1. According to Shao (1997) and Yang (2005), BIC would perform better than AIC when the true model has a simple finite dimensional structure. However, if one is concerned more about the correctness of the envelope model itself over its simplicity, AIC is more favorable since it is more conservative on choosing the dimensions.

In practice, we use cross-validation to determine the optimal dimensions for prediction, although the best predictive dimensions can be different from the true envelope dimensions d_X and d_Y .

7 Simulations

In practice, we found the estimator from the 1D algorithm essentially as efficient as the final estimator by alternating because the alternating procedure usually converged after only a few iterations and produced only small changes. To emphasize the 1D algorithm, we use it for all the simulations and examples in this section. In Section 7.2, we use a simulation example to demonstrate that estimators from the 1D manifold algorithm may have the same behavior as maximum likelihood estimators.

In Section 7.1, we compared the simultaneous envelope estimator to many other methods including OLS, PLS, CCA, RRR, X-envelope and Y-envelope, where the dimensions for each method were chosen purely based on the simulated data with no prior knowledge. In the Supplement Section J, the true dimensions were used for each method and we also used two different criterion for prediction and estimation. Given the true dimension, the simultaneous envelope estimators always had the best performance in both prediction and estimation among all estimators and always performed substantially better than OLS.

7.1 Prediction with cross-validation

Prediction error of a multivariate linear model is a composite of the variability in $\hat{\beta}$ and the intrinsic variance in ϵ . As illustrated by Table 1, the simultaneous envelope method can substantially increase estimation accuracy in the regression coefficient matrix and

thus it may also increase prediction accuracy. We simulated such situations in the following examples where we generated data from joint normal distribution of $\mathbf{X}$ and $\mathbf{Y}$.

We simulated data from two envelope models (E1 and E2) and two reduced-rank regression models (R1 and R2) as follows. The dimensions were $p = 10$, $r = 15$ and varied d , d_X and d_Y . For the envelope models, we generated $\mathbf{L}$, $\mathbf{R}$ and $\boldsymbol{\eta}$ by filling in random numbers from the uniform(0,1) distribution. Then we orthogonalized $\mathbf{L}$ and $\mathbf{R}$ and obtained corresponding $\mathbf{L}_0$ and $\mathbf{R}_0$. For the RRR models, we simulated $\boldsymbol{\beta} = c\mathbf{A}\mathbf{B}^T$ where $c \in \mathbb{R}^1$, $\mathbf{A} \in \mathbb{R}^{p \times d}$, $\mathbf{B} \in \mathbb{R}^{r \times d}$ were filled with uniform(0, 1) random numbers, then $\mathbf{A}$ and $\mathbf{B}$ were standardized to semi-orthogonal matrices. The covariance matrices $\boldsymbol{\Omega}$, $\boldsymbol{\Omega}_0$, $\boldsymbol{\Phi}$, $\boldsymbol{\Phi}_0$, $\boldsymbol{\Sigma}$ and $\boldsymbol{\Sigma}_X$ were generated as follows.

(E1) Isotropic envelope covariances. $\{\boldsymbol{\Omega}, \boldsymbol{\Omega}_0, \boldsymbol{\Phi}, \boldsymbol{\Phi}_0\} = \{5\mathbf{I}_{d_X}, \mathbf{I}_{p-d_X}, 5\mathbf{I}_{d_Y}, \mathbf{I}_{r-d_Y}\}$.

(E2) Randomized envelope covariances. All $\boldsymbol{\Omega}$, $\boldsymbol{\Omega}_0$, $\boldsymbol{\Phi}$ and $\boldsymbol{\Phi}_0$ were randomly generated as $\mathbf{M}\mathbf{M}^T$ where $\mathbf{M}$ had the same size as the corresponding covariances and were filled with uniform(0, 1) numbers.

(R1) Identity covariances. $\boldsymbol{\Sigma}_X = 5\mathbf{I}_{10}$ and $\boldsymbol{\Sigma} = 5\mathbf{I}_{15}$.

(R2) Randomized covariances. Both $\boldsymbol{\Sigma}_X$ and $\boldsymbol{\Sigma}$ were randomly generated as $\mathbf{M}\mathbf{M}^T$ in the similar way as in model (E2).

We simulated 200 data sets for each setting and used 100 for training and the others for testing. For each method, after obtaining the estimator $\hat{\boldsymbol{\beta}}$ from one training data set, we used the testing data set to compute the following quantity, which is an estimate of the variance of $\hat{\mathbf{Y}} - \mathbf{Y}_{\text{fit}} = (\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})^T \mathbf{X}_{\text{new}}$ for a new observation $\mathbf{X}_{\text{new}}$.

$$\boldsymbol{\alpha} = (\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})^T \hat{\boldsymbol{\Sigma}}_X (\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}) \in \mathbb{R}^{r \times r}, \quad (7.1)$$

where $\hat{\boldsymbol{\Sigma}}_X$ is the sample covariance matrix of $\mathbf{X}$ obtained from the testing set. Notice that $\text{cov}(\hat{\mathbf{Y}} - \mathbf{Y}) = \text{cov}(\hat{\mathbf{Y}} - \mathbf{Y}_{\text{fit}}) + \text{cov}(\boldsymbol{\epsilon})$ for all unbiased estimators, so we could make our comparison in the magnitudes of $\text{cov}(\hat{\mathbf{Y}} - \mathbf{Y}_{\text{fit}})$. Therefore the Frobenius norm $\|\boldsymbol{\alpha}\|_F$ was used as a measure of overall prediction error. We also considered the Frobenius norm

677 $\|\beta - \hat{\beta}\|_F$ as a measure of overall parameter estimation accuracy in the simulation studies
 678 in Supplement Section J.

679 We used five-fold cross-validation to find the best predictive dimensions for each of
 680 the methods. The results were summarized in Table 2. Our simulation results suggested
 681 that simultaneous envelopes with dimensions determined by cross-validation could be
 682 used in practice as a powerful prediction method. The simultaneous envelope methods
 683 had significantly better performance than all the other methods in most cases. For each of
 684 the other methods, it could perform as well as the simultaneous envelope in some cases
 685 but could have much worse performances in other cases.

686 In model (R2), where the true envelope dimensions were $(d_X, d_Y) = (p, r)$, the en-
 687 velope estimators with its true dimensions will be equivalent to OLS. Nevertheless, the
 688 $\mathbf{X}$ -envelope and the simultaneous envelope showed significant improvement over OLS
 689 in prediction by using the best predictive dimensions. This was probably due to the
 690 variance-bias tradeoff in finite samples. Model (R1) is both an envelope model and a
 691 RRR model with $d = d_X = d_Y$. It has the simplest structure, but the simultaneous en-
 692 velope estimators still produced significant improvement over OLS. RRR outperformed
 693 OLS but was still not as efficient as the simultaneous envelope estimators. Model (E2)
 694 produced immaterial covariances Φ_0 and Ω_0 that had much larger eigenvalues than the
 695 material parts Φ and Ω . The simultaneous envelope method efficiently discarded the im-
 696 material information and performed drastically better than other methods in models (E2).
 697 Moreover, CCA aimed for the most correlated pairs of $\hat{\mathbf{L}}^T \mathbf{X}$ and $\hat{\mathbf{R}}^T \mathbf{Y}$ and mistakenly
 698 found the immaterial pairs in model (E2).

699 To further study how the relative magnitude of immaterial variation over material
 700 variation can affect the performance of envelope methods, we repeated the above simu-
 701 lations in model (E2) for varying $d = d_X = d_Y$ from 1 to 8 with the same sample size
 702 $N = 200$. The true envelope dimensions were used instead of using cross-validation. We
 703 then plotted the averaged prediction error $\|\alpha\|_F$ over that of OLS in Figure 7.1. When the
 704 rank of β was smaller than 5, which is the half of its maximum possible rank, the immate-
 705 rial variance Φ_0 contained the larger eigenvalues of $\Sigma_{\mathbf{Y}|\mathbf{X}}$ hence the $\mathbf{Y}$ -envelope and the
 706 simultaneous envelope estimators had drastic gains in efficiency over OLS estimators.

(d_X, d_Y, N) or (d, N)		OLS	PLS	CCA	RRR	X-env.	Y-env.	S. env.	S. E. $\leq$
E1	(3, 3, 100)	3.15	1.55	1.14	2.46	1.50	2.53	0.63	0.06
	(3, 3, 400)	0.70	0.46	0.22	2.23	0.46	0.58	0.12	0.01
	(3, 3, 900)	0.30	0.14	0.08	2.28	0.14	0.25	0.05	0.006
E2	(2, 5, 400)	0.82	0.40	2.17	0.28	0.22	0.47	0.17	0.04
	(5, 3, 400)	1.00	0.51	1.76	0.36	0.55	0.65	0.27	0.07
	(7, 4, 400)	1.01	0.82	1.75	0.68	0.75	0.59	0.32	0.08
R1	(1, 200)	1.68	0.85	1.01	1.05	0.97	1.53	0.73	0.03
	(2, 200)	1.68	0.98	1.49	1.89	0.97	1.52	0.97	0.04
	(3, 200)	1.68	1.35	2.00	2.38	1.19	1.52	1.18	0.05
R2	(2, 200)	2.63	1.23	8.94	0.90	1.08	2.63	1.08	0.11*
	(3, 200)	2.86	1.54	26.6	1.52	1.08	2.79	1.08	0.15*
	(4, 200)	2.96	2.67	16.4	1.97	1.45	2.96	1.45	0.13*

Table 2: Prediction performances measured by $\|\alpha\|_F$, where * in the last column means the corresponding S.E. upper bound were computed without the S.E. of CCA. The best two methods for each setting were emphasized with boldface.

707 With increased rank of β , the immaterial variance parts were reduced but the simulta-
708 neous envelope reduction can still gain significantly over the OLS estimator. Even for
709 $d = 8$, the prediction error measurement $\|\alpha\|_F$ of the simultaneous envelope estimator
710 was only 70% of that of the OLS estimator.

7.2 Optimization with 1D manifold algorithm

712 The simulation results in Section 7.1 were all based on the 1D algorithm. The full Grass-
713 mann manifold optimization, which gives the MLE solution, would give indistinguishable
714 performances for small dimensions but could have trouble converging in larger dimen-
715 sions. To gain some insight on the computation cost of 1D algorithm versus the full
716 Grassmann manifold optimization, we simulated 100 data sets from model (E2) with
717 $N = 400$ and computed the averaged CPU time for estimating an envelope using the two
718 methods. When $d_X = d_Y$ increased from 2 to 3, the ratio of the averaged time costs on
719 full Grassmannian optimization over that of the 1D algorithm were increasing from 1.1
720 to 2.8 for estimating X-envelopes and from 1.4 to 3.8 for estimating Y-envelopes.

721 To examine how well the 1D manifold algorithm estimators can approximate the
722 MLE, we studied probability plots of the asymptotic χ^2 likelihood ratio test statistics.
723 Using model (E2), we simulated 100 dataset with sample size $N = 200$ with $p = r = 7$

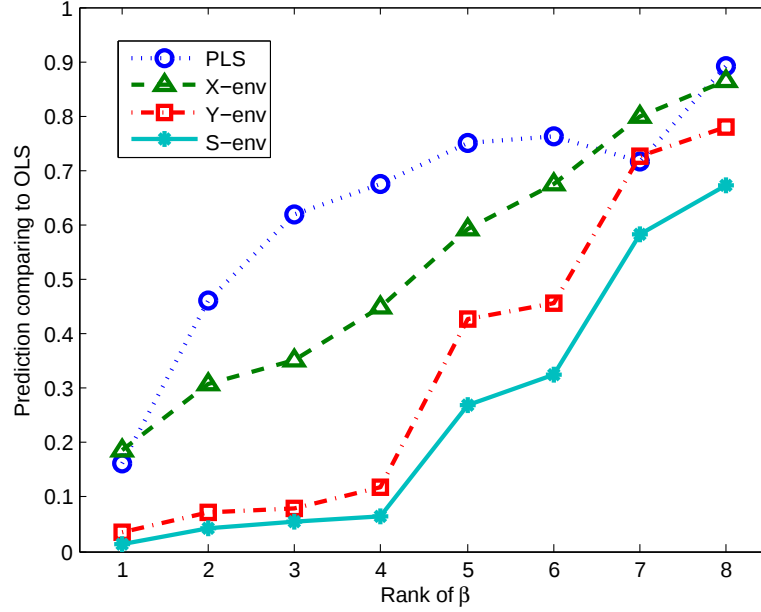


Figure 7.1: Prediction versus OLS for varying $d = d_X = d_Y$. Each point is the ratio of $\|\alpha\|_F$ for a specified estimator and for OLS estimator. The lines for CCA and RRR (not shown in the plot) were always larger than or close to 1.

and $d_X = d_Y = 3$. Consider the following hypothesis test given d_X and d_Y ,

$$H_0 : \mathcal{E}_{\Sigma_X}(\mathcal{L}) \oplus \mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R}) = \text{span}(\mathbf{L} \oplus \mathbf{R});$$

$$H_a : \mathcal{E}_{\Sigma_X}(\mathcal{L}) \oplus \mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R}) \neq \text{span}(\mathbf{L} \oplus \mathbf{R}).$$

The likelihood ratio test statistic for this hypothesis is $\Lambda = -2 \log(L(\mathbf{L}, \mathbf{R})) + 2 \log(L(\hat{\mathbf{L}}, \hat{\mathbf{R}}))$, where $\hat{\mathbf{L}}$ and $\hat{\mathbf{R}}$ are MLEs, is asymptotically χ^2 distributed. The degrees of freedom is the sum of the two Grassmann manifolds dimensions: $(p - d_X)d_X + (r - d_Y)d_Y = 24$. We computed the likelihood ratio test statistics using the 1D algorithm (4.9) to estimate $\hat{\mathbf{L}}$ and $\hat{\mathbf{R}}$. Figure 7.2 summarized the calculated sample quantiles of the one hundred statistics and the expected quantiles of the χ^2_{24} distribution. All the points fall near a straight line, indicating that the estimators from the 1D algorithm behave as MLEs.

7.3 Determining the dimensions (d_X, d_Y)

We simulated 100 datasets for each of the following settings, and performed the chi-square test for the rank of β with level 0.05, the likelihood ratio test with level 0.01, AIC model selection and BIC model selection. We used model (E2) in Section 7.1 with $p = 6$

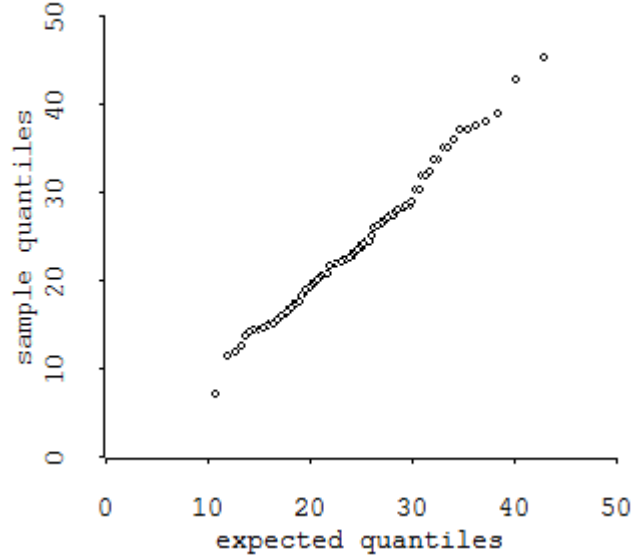


Figure 7.2: Probability plot based on the 1D algorithm

736 predictors and $r = 8$ responses in two settings: (I) $N = 100$ observations for each dataset
 737 with $\mathbf{X}$ - and $\mathbf{Y}$ -envelope dimensions $d_X = d_Y = 2$. (II) $N = 300$ observations for each
 738 dataset with $\mathbf{X}$ - and $\mathbf{Y}$ -envelope dimensions $d_X = 2$, $d_Y = 3$. Table 3 summarizes
 739 the simulation performances of various tests. Clearly, the χ^2 test quite often reveals the
 740 correct dimension of β and BIC has the best performance on determining the envelope
 741 dimensions.

742 We further studied the likelihood ratio test statistics obtained by the 1D algorithm
 743 and SIMPLS. We used $p = r = 7$, $d_X = d_Y = 3$ and simulated 100 dataset with
 744 $N = 300$ observations from model (E2) in Section 7.1. The likelihood ratio test was
 745 computed for the hypothesis $d_X = d_Y = 3$. The test statistic follows a χ^2 -distribution
 746 with $pr - d_X d_Y = 40$ degrees of freedom as shown in Section 6.2. The probability plot of
 747 expected quantiles versus sample quantiles from the 100 datasets are given in Figure 7.3.
 748 Since the test statistic depends on our estimators of $\mathbf{L}$ and $\mathbf{R}$, the probability plot should
 749 follow a straight line if the estimators behave as MLEs. From this plot, we see again
 750 that the behavior of the 1D algorithm is as expected for MLEs. However, the SIMPLS
 751 algorithm produced noticeable curvature in its probability plot.

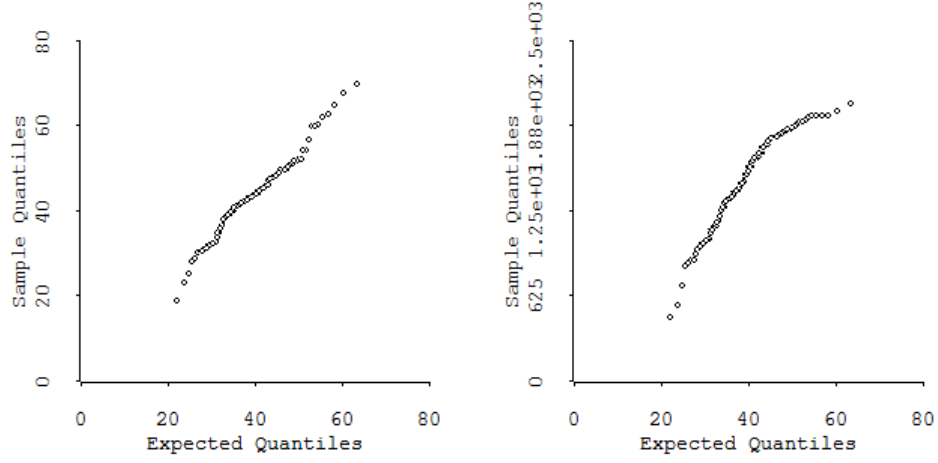


Figure 7.3: Chi-squared distribution probability plots. The 1D algorithm results the left plot, and the SIMPLS algorithm results the right plot. The reference distribution is the χ^2_{40} -distribution.

	correctness on	setting (I)	setting (II)
χ^2	d	89%	97%
LRT	d_X	98%	97%
	d_Y	100%	90%
	d_X and d_Y	98%	89%
AIC	d_X	75%	92%
	d_Y	67%	63%
	d_X and d_Y	56%	61%
BIC	d_X	98%	99%
	d_Y	100%	100%
	d_X and d_Y	98%	99%

Table 3: The numbers indicate how many times each true/estimated dimensions occurs among the 100 datasets.

8 Biscuit NIR spectroscopy data

This experiment involved varying the composition of biscuit dough. Two sets of dough pieces were measured, a calibration set and a prediction set. They were created and measured as two distinct sets, on separate occasions, and do not result from a random (or any other) split of a larger set. We have $r = 4$ response variables: fat, sucrose, flour and water percentages of biscuit dough. To keep the number of predictors less than the sample size of the training set, we used the spectral range 1200-1280nm with 4nm steps, which gave $p = 20$ predictors with $n_1 = 39$ training samples and $n_2 = 31$ testing samples. The prediction accuracy was based on the sum of squared error for predicting the given testing set.

The rank test of β indicated that $d = \text{rank}(\beta) = 2$ at the 0.05 significant level. Simultaneous envelope dimension $d_Y = 2$ was selected by AIC, BIC and LRT; $d_X = 10, 3, 8$ was selected by AIC, BIC and LRT respectively. Based on the prediction error, AIC chose the best dimension, while BIC made a bad choice, which was probably due to the small sample size of the training set.

The sum of squared prediction errors (prediction SSE) on the testing sets are summarized in the Figure 8.1. When the number of components was less than 8, the envelope and SIMPLS estimators had similar performances and were both inferior to OLS because the dimension was too small to cover all the material information. Starting at 8 components, the envelope estimator performed much better than the SIMPLS and OLS estimators. It converged to the OLS estimator at 13 components and stayed the same thereafter. However, SIMPLS performed much worse than OLS unless the number of component was larger than 16. To aid visualization, prediction errors from CCA and RRR were not plotted because they were more than twice of that of OLS. We also considered predictions based on X- or Y-envelope reduction alone. The prediction errors for the X-envelope estimator traced the errors from the simultaneous envelope estimator but were uniformly larger. The prediction errors for the Y-envelope estimators were indistinguishable to that of the OLS estimator for $d_Y > 1$ and were larger than that for $d_Y = 1$.

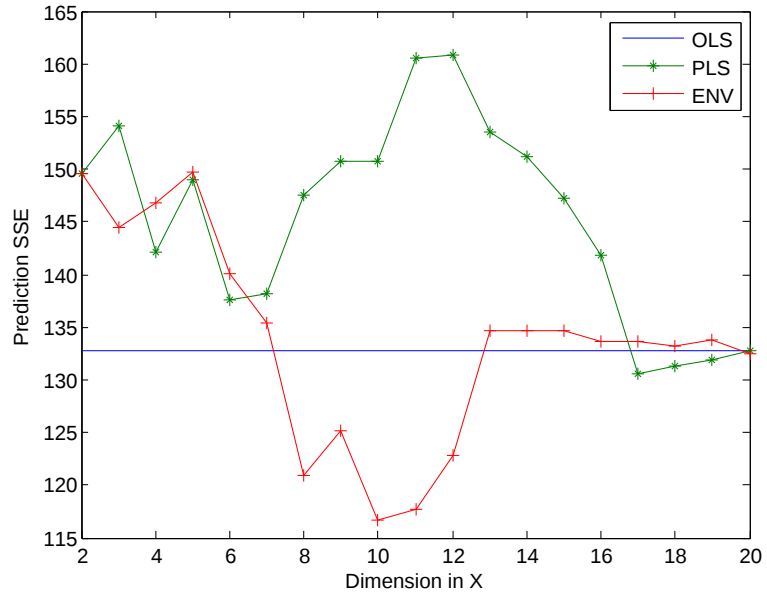


Figure 8.1: Prediction sum of square errors on the testing set. X-axis denote the numbers of components for PLS and simultaneous envelope X-dimension, d_X , where the Y-dimension of simultaneous envelope is fixed at $d_Y = 2$. To help visualization, the CCA performance is not included and the SSE for CCA are all greater than OLS no matter how many components to use.

9 Discussion

The partial envelope model (Su and Cook, 2011) is an extension of the envelope model in which some predictors are of special interest. Suppose $\mathbf{X}$ can be partitioned into $\mathbf{X}_1 \in \mathbb{R}^{p_1}$ and $\mathbf{X}_2 \in \mathbb{R}^{p_2}$, $p_1 + p_2 = p$, and the regression coefficient becomes $\boldsymbol{\beta} = (\boldsymbol{\beta}_1^T, \boldsymbol{\beta}_2^T)^T$ correspondingly. The partial envelope model focus on estimating $\boldsymbol{\beta}_1$ instead of $\boldsymbol{\beta}$, and can be written as,

$$\mathbf{Y} = \boldsymbol{\mu}_Y + \boldsymbol{\beta}_1^T(\mathbf{X}_1 - \boldsymbol{\mu}_{X_1}) + \boldsymbol{\beta}_2^T(\mathbf{X}_2 - \boldsymbol{\mu}_{X_2}) + \boldsymbol{\epsilon}, \quad (9.1)$$

where $\boldsymbol{\mu}_{X_1} = E(\mathbf{X}_1)$ and $\boldsymbol{\mu}_{X_2} = E(\mathbf{X}_2)$. The partial envelope in this model is $\mathcal{E}_{\Sigma_{Y|X}}(\text{span}(\boldsymbol{\beta}_1^T))$, which is contained in the envelope $\mathcal{E}_{\Sigma_{Y|X}}(\mathcal{R}) = \mathcal{E}_{\Sigma_{Y|X}}(\text{span}(\boldsymbol{\beta}^T))$. Since the partial envelope model focuses on a smaller subspace, there is potentially more efficiency gain than with the envelope model.

It is straightforward to extent partial envelope to simultaneous envelopes. Let $\Sigma_{X_1} \in \mathbb{R}^{p_1 \times p_1}$ be the covariance matrix of $\mathbf{X}_1$, and let $\Sigma_{X_2} \in \mathbb{R}^{p_2 \times p_2}$ be the covariance matrix of $\mathbf{X}_2$. We can then study the simultaneous partial envelope $\mathcal{E}_{\Sigma_{Y|X}}(\text{span}(\boldsymbol{\beta}_1^T)) \oplus \mathcal{E}_{\Sigma_{X_1}}(\text{span}(\boldsymbol{\beta}_1))$ to reduce $\mathbf{Y}$ and to reduce $\mathbf{X}_1$.

The simultaneous envelope model contains the envelope models in Cook et al. (2010) and Cook et al. (2013) as special cases. Extensions to the simultaneous envelope methods provide unified treatments for these methods. For instance, variable selection in multivariate linear regression has started to become an interesting and important research area. Multivariate linear methods that have well performances in classical problems were extended to high-dimensional situations. To deal with high-dimensional problems where $p, r > n$, Chun and Keles (2010) introduced sparse PLS and Chen and Huang (2012) proposed sparse RRR. We conjecture that an variable selection extension of the simultaneous envelope methodology will have nice applications in high-dimensional settings. Alternatively, we can first apply the SIMPLS algorithm, which is actually an envelope algorithm shown by Cook et al. (2013), to reduce the dimensionality and then apply the simultaneous envelope methods to gain efficiency in estimation.

806 Acknowledgment

807 We thank the Editor, the Associate Editor and three referees for their insightful comments
808 and constructive suggestions that have greatly improved the presentation of the paper.
809 Research for this article was supported in part by grant DMS-1007547 from the National
810 Science Foundation.

811 Supplementary Materials

812 **Proofs and Technical Details:** Detailed proofs for all Lemmas and Propositions are pro-
813 vided in the online supplement to this article. (PDF file)

814 References

- 815 [1] ADRAGNI, K., COOK, R.D. AND WU, S. (2012), GrassmannOptim: An
816 R Package for Grassmann Manifold Optimization, *Journal of Statistical*
817 *Software*, **50**, 1–18.
- 818 [2] BACH, F.R. AND JORDAN, M.I. (2005), A probabilistic interpretation of
819 canonical correlation analysis, *Technical Report*, **688**, 1–11.
- 820 [3] BURA, E. AND COOK, R.D. (2003), Rank estimation in reduced-rank re-
821 gression, *Journal of Multivariate Analysis*, **87**, 159–176.
- 822 [4] CHUN, H. AND KELES, S. (2010), Sparse partial least squares regression
823 for simultaneous dimension reduction and variable selection, *Journal of the*
824 *Royal Statistical Society: Series B*, **72**, 3–25.
- 825 [5] CHEN, L. AND HUANG, J.Z. (2012), Sparse reduced-rank regression for
826 simultaneous dimension reduction and variable selection, *Journal of the*
827 *American Statistical Association*, **107**, 1533–1545.
- 828 [6] CONWAY, J. (1990). A course in functional analysis. Second edition.
829 *Springer, New York.*

- 830 [7] COOK, R.D., HELLAND, I.S. AND SU, Z. (2013), Envelopes and partial
831 least squares regression. *To appear in JRSS-B*.
- 832 [8] COOK, R.D., LI, B. AND CHIAROMONTE, F. (2010). Envelope models
833 for parsimonious and efficient multivariate linear regression (with discus-
834 sion). *Statistica Sinica*, **20**, 927–1010.
- 835 [9] EDELMAN, A., TOMAS, A.A. AND SMITH, S.T. (1998), The geome-
836 try of algorithms with orthogonality constraints, *SIAM Journal of Matrix*
837 *Analysis and Applications*, **20**, 303–353.
- 838 [10] HENDERSON, H.V. AND SEARLE, S.R. (1979), Vec and Vech operators
839 for matrices, with some uses in Jacobians and multivariate statistics. *Cana-
840 dian Journal of Statistics*, **7**, 65–81.
- 841 [11] HOTELLING, H. (1936), Relations between two sets of variates.
842 *Biometrika*. **28**, 321–377.
- 843 [12] IZENMAN, A. J. (1975), Reduced-rank regression for the multivariate lin-
844 ear model. *Journal of Multivariate Analysis*. **5**, 248–264.
- 845 [13] JOLLIFFE, I. (2005), Principal Component Analysis. *Encyclopedia of*
846 *Statistics in Behavioral Science*.
- 847 [14] JOLLIFFE, I. (1986), Principal component analysis, *Springer Verlag, New*
848 *York*.
- 849 [15] DE JONG, S. (1993), SIMPLS: an alternative approach to partial least
850 squares regression. *Chemometr. Intell. Lab. Syst.*, **18**, 251–26.
- 851 [16] LEHMANN, E.L. AND CASELLA G. (1998) Theory of Point Estimation,
852 Second Edition *New York: Springer*.
- 853 [17] LI, B., WEN, S. AND ZHU, L.X. (2008) On a Projective Resampling
854 Method for Dimension Reduction With Multivariate Responses, *Journal*
855 *of the American Statistical Association*, **103**, 1177–1186.

- 856 [18] LI, K.C., ARAGON, Y., SHEDDEN, K. AND AGAN, C.T. (2003), Di-
857 mension reduction for multivariate response data, *Journal of the American*
858 *Statistical Association*, **98**, 99–109.
- 859 [19] REINSEL, G. C. AND VELU, R. P. (1998), Multivariate reduced-rank re-
860 gression: theory and applications. *Springer, New York*.
- 861 [20] SHAO, J. (1997), An asymptotic theory for linear model selection (with
862 discussion). *Statistica sinica*, **7**, 221–264.
- 863 [21] SHAPIRO, A. (1986), Asymptotic theory of overparameterized structural
864 models, *Journal of the American Statistical Association*, **81**, 142–149.
- 865 [22] SMALL, C.G., WANG, J. AND YANG, Z. (2000) Eliminating multiple root
866 problems in estimation. *Statistical Science*, **15**, 313–341.
- 867 [23] SU, Z. AND COOK, R.D. (2011). Partial envelopes for efficient estimation
868 in multivariate linear regression. *Biometrika*, **98**, 133–146.
- 869 [24] TYLER, D.E. (1981). Asymptotic inference for eigenvectors. *Annals of*
870 *Statistics*, **9**, 725–736.
- 871 [25] WOLD, H. (1966), Estimation of Principal Components and Related Mod-
872 els by Iterative Least Squares. *New York: Academic Press*.
- 873 [26] YANG, Y. (2005), Can the strengths of AIC and BIC be shared? A conflict
874 between model identification and regression estimation. *Biometrika*. **92**,
875 934–950.

876

877

878

879

880

Supplement: Proofs and Technical Details for “Simultaneous envelopes for multivariate linear regression”

881

R. Dennis Cook and Xin Zhang

882

883

884

885

886

The referred equations in the Supplement are labeled as (A1), (A2) and so on, whereas labels such as (1), (2), etc. refer to equations the main text. We also include here additional lemmas and corollaries (and when necessary, their proofs), sometimes within the proof of an assertion in the main text. In the following proofs, we use f and F as generic functions whose definitions change.

887

A Proof of Lemma 1 and Lemma 2

888

889

Lemma 1 and Lemma 2 follow directly from the discussions in Section 3 and the definition of envelopes.

890

B Proof of Proposition 2

891

892

893

Proof. Let $\mathbb{X} \in \mathbb{R}^{n \times p}$ and $\mathbb{Y} \in \mathbb{R}^{n \times r}$ be the centered data matrices from i.i.d. random vectors $\mathbf{X}$ and $\mathbf{Y}$. Then $\hat{\beta}_{\text{OLS}} = (\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \mathbb{Y}$ and $\hat{\beta}_{\text{L,R}} = \mathbf{L}(\mathbf{L}^T \mathbb{X}^T \mathbb{X} \mathbf{L})^{-1} \mathbf{L}^T \mathbb{X}^T \mathbb{Y} \mathbf{R} \mathbf{R}^T$.

For the variance of $\text{vec}(\hat{\beta}_{\text{OLS}}^T)$, we decompose it into two terms, conditioning on $\mathbb{X}$.

$$\begin{aligned}
 \text{var}(\text{vec}(\hat{\beta}_{\text{OLS}}^T)) &= \text{var} \left\{ \text{vec}(\mathbb{Y}^T \mathbb{X} (\mathbb{X}^T \mathbb{X})^{-1}) \right\} \\
 &= \text{var} \left\{ ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) \text{vec}(\mathbb{Y}^T) \right\} \\
 &= \text{var} \left\{ \mathbb{E} \{ ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) \text{vec}(\mathbb{Y}^T) | \mathbb{X} \} \right\} \\
 &\quad + \mathbb{E} \left\{ \text{var}((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) \text{vec}(\mathbb{Y}^T) | \mathbb{X} \right\} \\
 &\equiv VE + EV.
 \end{aligned}$$

894 The first term is

$$\begin{aligned}
VE &= \text{var} \left\{ E \left\{ ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) \text{vec}(\mathbb{Y}^T) | \mathbb{X} \right\} \right\} \\
&= \text{var} \left\{ ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) E \{ \text{vec}(\mathbb{Y}^T) | \mathbb{X} \} \right\} \\
&= \text{var} \left\{ ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) \text{vec}(\boldsymbol{\beta}^T \mathbb{X}^T) \right\} \\
&= \text{var} \left\{ \text{vec}(\boldsymbol{\beta}^T \mathbb{X}^T \mathbb{X} (\mathbb{X}^T \mathbb{X})^{-1}) \right\} = 0.
\end{aligned}$$

895 Then the variance is just the second term which is

$$\begin{aligned}
\text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\text{OLS}}^T)) &= EV = E \left\{ \text{var}((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) \text{vec}(\mathbb{Y}^T) | \mathbb{X} \right\} \\
&= E \left\{ ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) \text{var}(\text{vec}(\mathbb{Y}^T) | \mathbb{X}) ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r)^T \right\} \\
&= E \left\{ ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r) (\mathbf{I}_n \otimes \boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}}) ((\mathbb{X}^T \mathbb{X})^{-1} \mathbb{X}^T \otimes \mathbf{I}_r)^T \right\} \\
&= E \left\{ (\mathbb{X}^T \mathbb{X})^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}} \right\} \\
&= f_p^{-1} \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}},
\end{aligned}$$

896 where the last equality come from the fact that $\mathbb{X}^T \mathbb{X} \sim W_p(\boldsymbol{\Sigma}_{\mathbf{X}}, n - 1)$. The degrees
897 of freedom $n - 1$ are greater than $p - 1$. Then $(\mathbb{X}^T \mathbb{X})^{-1}$ follows an inverse Wishart
898 distribution $W_p^{-1}(\boldsymbol{\Sigma}_{\mathbf{X}}^{-1}, n - 1)$, and its mean equals $\boldsymbol{\Sigma}_{\mathbf{X}}^{-1}/(n - p - 2) = f_p^{-1} \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}$.

899 Similarly for $\hat{\boldsymbol{\beta}}_{\mathbf{L}, \mathbf{R}}$, we have the following expressions by noticing that $\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{L}, \mathbf{R}}^T) =$
900 $(\mathbf{L} \oplus \mathbf{R}) \text{vec}(\hat{\boldsymbol{\eta}}^T)$ and $\hat{\boldsymbol{\eta}}$ is the OLS estimator of $\mathbf{R}^T \mathbf{Y}$ on $\mathbf{L}^T \mathbf{X}$.

$$\begin{aligned}
\text{var}(\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{L}, \mathbf{R}}^T)) &= f_x^{-1} (\mathbf{L} \otimes \mathbf{R}) (\boldsymbol{\Sigma}_{\mathbf{L}^T \mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{R}^T \mathbf{Y} | \mathbf{L}^T \mathbf{X}}) (\mathbf{L}^T \otimes \mathbf{R}^T) \\
&= f_x^{-1} (\mathbf{L} \otimes \mathbf{R}) ((\mathbf{L}^T \boldsymbol{\Sigma}_{\mathbf{X}} \mathbf{L})^{-1} \otimes (\mathbf{R}^T \boldsymbol{\Sigma}_{\mathbf{Y} | \mathbf{X}} \mathbf{R})) (\mathbf{L}^T \otimes \mathbf{R}^T) \\
&= f_x^{-1} (\mathbf{L} (\mathbf{L}^T \boldsymbol{\Sigma}_{\mathbf{X}} \mathbf{L})^{-1} \mathbf{L}^T) \otimes (\mathbf{R} \mathbf{R}^T \boldsymbol{\Sigma}_{\mathbf{Y} | \mathbf{X}} \mathbf{R} \mathbf{R}^T) \\
&= f_x^{-1} (\mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T) \otimes (\mathbf{P}_{\mathbf{R}} \boldsymbol{\Sigma}_{\mathbf{Y} | \mathbf{X}} \mathbf{P}_{\mathbf{R}}).
\end{aligned}$$

901 Next, by noticing that

$$\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} = \mathbf{L} (\mathbf{L}^T \boldsymbol{\Sigma}_{\mathbf{X}} \mathbf{L})^{-1} \mathbf{L}^T + \mathbf{L}_0 (\mathbf{L}_0^T \boldsymbol{\Sigma}_{\mathbf{X}} \mathbf{L}_0)^{-1} \mathbf{L}_0^T, \quad (\text{B.1})$$

$$\boldsymbol{\Sigma}_{\mathbf{Y} | \mathbf{X}} = \mathbf{P}_{\mathbf{R}} \boldsymbol{\Sigma}_{\mathbf{Y} | \mathbf{X}} \mathbf{P}_{\mathbf{R}} + \mathbf{P}_{\mathbf{R}_0} \boldsymbol{\Sigma}_{\mathbf{Y} | \mathbf{X}} \mathbf{P}_{\mathbf{R}_0}, \quad (\text{B.2})$$

$$\boldsymbol{\Sigma}_{\mathbf{Y} | \mathbf{X}} = \mathbf{R} \boldsymbol{\Phi} \mathbf{R}^T + \mathbf{R}_0 \boldsymbol{\Phi}_0 \mathbf{R}_0^T, \quad (\text{B.3})$$

we have

$$\begin{aligned}
\text{var}(\text{vec}(\hat{\beta}_{\text{OLS}}^T)) &= f_p^{-1} \Sigma_{\mathbf{X}}^{-1} \otimes \Sigma_{\mathbf{Y}|\mathbf{X}} \\
&= f_x f_p^{-1} \text{var}(\text{vec}(\hat{\beta}_{\mathbf{L}, \mathbf{R}}^T)) + f_p^{-1} \mathbf{L} \Omega^{-1} \mathbf{L}^T \otimes \mathbf{R}_0 \Phi_0 \mathbf{R}_0^T \\
&\quad + f_p^{-1} \mathbf{L}_0 \Omega_0^{-1} \mathbf{L}_0^T \otimes \mathbf{R} \Phi \mathbf{R}^T + f_p^{-1} \mathbf{L}_0 \Omega_0^{-1} \mathbf{L}_0^T \otimes \mathbf{R}_0 \Phi_0 \mathbf{R}_0^T,
\end{aligned}$$

where $\Omega = \mathbf{L}^T \Sigma_{\mathbf{X}} \mathbf{L}$, $\Omega_0 = \mathbf{L}_0^T \Sigma_{\mathbf{X}} \mathbf{L}_0$, $\Phi_0 = \mathbf{R}_0^T \Sigma_{\mathbf{Y}|\mathbf{X}} \mathbf{R}_0$ and $\Phi = \mathbf{R}^T \Sigma_{\mathbf{Y}|\mathbf{X}} \mathbf{R}$. Then the result follows directly. $\square$

C Proof for Lemma 3

From Proposition 2.2 in Cook et al. (2010), we have $\mathcal{E}_{\mathbf{M}_i}(\mathcal{S}_i) = \sum_{j=1}^{q_i} \mathbf{P}_{ij} \mathcal{S}_i$, $i = 1, 2$, where $\mathbf{P}_{ij}$ is the projection onto the j -th eigenspace of $\mathbf{M}_i$. Since $\mathbf{M} = \mathbf{M}_1 \oplus \mathbf{M}_2$, then the eigen-projections of $\mathbf{M}$ will be $\mathbf{P}_{1j} \oplus \mathbf{0}_{p_2}$, $j = 1, \dots, q_1$, and $\mathbf{0}_{p_1} \oplus \mathbf{P}_{2k}$, $k = 1, \dots, q_2$. Therefore, applying Proposition 2.2 of Cook et al. (2010) again, we have

$$\begin{aligned}
\mathcal{E}_{\mathbf{M}}(\mathcal{S}) &= \left\{ \sum_{j=1}^{q_1} [(\mathbf{P}_{1j} \oplus \mathbf{0})(\mathcal{S}_1 \oplus \mathcal{S}_2)] \right\} + \left\{ \sum_{k=1}^{q_2} [(\mathbf{0} \oplus \mathbf{P}_{2k})(\mathcal{S}_1 \oplus \mathcal{S}_2)] \right\} \\
&= \{ \mathcal{E}_{\mathbf{M}_1}(\mathcal{S}_1) \oplus \mathbf{0} \} + \{ \mathbf{0} \oplus \mathcal{E}_{\mathbf{M}_2}(\mathcal{S}_2) \} \\
&= \mathcal{E}_{\mathbf{M}_1}(\mathcal{S}_1) \oplus \mathcal{E}_{\mathbf{M}_2}(\mathcal{S}_2).
\end{aligned}$$

D Proof for Lemma 4

D.1 The log-det term: $\log |\Sigma_{\mathbf{C}}|$

Apply the following orthogonal transformation to $\Sigma_{\mathbf{C}}$,

$$\begin{aligned}
\mathbf{O}^T \Sigma_{\mathbf{C}} \mathbf{O} &= \begin{pmatrix} (\mathbf{L} \ \mathbf{L}_0)^T & \mathbf{0} \\ \mathbf{0} & (\mathbf{R} \ \mathbf{R}_0)^T \end{pmatrix} \begin{pmatrix} \Sigma_{\mathbf{X}} & \Sigma_{\mathbf{XY}} \\ \Sigma_{\mathbf{XY}}^T & \Sigma_{\mathbf{Y}} \end{pmatrix} \begin{pmatrix} (\mathbf{L} \ \mathbf{L}_0) & \mathbf{0} \\ \mathbf{0} & (\mathbf{R} \ \mathbf{R}_0) \end{pmatrix} \\
&= \begin{pmatrix} \begin{pmatrix} \Omega & \mathbf{0} \\ \mathbf{0} & \Omega_0 \end{pmatrix} & \begin{pmatrix} \Omega \eta & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{pmatrix} \\ \begin{pmatrix} \eta^T \Omega & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{pmatrix} & \begin{pmatrix} \Phi + \eta^T \Omega \eta & \mathbf{0} \\ \mathbf{0} & \Phi_0 \end{pmatrix} \end{pmatrix}.
\end{aligned}$$

We can further transform it into a block diagonal matrix as in the following general case:

$$\begin{pmatrix} \mathbf{I} & \mathbf{0} \\ -\mathbf{B}^T \mathbf{A}^{-1} & \mathbf{I} \end{pmatrix} \begin{pmatrix} \mathbf{A} & \mathbf{B} \\ \mathbf{B}^T & \mathbf{C} \end{pmatrix} \begin{pmatrix} \mathbf{I} & -\mathbf{A}^{-1} \mathbf{B} \\ \mathbf{0} & \mathbf{I} \end{pmatrix} = \begin{pmatrix} \mathbf{A} & \mathbf{0} \\ \mathbf{0} & \mathbf{C} - \mathbf{B}^T \mathbf{A}^{-1} \mathbf{B} \end{pmatrix}. \quad (\text{D.1})$$

915 By taking the transformation matrix $\mathbf{T} = \begin{pmatrix} \mathbf{I} & \mathbf{T}_{12} \\ \mathbf{0} & \mathbf{I} \end{pmatrix}$ where $\mathbf{T}_{12} = \begin{pmatrix} -\boldsymbol{\eta} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{pmatrix}$, we
 916 have the block diagonal matrix

$$\mathbf{T}^T \mathbf{O}^T \boldsymbol{\Sigma}_C \mathbf{O} \mathbf{T} = \text{diag}\{\boldsymbol{\Omega}, \boldsymbol{\Omega}_0, \boldsymbol{\Phi}, \boldsymbol{\Phi}_0\} \equiv \boldsymbol{\Sigma}_{\text{diag}}. \quad (\text{D.2})$$

917 Also, because of the determinants of $\mathbf{T}$ and $\mathbf{O}$ are both 1, $|\boldsymbol{\Sigma}_C| = |\boldsymbol{\Sigma}_{\text{diag}}|$, and

$$\log |\boldsymbol{\Sigma}_C| = \log |\boldsymbol{\Omega}| + \log |\boldsymbol{\Omega}_0| + \log |\boldsymbol{\Phi}| + \log |\boldsymbol{\Phi}_0|. \quad (\text{D.3})$$

918 **D.2 The trace term: $\text{trace}(\mathbf{S}_C \boldsymbol{\Sigma}_C^{-1})$**

$$\text{trace}(\mathbf{S}_C \boldsymbol{\Sigma}_C^{-1}) = \text{trace}(\mathbf{S}_C \mathbf{O} \mathbf{T} \boldsymbol{\Sigma}_{\text{diag}}^{-1} \mathbf{T}^T \mathbf{O}^T) = \text{trace}(\mathbf{T}^T \mathbf{O}^T \mathbf{S}_C \mathbf{O} \mathbf{T} \boldsymbol{\Sigma}_{\text{diag}}^{-1}).$$

919 Write $\mathbf{O} \equiv \begin{pmatrix} \mathbf{O}_g & \mathbf{0} \\ \mathbf{0} & \mathbf{O}_h \end{pmatrix}$, where $\mathbf{O}_g = (\mathbf{L}, \mathbf{L}_0)$ and $\mathbf{O}_h = (\mathbf{R}, \mathbf{R}_0)$. Next, we compute
 920 the matrix

$$\mathbf{T}^T \mathbf{O}^T \mathbf{S}_C \mathbf{O} \mathbf{T} = \mathbf{T}^T \begin{pmatrix} \mathbf{O}_g^T \mathbf{S}_X \mathbf{O}_g & \mathbf{O}_g^T \mathbf{S}_{XY} \mathbf{O}_h \\ \mathbf{O}_h^T \mathbf{S}_{XY}^T \mathbf{O}_g & \mathbf{O}_h^T \mathbf{S}_Y \mathbf{O}_h \end{pmatrix} \mathbf{T} = \begin{pmatrix} \mathbf{O}_g^T \mathbf{S}_X \mathbf{O}_g & * \\ * & \mathbf{M}_1 \end{pmatrix}.$$

921 Since we only need the trace of $\mathbf{T}^T \mathbf{O}^T \mathbf{S}_C \mathbf{O} \mathbf{T} \boldsymbol{\Sigma}_{\text{diag}}^{-1}$, and $\boldsymbol{\Sigma}_{\text{diag}}$ is block diagonal, so
 922 the off-diagonal blocks of $\mathbf{T}^T \mathbf{O}^T \mathbf{S}_C \mathbf{O} \mathbf{T}$ are not needed. And the matrix $\mathbf{M}_1$ can be
 923 computed to be

$$\mathbf{M}_1 = \mathbf{O}_h^T \mathbf{S}_Y \mathbf{O}_h + (\mathbf{O}_g \mathbf{T}_{12})^T \mathbf{S}_X (\mathbf{O}_g \mathbf{T}_{12}) + (\mathbf{O}_g \mathbf{T}_{12})^T \mathbf{S}_{XY} \mathbf{O}_h + [(\mathbf{O}_g \mathbf{T}_{12})^T \mathbf{S}_{XY} \mathbf{O}_h]^T.$$

924 Therefore,

$$\begin{aligned} \text{trace}(\mathbf{S}_C \boldsymbol{\Sigma}_C^{-1}) &= \text{trace}(\mathbf{T}^T \mathbf{O}^T \mathbf{S}_C \mathbf{O} \mathbf{T} \boldsymbol{\Sigma}_{\text{diag}}^{-1}) \\ &= \text{trace}(\mathbf{O}_g^T \mathbf{S}_X \mathbf{O}_g \text{diag}(\boldsymbol{\Omega}^{-1}, \boldsymbol{\Omega}_0^{-1})) + \text{trace}(\mathbf{M}_1 \text{diag}(\boldsymbol{\Phi}^{-1}, \boldsymbol{\Phi}_0^{-1})) \\ &\equiv \text{trace}_1 + \text{trace}_2. \end{aligned}$$

925 The first trace term,

$$\begin{aligned} \text{trace}_1 &= \text{trace}(\mathbf{O}_g^T \mathbf{S}_X \mathbf{O}_g \text{diag}(\boldsymbol{\Omega}^{-1}, \boldsymbol{\Omega}_0^{-1})) \\ &= \text{trace} \left\{ \begin{pmatrix} \mathbf{L}^T \\ \mathbf{L}_0^T \end{pmatrix} \mathbf{S}_X \begin{pmatrix} \mathbf{L} & \mathbf{L}_0 \end{pmatrix} \begin{pmatrix} \boldsymbol{\Omega}^{-1} & \mathbf{0} \\ \mathbf{0} & \boldsymbol{\Omega}_0^{-1} \end{pmatrix} \right\} \\ &= \text{trace}\{\mathbf{L}^T \mathbf{S}_X \mathbf{L} \boldsymbol{\Omega}^{-1}\} + \text{trace}\{\mathbf{L}_0^T \mathbf{S}_X \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1}\} = \text{trace}\{\mathbf{S}_X \boldsymbol{\Sigma}_X^{-1}\}. \end{aligned}$$

926 Similarly for the second trace term,

$$\begin{aligned}
\text{trace}_2 &= \text{trace}(\mathbf{M}_1 \text{diag}(\Phi^{-1}, \Phi_0^{-1})) \\
&= \text{trace}\{\mathbf{R}^T \mathbf{S}_Y \mathbf{R} \Phi^{-1}\} + \text{trace}\{\mathbf{R}_0^T \mathbf{S}_Y \mathbf{R}_0 \Phi_0^{-1}\} \\
&\quad + \text{trace}\{(\mathbf{L}\boldsymbol{\eta})^T \mathbf{S}_X (\mathbf{L}\boldsymbol{\eta}) \Phi^{-1}\} - 2 \times \text{trace}\{(\mathbf{L}\boldsymbol{\eta})^T \mathbf{S}_{XY} \mathbf{R} \Phi^{-1}\}.
\end{aligned}$$

927 **D.3 Partial derivatives of the objective function**

928 The objective function $F(\Sigma_C)$ now becomes an objective function of all the parameters

$$\begin{aligned}
F(\mathbf{L}, \mathbf{R}, \boldsymbol{\eta}, \Phi, \Phi_0, \Omega, \Omega_0) &= \log |\Omega| + \log |\Omega_0| + \log |\Phi| + \log |\Phi_0| \\
&\quad + \text{trace}\{\mathbf{L}^T \mathbf{S}_X \mathbf{L} \Omega^{-1}\} + \text{trace}\{\mathbf{L}_0^T \mathbf{S}_X \mathbf{L}_0 \Omega_0^{-1}\} \\
&\quad + \text{trace}\{\mathbf{R}^T \mathbf{S}_Y \mathbf{R} \Phi^{-1}\} + \text{trace}\{\mathbf{R}_0^T \mathbf{S}_Y \mathbf{R}_0 \Phi_0^{-1}\} \\
&\quad + \text{trace}\{(\mathbf{L}\boldsymbol{\eta})^T \mathbf{S}_X (\mathbf{L}\boldsymbol{\eta}) \Phi^{-1}\} - 2 \text{trace}\{(\mathbf{L}\boldsymbol{\eta})^T \mathbf{S}_{XY} \mathbf{R} \Phi^{-1}\}.
\end{aligned}$$

929 We will apply the following result repeatedly in this section.

$$\arg \min_{\mathbf{A}} \{\log |\mathbf{A}| + \text{trace}(\mathbf{A}^{-1} \mathbf{B})\} = \mathbf{B}, \quad (\text{D.4})$$

930 where $\mathbf{B}$ is positive definite symmetric matrix and the minimization is over all positive
931 definite symmetric matrices.

932 **D.3.1 Minimization over Ω, Ω_0, Φ and Φ_0**

933 Applying (D.4), we have

$$\hat{\Omega} = \hat{\mathbf{L}}^T \mathbf{S}_X \hat{\mathbf{L}}; \quad (\text{D.5})$$

$$\hat{\Omega}_0 = \hat{\mathbf{L}}_0^T \mathbf{S}_X \hat{\mathbf{L}}_0; \quad (\text{D.6})$$

$$\hat{\Phi}_0 = \hat{\mathbf{R}}_0^T \mathbf{S}_Y \hat{\mathbf{R}}_0. \quad (\text{D.7})$$

936 After we obtain the minimizer of $\hat{\boldsymbol{\eta}}$ and substitute it into the objective function, we can
937 also use (D.4) to get the following.

$$\hat{\Phi} = \hat{\mathbf{R}}^T \left(\mathbf{S}_Y - \mathbf{S}_{XY}^T \hat{\mathbf{L}} \left(\hat{\mathbf{L}}^T \mathbf{S}_X \hat{\mathbf{L}} \right)^{-1} \hat{\mathbf{L}}^T \mathbf{S}_{XY} \right) \hat{\mathbf{R}}. \quad (\text{D.8})$$

938 D.3.2 Partial derivative w.r.t. η

$$\begin{aligned}
\frac{\partial}{\partial \eta} F(\mathbf{L}, \mathbf{R}, \eta, \Phi, \Phi_0, \Omega, \Omega_0) &= \frac{\partial}{\partial \eta} (\text{trace}\{(\mathbf{L}\eta)^T \mathbf{S}_X (\mathbf{L}\eta) \Phi^{-1}\}) \\
&\quad - 2 \frac{\partial}{\partial \eta} (\text{trace}\{(\mathbf{L}\eta)^T \mathbf{S}_{XY} \mathbf{R} \Phi^{-1}\}) \\
&= 2\mathbf{L}^T \mathbf{S}_X \mathbf{L} \eta \Phi^{-1} - 2\mathbf{L}^T \mathbf{S}_{XY} \mathbf{R} \Phi^{-1}.
\end{aligned}$$

939 Set it to 0, we have

$$\hat{\eta} = (\hat{\mathbf{L}}^T \mathbf{S}_X \hat{\mathbf{L}})^{-1} (\hat{\mathbf{L}}^T \mathbf{S}_{XY} \hat{\mathbf{R}}). \quad (\text{D.9})$$

940 Substituting (D.5), (D.6), (D.7) and (D.9) into the full objective function, we have the
941 following partially optimized objective function

$$\begin{aligned}
F(\mathbf{L}, \mathbf{R}, \Phi) &= \log |\mathbf{L}^T \mathbf{S}_X \mathbf{L}| + \log |\mathbf{L}_0^T \mathbf{S}_X \mathbf{L}_0| + \log |\Phi| + \log |\mathbf{R}_0^T \mathbf{S}_Y \mathbf{R}_0| \\
&\quad + \text{trace}\{\mathbf{I}_l\} + \text{trace}\{\mathbf{I}_{p-l}\} + \text{trace}\{\mathbf{R}^T \mathbf{S}_Y \mathbf{R} \Phi^{-1}\} + \text{trace}\{\mathbf{I}_{r-m}\} \\
&\quad + \text{trace}\{(\mathbf{L}^T \mathbf{S}_{XY} \mathbf{R})^T (\mathbf{L}^T \mathbf{S}_X \mathbf{L})^{-1} (\mathbf{L}^T \mathbf{S}_{XY} \mathbf{R}) \Phi^{-1}\} \\
&\quad - 2 \times \text{trace}\{(\mathbf{L}^T \mathbf{S}_{XY} \mathbf{R})^T (\mathbf{L}^T \mathbf{S}_X \mathbf{L})^{-1} (\mathbf{L}^T \mathbf{S}_{XY} \mathbf{R}) \Phi^{-1}\}.
\end{aligned}$$

942 If we ignore the terms that do not change with Φ , the rest part of the objective function
943 becomes

$$F(\Phi) = \log |\Phi| + \text{trace}\{\mathbf{R}^T \mathbf{S}_Y \mathbf{R} \Phi^{-1}\} - \text{trace}\{(\mathbf{L}^T \mathbf{S}_{XY} \mathbf{R})^T (\mathbf{L}^T \mathbf{S}_X \mathbf{L})^{-1} (\mathbf{L}^T \mathbf{S}_{XY} \mathbf{R}) \Phi^{-1}\},$$

944 which leads us to (D.8)

945 D.4 MLE

946 The MLE for $\hat{\mathbf{L}}$ and $\hat{\mathbf{R}}$ are obtained by minimizing over the Grassmann manifolds $\mathcal{G}_{(d_X, p)}$
947 and $\mathcal{G}_{(d_Y, r)}$, $(\hat{\mathbf{L}}, \hat{\mathbf{R}}) = \arg \min_{(\mathbf{L}, \mathbf{R})} \{F(\mathbf{L}, \mathbf{R})\}$, where the objective function $F(\mathbf{L}, \mathbf{R})$ is

$$\begin{aligned}
F(\mathbf{L}, \mathbf{R}) &= \log |\mathbf{L}^T \mathbf{S}_X \mathbf{L}| + \log |\mathbf{L}_0^T \mathbf{S}_X \mathbf{L}_0| + \log |\mathbf{R}_0^T \mathbf{S}_Y \mathbf{R}_0| \\
&\quad + \log |\mathbf{R}^T (\mathbf{S}_Y - \mathbf{S}_{XY}^T \mathbf{L} (\mathbf{L}^T \mathbf{S}_X \mathbf{L})^{-1} \mathbf{L}^T \mathbf{S}_{XY}) \mathbf{R}| \\
&= \log \begin{vmatrix} \mathbf{L}^T \mathbf{S}_X \mathbf{L} & \mathbf{L}^T \mathbf{S}_{XY} \mathbf{R} \\ \mathbf{R}^T \mathbf{S}_{XY}^T \mathbf{L} & \mathbf{R}^T \mathbf{S}_Y \mathbf{R} \end{vmatrix} + \log |\mathbf{L}_0^T \mathbf{S}_X \mathbf{L}_0| + \log |\mathbf{R}_0^T \mathbf{S}_Y \mathbf{R}_0| \\
&= \log \left| (\mathbf{L}^T \oplus \mathbf{R}^T) \begin{pmatrix} \mathbf{S}_X & \mathbf{S}_{XY} \\ \mathbf{S}_{XY}^T & \mathbf{S}_Y \end{pmatrix} (\mathbf{L} \oplus \mathbf{R}) \right| \\
&\quad + \log |(\mathbf{L}^T \oplus \mathbf{R}^T) (\mathbf{S}_X^{-1} \oplus \mathbf{S}_Y^{-1}) (\mathbf{L} \oplus \mathbf{R})| \equiv F(\mathbf{L} \oplus \mathbf{R}).
\end{aligned}$$

948 Then the MLEs for $\eta, \Omega, \Omega_0, \Phi, \Phi_0$ are as given in the lemma.

949 E Proof for Lemma 5

950 From our set-up, we know that $\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0 > 0$ thus $\mathcal{E}_{\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0}(\mathbf{B}_0^T \mathcal{S})$ exists. Let Γ be a basis
 951 of $\mathcal{E}_{\mathbf{M}}(\mathcal{S})$, and (Γ, Γ_0) be a orthogonal basis of $\mathbb{R}^p$, then $\mathbf{M} = \Gamma \Omega \Gamma^T + \Gamma_0 \Omega_0 \Gamma_0^T$ and
 952 $\mathcal{S} \subseteq \text{span}(\Gamma)$ for some $\Omega > 0$ and $\Omega_0 > 0$. Therefore,

$$\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0 = (\mathbf{B}_0^T \Gamma) \Omega (\mathbf{B}_0^T \Gamma)^T + (\mathbf{B}_0^T \Gamma_0) \Omega_0 (\mathbf{B}_0^T \Gamma_0)^T \quad (\text{E.1})$$

$$\mathbf{B}_0^T \mathcal{S} \subseteq \text{span}(\mathbf{B}_0^T \Gamma), \quad (\text{E.2})$$

953 where $\text{span}(\mathbf{B}_0^T \Gamma)$ is the orthogonal compliment of $\text{span}(\mathbf{B}_0^T \Gamma_0)$ in $\mathbb{R}^{p-q}$ since $\text{span}(\mathbf{B}) \subseteq$
 954 $\text{span}(\Gamma)$. Then we see

$$\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0 = \mathbf{P}_{\mathbf{B}_0^T \Gamma} \mathbf{B}_0^T \mathbf{M} \mathbf{B}_0 \mathbf{P}_{\mathbf{B}_0^T \Gamma} + \mathbf{Q}_{\mathbf{B}_0^T \Gamma} \mathbf{B}_0^T \mathbf{M} \mathbf{B}_0 \mathbf{Q}_{\mathbf{B}_0^T \Gamma},$$

955 which means $\text{span}(\mathbf{B}_0^T \Gamma)$ is a reducing subspace of $\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0$ which also contains $\mathbf{B}_0^T \mathcal{S}$
 956 by (E.2). By definition, we know $\mathcal{E}_{\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0}(\mathbf{B}_0^T \mathcal{S})$ is the smallest reducing subspace of
 957 $\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0$ that contains $\mathbf{B}_0^T \mathcal{S}$. Hence, $\mathcal{E}_{\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0}(\mathbf{B}_0^T \mathcal{S}) \subseteq \text{span}(\mathbf{B}_0^T \Gamma)$. Therefore, $\mathbf{v} \in$
 958 $\mathcal{E}_{\mathbf{B}_0^T \mathbf{M} \mathbf{B}_0}(\mathbf{B}_0^T \mathcal{S})$ implies $\mathbf{B}_0 \mathbf{v} \in \mathbf{B}_0 \text{span}(\mathbf{B}_0^T \Gamma) \subseteq \text{span}(\Gamma) = \mathcal{E}_{\mathbf{M}}(\mathcal{S})$.

959 F Proof for Proposition 3

960 Let $\mathbf{L} \in \mathbb{R}^{p \times d_X}$ be a semi-orthogonal basis of $\mathcal{E}_{\Sigma_X}(\mathcal{L})$, then we have $\Sigma_X = \mathbf{L} \Omega \mathbf{L}^T +$
 961 $\mathbf{L}_0 \Omega_0 \mathbf{L}_0^T$, $\Sigma_{\mathbf{X}\mathbf{Y}} \Sigma_{\mathbf{Y}}^{-1/2} = \mathbf{L} \boldsymbol{\xi}$ and $\Sigma_{\mathbf{X}|\mathbf{Y}} = \mathbf{L}(\Omega - \boldsymbol{\xi} \boldsymbol{\xi}^T) \mathbf{L}^T + \mathbf{L}_0 \Omega_0 \mathbf{L}_0^T$, where $\boldsymbol{\xi}$ is $d_X \times r$
 962 matrix with rank d . We also assume that $\Omega > 0$, $\Omega_0 > 0$ and $\Omega - \boldsymbol{\xi} \boldsymbol{\xi}^T > 0$. In this
 963 section, we first prove that the population 1D algorithm is estimating the envelope, then
 964 we show the $\sqrt{n}$ -consistency of the sample algorithm as stated in Proposition 3.

965 F.1 The first direction ℓ_1 in the population

966 We first solve the optimization for the first direction $\ell_1 = \arg \min_{\ell \in \mathbb{R}^p} f_0(\ell)$ subject to
 967 $\ell^T \ell = 1$, where $f_0(\ell) = \log |\ell^T \Sigma_{\mathbf{X}|\mathbf{Y}}| + \log |\ell^T \Sigma_{\mathbf{X}}^{-1}|$. We want to show that $\ell_1 \in$
 968 $\text{span}(\mathbf{L})$.

969 Let $\ell = \mathbf{L} \mathbf{h} + \mathbf{L}_0 \mathbf{h}_0$ for some $\mathbf{h} \in \mathbb{R}^{d_X}$ and $\mathbf{h}_0 \in \mathbb{R}^{(p-d_X)}$. Consider the optimization
 970 problem as unconstrained problem, $\ell_1 = \arg \min_{\ell \in \mathbb{R}^p} \{f_0(\ell) - 2 \log |\ell^T \ell|\}$. Then we

will have the same solution as the original problem up to an arbitrary scaling constant.

Next, we plug-in these expressions for ℓ , $\Sigma_{\mathbf{X}}$ and $\Sigma_{\mathbf{X}|\mathbf{Y}}$,

$$\begin{aligned} f_0(\ell) - 2 \log |\ell^T \ell| &= \log\{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h} + \mathbf{h}_0^T\Omega_0\mathbf{h}_0\} + \log\{\mathbf{h}^T\Omega^{-1}\mathbf{h} + \mathbf{h}_0^T\Omega_0^{-1}\mathbf{h}_0\} \\ &\quad - 2 \log\{\mathbf{h}^T\mathbf{h} + \mathbf{h}_0^T\mathbf{h}_0\} \\ &\equiv f(\mathbf{h}, \mathbf{h}_0) \end{aligned}$$

Taking partial derivative with respect to $\mathbf{h}_0$ and $\mathbf{h}$, we have

$$\begin{aligned} \frac{\partial}{\partial \mathbf{h}_0} f(\mathbf{h}, \mathbf{h}_0) &= \frac{2\Omega_0\mathbf{h}_0}{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h} + \mathbf{h}_0^T\Omega_0\mathbf{h}_0} + \frac{2\Omega_0^{-1}\mathbf{h}_0}{\mathbf{h}^T\Omega^{-1}\mathbf{h} + \mathbf{h}_0^T\Omega_0^{-1}\mathbf{h}_0} - \frac{4\mathbf{h}_0}{\mathbf{h}^T\mathbf{h} + \mathbf{h}_0^T\mathbf{h}_0} \\ \frac{\partial}{\partial \mathbf{h}} f(\mathbf{h}, \mathbf{h}_0) &= \frac{2(\Omega - \xi\xi^T)\mathbf{h}}{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h} + \mathbf{h}_0^T\Omega_0\mathbf{h}_0} + \frac{2\Omega^{-1}\mathbf{h}}{\mathbf{h}^T\Omega^{-1}\mathbf{h} + \mathbf{h}_0^T\Omega_0^{-1}\mathbf{h}_0} - \frac{4\mathbf{h}}{\mathbf{h}^T\mathbf{h} + \mathbf{h}_0^T\mathbf{h}_0} \end{aligned}$$

To get local minimums we need to set $\frac{\partial}{\partial \mathbf{h}_0} f(\mathbf{h}, \mathbf{h}_0) = 0$ and $\frac{\partial}{\partial \mathbf{h}} f(\mathbf{h}, \mathbf{h}_0) = 0$ which become the following.

$$\begin{aligned} \left\{ \frac{2\Omega_0}{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h} + \mathbf{h}_0^T\Omega_0\mathbf{h}_0} + \frac{2\Omega_0^{-1}}{\mathbf{h}^T\Omega^{-1}\mathbf{h} + \mathbf{h}_0^T\Omega_0^{-1}\mathbf{h}_0} \right\} \mathbf{h}_0 &= \left\{ \frac{4}{\mathbf{h}^T\mathbf{h} + \mathbf{h}_0^T\mathbf{h}_0} \right\} \mathbf{h}_0 \\ \left\{ \frac{2(\Omega - \xi\xi^T)}{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h} + \mathbf{h}_0^T\Omega_0\mathbf{h}_0} + \frac{2\Omega^{-1}}{\mathbf{h}^T\Omega^{-1}\mathbf{h} + \mathbf{h}_0^T\Omega_0^{-1}\mathbf{h}_0} \right\} \mathbf{h} &= \left\{ \frac{4}{\mathbf{h}^T\mathbf{h} + \mathbf{h}_0^T\mathbf{h}_0} \right\} \mathbf{h} \end{aligned}$$

Since the matrices involved in the above equations are all positive definite: $\Omega^{-1} > 0$, $\Omega - \xi\xi^T > 0$, $\Omega_0 > 0$ and $\Omega_0^{-1} > 0$. The solutions to the above two equations can only be eigenvectors of the matrices

$$\mathbf{A}_0 = \frac{2\Omega_0}{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h} + \mathbf{h}_0^T\Omega_0\mathbf{h}_0} + \frac{2\Omega_0^{-1}}{\mathbf{h}^T\Omega^{-1}\mathbf{h} + \mathbf{h}_0^T\Omega_0^{-1}\mathbf{h}_0} \quad (\text{F.1})$$

$$\mathbf{A} = \frac{2(\Omega - \xi\xi^T)}{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h} + \mathbf{h}_0^T\Omega_0\mathbf{h}_0} + \frac{2\Omega^{-1}}{\mathbf{h}^T\Omega^{-1}\mathbf{h} + \mathbf{h}_0^T\Omega_0^{-1}\mathbf{h}_0}. \quad (\text{F.2})$$

The eigenvectors of $\mathbf{A}_0$ are the same as those of Ω_0 . Hence, we have $\mathbf{h}_0$ equals 0 or any eigenvector of Ω_0 (and of course Ω_0^{-1}). Therefore, the minimum value of $f(\mathbf{h}, \mathbf{h}_0)$ has to be obtained by 0 or some eigenvector of Ω_0 (since $\mathbf{h}_0 = \infty$ can be easily eliminated). If $\mathbf{h}_0 = 0$ then our conclusion follows.

Let's assume $\mathbf{h}_0 \neq 0$ and $\Omega_0\mathbf{h}_0 = \lambda_k\mathbf{h}_0$.

Then,

$$\begin{aligned} f(\mathbf{h}, \mathbf{h}_0) &= \log\left\{ \frac{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h} + \lambda_k\mathbf{h}_0^T\mathbf{h}_0}{\mathbf{h}^T\mathbf{h} + \mathbf{h}_0^T\mathbf{h}_0} \right\} + \log\left\{ \frac{\mathbf{h}^T\Omega^{-1}\mathbf{h} + \frac{1}{\lambda_k}\mathbf{h}_0^T\mathbf{h}_0}{\mathbf{h}^T\mathbf{h} + \mathbf{h}_0^T\mathbf{h}_0} \right\} \\ &= \log\left\{ \frac{\mathbf{h}^T(\Omega - \xi\xi^T)\mathbf{h}}{\mathbf{h}^T\mathbf{h}} W_h + \lambda_k(1 - W_h) \right\} + \log\left\{ \frac{\mathbf{h}^T\Omega^{-1}\mathbf{h}}{\mathbf{h}^T\mathbf{h}} W_h + \frac{1}{\lambda_k}(1 - W_h) \right\} \end{aligned}$$

985 where $W_h = \frac{\mathbf{h}^T \mathbf{h}}{\mathbf{h}^T \mathbf{h} + \mathbf{h}_0^T \mathbf{h}_0}$ is the weight between 0 and 1. Because $\log(\cdot)$ is concave, we have

986 $\log(aW_h + b(1 - W_h)) \geq W_h \log(a) + (1 - W_h) \log(b)$, hence,

$$\begin{aligned}
f(\mathbf{h}, \mathbf{h}_0) &\geq W_h \left\{ \log \frac{\mathbf{h}^T (\mathbf{\Omega} - \boldsymbol{\xi} \boldsymbol{\xi}^T) \mathbf{h}}{\mathbf{h}^T \mathbf{h}} + \log \frac{\mathbf{h}^T \mathbf{\Omega}^{-1} \mathbf{h}}{\mathbf{h}^T \mathbf{h}} \right\} + (1 - W_h) \left\{ \log(\lambda_k) + \log\left(\frac{1}{\lambda_k}\right) \right\} \\
&= W_h \left\{ \log \frac{\mathbf{h}^T (\mathbf{\Omega} - \boldsymbol{\xi} \boldsymbol{\xi}^T) \mathbf{h}}{\mathbf{h}^T \mathbf{h}} + \log \frac{\mathbf{h}^T \mathbf{\Omega}^{-1} \mathbf{h}}{\mathbf{h}^T \mathbf{h}} \right\} \\
&\geq W_h \cdot \min_{\mathbf{h} \in \mathbb{R}^{d_X}} \left\{ \log \frac{\mathbf{h}^T (\mathbf{\Omega} - \boldsymbol{\xi} \boldsymbol{\xi}^T) \mathbf{h}}{\mathbf{h}^T \mathbf{h}} + \log \frac{\mathbf{h}^T \mathbf{\Omega}^{-1} \mathbf{h}}{\mathbf{h}^T \mathbf{h}} \right\} \\
&\geq \min_{\mathbf{h} \in \mathbb{R}^{d_X}} \left\{ \log \frac{\mathbf{h}^T (\mathbf{\Omega} - \boldsymbol{\xi} \boldsymbol{\xi}^T) \mathbf{h}}{\mathbf{h}^T \mathbf{h}} + \log \frac{\mathbf{h}^T \mathbf{\Omega}^{-1} \mathbf{h}}{\mathbf{h}^T \mathbf{h}} \right\}
\end{aligned}$$

987 The last inequality can be seen by noticing $\min_{\mathbf{h} \in \mathbb{R}^{d_X}} \left\{ \log \frac{\mathbf{h}^T (\mathbf{\Omega} - \boldsymbol{\xi} \boldsymbol{\xi}^T) \mathbf{h}}{\mathbf{h}^T \mathbf{h}} + \log \frac{\mathbf{h}^T \mathbf{\Omega}^{-1} \mathbf{h}}{\mathbf{h}^T \mathbf{h}} \right\} <$

988 0. This is because that

$$\log \frac{\mathbf{h}^T (\mathbf{\Omega} - \boldsymbol{\xi} \boldsymbol{\xi}^T) \mathbf{h}}{\mathbf{h}^T \mathbf{h}} + \log \frac{\mathbf{h}^T \mathbf{\Omega}^{-1} \mathbf{h}}{\mathbf{h}^T \mathbf{h}} < \log \frac{\mathbf{h}^T \mathbf{\Omega} \mathbf{h}}{\mathbf{h}^T \mathbf{h}} + \log \frac{\mathbf{h}^T \mathbf{\Omega}^{-1} \mathbf{h}}{\mathbf{h}^T \mathbf{h}},$$

989 where the minimum of the right hand side is zero by taking $\mathbf{h}$ equals to any eigenvector of
990 $\mathbf{\Omega}$. The inequality holds because the strict concaveness of logarithm function and because
991 we have $\mathbf{h}^T \boldsymbol{\xi} \boldsymbol{\xi}^T \mathbf{h} > 0$ for all $\mathbf{h} \neq 0$. (If $\mathbf{h}^T \boldsymbol{\xi} \boldsymbol{\xi}^T \mathbf{h} = 0$ for some $\mathbf{h}$ then it will contradict
992 with the dimension of the envelope.)

993 Moreover, the lower bound of $f(\mathbf{h}, \mathbf{h}_0)$, which is negative, will be attained if we let
994 $W_h = 1$ and let $\mathbf{h} = \arg \min_{\mathbf{h} \in \mathbb{R}^{d_X}} \left\{ \log \frac{\mathbf{h}^T (\mathbf{\Omega} - \boldsymbol{\xi} \boldsymbol{\xi}^T) \mathbf{h}}{\mathbf{h}^T \mathbf{h}} + \log \frac{\mathbf{h}^T \mathbf{\Omega}^{-1} \mathbf{h}}{\mathbf{h}^T \mathbf{h}} \right\}$. So we have the
995 minimum found at $W_h = \frac{\mathbf{h}^T \mathbf{h}}{\mathbf{h}^T \mathbf{h} + \mathbf{h}_0^T \mathbf{h}_0} = 1$, or equivalently, $\ell = \mathbf{L} \mathbf{h} \in \text{span}(\mathbf{L})$.

996 **F.2 The $(k + 1)$ -th direction ℓ_{k+1} in the population**

997 Suppose $\mathbf{L}_k = (\ell_1, \dots, \ell_k) \in \text{span}(\mathbf{L})$, $k < d_X$, then we want to show that $\ell_{k+1} \in$
998 $\mathcal{E}_{\Sigma_X}(\mathcal{L})$. Since the minimization is over

$$f_k(\mathbf{v}) = \log(\mathbf{v}^T \mathbf{L}_{0k}^T \Sigma_{\mathbf{X}|\mathbf{Y}} \mathbf{L}_{0k} \mathbf{v}) + \log(\mathbf{v}^T (\mathbf{L}_{0k}^T \Sigma_{\mathbf{X}} \mathbf{L}_{0k})^{-1} \mathbf{v}), \quad (\text{F.3})$$

999 subject to $\mathbf{v}^T \mathbf{v} = 1$. We know that this gives $\mathbf{v}_{k+1} \in \mathcal{E}_{\mathbf{L}_{0k}^T \Sigma_{\mathbf{X}} \mathbf{L}_{0k}}(\mathbf{L}_{0k}^T \mathcal{L})$ which will lead
1000 to $\ell_{k+1} = \mathbf{L}_{0k} \mathbf{v}_{k+1} \in \mathcal{E}_{\Sigma_X}(\mathcal{L})$ by Lemma 5.

1001 F.3 The $\sqrt{n}$ -consistency for the sample algorithm

1002 The proof hinges on Amemiya's (1985) results on the asymptotic properties of extremum
1003 estimators. We first state these results and then sketch how they can be used to prove the
1004 $\sqrt{n}$ -consistency for our algorithm.

1005 Let $Q_n(\mathbf{y}, \boldsymbol{\theta})$ be a real-valued function of the random variables $\mathbf{y} = (\mathbf{y}_1, \dots, \mathbf{y}_n)^T$ and
1006 the parameters $\boldsymbol{\theta} = (\theta_1, \dots, \theta_K)^T$. We shall sometimes write $Q_n(\mathbf{y}, \boldsymbol{\theta})$ more compactly
1007 as $Q_n(\boldsymbol{\theta})$. Let the parameter space be Θ and let the true value of $\boldsymbol{\theta}$ be $\boldsymbol{\theta}_0$ which is in Θ .
1008 Then Theorem 4.1.1 in Amemiya (1985) gave the consistency of the extremum estimator,
1009 $\hat{\boldsymbol{\theta}}_n = \arg \max_{\boldsymbol{\theta} \in \Theta} Q_n(\mathbf{y}, \boldsymbol{\theta})$, as follows.

1010 **Theorem 1.** *Under the following assumptions, $\hat{\boldsymbol{\theta}}_n$ converges to $\boldsymbol{\theta}_0$ in probability.*

1011 (A) *The parameter space Θ is a compact subset of $\mathbb{R}^K$.*

1012 (B) *$Q_n(\mathbf{y}, \boldsymbol{\theta})$ is continuous in $\boldsymbol{\theta} \in \Theta$ for all $\mathbf{y}$ and is a measurable function of $\mathbf{y}$ for all
1013 $\boldsymbol{\theta} \in \Theta$.*

1014 (C) *$n^{-1}Q_n(\boldsymbol{\theta})$ converges to a nonstochastic function $Q(\boldsymbol{\theta})$ in probability uniformly in
1015 $\boldsymbol{\theta} \in \Theta$ as n goes to infinity, and $Q(\boldsymbol{\theta})$ attains a unique global maximum at $\boldsymbol{\theta}_0$.*

1016 The continuity of $Q(\boldsymbol{\theta})$ then follows from assumptions (A) – (C). For a consistent
1017 $\hat{\boldsymbol{\theta}}_n$, Theorem 4.1.3 in Amemiya (1985) established the asymptotic normality of $\hat{\boldsymbol{\theta}}_n$ as
1018 follows.

Theorem 2. *In addition to the assumptions in Theorem 1, we assume the following con-
ditions.*

(D) *$\partial^2 Q_n(\boldsymbol{\theta}) / \partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T$ exists and is continuous in an open, convex neighborhood of $\boldsymbol{\theta}_0$.*

(E) *$n^{-1} \{ \partial^2 Q_n(\mathbf{y}, \boldsymbol{\theta}) / \partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T \}_{\boldsymbol{\theta}=\boldsymbol{\theta}_n^*}$ converges to a finite nonsingular matrix*

$$\mathbf{A}(\boldsymbol{\theta}_0) = \lim_{n \rightarrow \infty} \mathbb{E} \left\{ n^{-1} \{ \partial^2 Q_n(\boldsymbol{\theta}) / \partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T \}_{\boldsymbol{\theta}=\boldsymbol{\theta}_0} \right\},$$

for any random sequences $\boldsymbol{\theta}_n^$ such that $\text{plim}(\boldsymbol{\theta}_n^*) = \boldsymbol{\theta}_0$.*

(F) *$n^{-1/2} \{ \partial Q_n(\boldsymbol{\theta}) / \partial \boldsymbol{\theta} \}_{\boldsymbol{\theta}=\boldsymbol{\theta}_0} \rightarrow N(0, \mathbf{B}(\boldsymbol{\theta}_0))$, where*

$$\mathbf{B}(\boldsymbol{\theta}_0) = \lim_{n \rightarrow \infty} \mathbb{E} \left\{ n^{-1} \{ \partial Q_n(\boldsymbol{\theta}) / \partial \boldsymbol{\theta} \}_{\boldsymbol{\theta}=\boldsymbol{\theta}_0} \{ \partial Q_n(\boldsymbol{\theta}) / \partial \boldsymbol{\theta}^T \}_{\boldsymbol{\theta}=\boldsymbol{\theta}_0} \right\}.$$

1019 *Then $\sqrt{n}(\hat{\boldsymbol{\theta}}_n - \boldsymbol{\theta}_0) \rightarrow N(0, \mathbf{A}(\boldsymbol{\theta}_0)^{-1} \mathbf{B}(\boldsymbol{\theta}_0) \mathbf{A}(\boldsymbol{\theta}_0)^{-1})$.*

1020 In our adaptation of Theorem 1 and Theorem 2, we let $\boldsymbol{\theta} \equiv \ell$ whose true value
 1021 is denoted by ℓ_0 and let the random variables $\mathbf{y} = \{(\mathbf{X}_1, \mathbf{Y}_1), \dots, (\mathbf{X}_n, \mathbf{Y}_n)\}$ where
 1022 $(\mathbf{X}_i, \mathbf{Y}_i)$'s, $i = 1, \dots, n$, are i.i.d. random variables. The parameter space is the 1D
 1023 manifold $\boldsymbol{\Theta} = \mathcal{G}_{(p,1)}$ which is a compact set of $\mathbb{R}^{p \times 1}$, so condition (A) in Theorem 1 is
 1024 satisfied. The function to be maximized is defined as follows.

$$Q_n(\ell) = -n/2 \log(\ell^T \mathbf{S}_{\mathbf{X}}^{-1} \ell) - n/2 \log(\ell^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell) + n \log(\ell^T \ell). \quad (\text{F.4})$$

1025 Condition (B) holds because that $\mathbf{S}_{\mathbf{X}|\mathbf{Y}}$ and $\mathbf{S}_{\mathbf{X}}^{-1}$ are measurable functions of the i.i.d.
 1026 variables $\{(\mathbf{X}_1, \mathbf{Y}_1), \dots, (\mathbf{X}_n, \mathbf{Y}_n)\}$, and because Q_n is continuous in $\boldsymbol{\theta}$, $\mathbf{S}_{\mathbf{X}}^{-1}$ and $\mathbf{S}_{\mathbf{X}|\mathbf{Y}}$.

1027 We next verify condition (C) that $n^{-1}Q_n(\ell)$ converges uniformly to

$$Q(\ell) = -1/2 \log(\ell^T \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \ell) - 1/2 \log(\ell^T \boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}} \ell) + \log(\ell^T \ell). \quad (\text{F.5})$$

1028 It follows from Section F.1 that $Q(\ell)$ attains the unique global maximum at ℓ_0 . For sim-
 1029 plicity, we assume $\boldsymbol{\Sigma}_{\mathbf{X}}$ and $\boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}}$ both have distinct eigenvalues so that ℓ_0 is the unique
 1030 maximum of $Q_n(\ell)$ in the 1D manifold $\boldsymbol{\Theta}$. For the case where there are multiple local
 1031 maxima of $Q(\ell)$, we can obtain similar results by applying Theorem 4.1.2 in Amemiya
 1032 (1985) as an alternative of Theorem 1. When p and d_X are fixed and n goes to infinity,
 1033 the eigenvectors and eigenvalues of $\mathbf{S}_{\mathbf{X}}$, $\mathbf{S}_{\mathbf{X}}^{-1}$ and $\mathbf{S}_{\mathbf{X}|\mathbf{Y}}$ are $\sqrt{n}$ -consistent for the eigen-
 1034 vectors and eigenvalues of $\boldsymbol{\Sigma}_{\mathbf{X}}$, $\boldsymbol{\Sigma}_{\mathbf{X}}^{-1}$ and $\boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}}$. Then $n^{-1}Q_n(\ell)$ converge in probability
 1035 to $Q(\ell)$ uniformly in ℓ , as can be seen from the following argument.

$$\begin{aligned} n^{-1}Q_n(\ell) - Q(\ell) &= -1/2 (\log(\ell^T \mathbf{S}_{\mathbf{X}}^{-1} \ell) - \log(\ell^T \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \ell)) \\ &\quad -1/2 (\log(\ell^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell) - \log(\ell^T \boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}} \ell)) \\ &= -1/2 \log(\ell^T \mathbf{S}_{\mathbf{X}}^{-1} \ell / \ell^T \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \ell) - 1/2 \log(\ell^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell / \ell^T \boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}} \ell). \end{aligned}$$

1036 Hence, $\sup_{\ell \in \boldsymbol{\Theta}} \log(\ell^T \mathbf{S}_{\mathbf{X}}^{-1} \ell / \ell^T \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \ell) = \sup_{\ell \in \boldsymbol{\Theta}} \log(\ell^T \boldsymbol{\Sigma}_{\mathbf{X}}^{1/2} \mathbf{S}_{\mathbf{X}}^{-1} \boldsymbol{\Sigma}_{\mathbf{X}}^{1/2} \ell / \ell^T \ell)$. Notice that
 1037 $\sup_{\ell \in \boldsymbol{\Theta}} \log(\ell^T \boldsymbol{\Sigma}_{\mathbf{X}}^{1/2} \mathbf{S}_{\mathbf{X}}^{-1} \boldsymbol{\Sigma}_{\mathbf{X}}^{1/2} \ell / \ell^T \ell)$ equals to the logarithm of the largest eigenvalue of
 1038 $\boldsymbol{\Sigma}_{\mathbf{X}}^{1/2} \mathbf{S}_{\mathbf{X}}^{-1} \boldsymbol{\Sigma}_{\mathbf{X}}^{1/2}$ which converges to 0 in probability. Similarly, $\sup_{\ell \in \boldsymbol{\Theta}} \log(\ell^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell / \ell^T \boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}} \ell)$
 1039 converges to zero in probability. Therefore, $n^{-1}Q_n(\ell)$ converges to $Q(\ell)$ in probability
 1040 uniformly in $\ell \in \boldsymbol{\Theta}$. Note that we have assumed $\boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}} > 0$ and $\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} > 0$, so their
 1041 eigenvalues will be bounded away from zero.

1042 We next verify conditions (D) – (F). By straightforward calculation, condition (D)
 1043 follows from the second derivative matrix

$$\begin{aligned}
 n^{-1} \frac{\partial^2 Q_n(\ell)}{\partial \ell \partial \ell^T} &= 2(\ell^T \mathbf{S}_{\mathbf{X}}^{-1} \ell)^{-2} (\mathbf{S}_{\mathbf{X}}^{-1} \ell \ell^T \mathbf{S}_{\mathbf{X}}^{-1}) - (\ell^T \mathbf{S}_{\mathbf{X}}^{-1} \ell)^{-1} \mathbf{S}_{\mathbf{X}}^{-1} \\
 &\quad + 2(\ell^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell)^{-2} (\mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell \ell^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}}) - (\ell^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell)^{-1} \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \\
 &\quad - 2(\ell^T \ell)^{-2} \mathbf{P}_\ell + (\ell^T \ell)^{-1} \mathbf{I}_p.
 \end{aligned} \tag{F.6}$$

1044 Condition (E) holds because the above quantity is a smooth function of ℓ , $\mathbf{S}_{\mathbf{X}}^{-1}$ and $\mathbf{S}_{\mathbf{X}|\mathbf{Y}}$.

1045 Last, we need to verify condition (F). From the proof of Theorem 2, we only need to
 1046 show that $n^{-1} \{\partial Q_n(\boldsymbol{\theta}) / \partial \boldsymbol{\theta}\}_{\boldsymbol{\theta}=\boldsymbol{\theta}_0} = O_p(1/\sqrt{n})$ for $\sqrt{n}$ -consistency of the estimator $\hat{\boldsymbol{\theta}}_n$.

$$n^{-1} \{\partial Q_n(\ell) / \partial \ell\}_{\ell=\ell_0} = -(\ell_0^T \mathbf{S}_{\mathbf{X}}^{-1} \ell_0)^{-1} \mathbf{S}_{\mathbf{X}}^{-1} \ell_0 - (\ell_0^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell_0)^{-1} \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \ell_0 + 2\ell_0. \tag{F.7}$$

1047 Following the derivation in Section F.1, we know that $\{\partial Q(\ell) / \partial \ell\}_{\ell=\ell_0} = 0$. Then the
 1048 result follows from the fact that $n^{-1} \partial Q_n(\ell) / \partial \ell$ is a smooth function of $\mathbf{S}_{\mathbf{X}}^{-1}$ and $\mathbf{S}_{\mathbf{X}|\mathbf{Y}}$
 1049 which are $\sqrt{n}$ -consistent estimators.

1050 So far, we have verified the conditions (A) – (F) so that the sample estimator $\hat{\ell}$ will
 1051 be $\sqrt{n}$ -consistent for the population estimator. For the $(k+1)$ -th direction, $k < d_X$,
 1052 let $\hat{\mathbf{L}}_k$ denote an $\sqrt{n}$ -consistent estimator of the first k directions and let $(\hat{\mathbf{L}}_k, \hat{\mathbf{L}}_{0k})$ be
 1053 an orthogonal matrix. The $(k+1)$ -th direction is defined by $\ell_{k+1} = \hat{\mathbf{L}}_{0k} \mathbf{v}$ where the
 1054 parameters are $\mathbf{v} \in \boldsymbol{\Theta} \subset \mathbb{R}^{p-k}$ and the parameter space is $\boldsymbol{\Theta} = \mathcal{G}_{p-k,1}$. We show that we
 1055 can obtain a $\sqrt{n}$ -consistent estimator $\hat{\mathbf{v}}_n$, so the $\sqrt{n}$ -consistency of $\hat{\ell}_{k+1} = \hat{\mathbf{L}}_{0k} \hat{\mathbf{v}}_n$ then
 1056 follows. From Section F.2, we define our function $Q_n(\mathbf{v})$ and $Q(\mathbf{v})$ as

$$\begin{aligned}
 Q_n(\mathbf{v}) &= -n/2 \log(\mathbf{v}^T \hat{\mathbf{L}}_{0k}^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \hat{\mathbf{L}}_{0k} \mathbf{v}) - n/2 \log(\mathbf{v}^T \hat{\mathbf{L}}_{0k}^T \mathbf{S}_{\mathbf{X}}^{-1} \hat{\mathbf{L}}_{0k} \mathbf{v}) + n \log(\mathbf{v}^T \mathbf{v}) \\
 Q(\mathbf{v}) &= -1/2 \log(\mathbf{v}^T \mathbf{L}_{0k}^T \boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}} \mathbf{L}_{0k} \mathbf{v}) - 1/2 \log(\mathbf{v}^T \mathbf{L}_{0k}^T \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \mathbf{L}_{0k} \mathbf{v}) + \log(\mathbf{v}^T \mathbf{v})
 \end{aligned}$$

1057 Following the same logic as verifying the conditions for the first direction, we can see
 1058 that $\hat{\mathbf{v}}_n = \arg \max Q_n(\mathbf{v})$ will be $\sqrt{n}$ -consistent for $\mathbf{v}_0 = \arg \max Q(\mathbf{v})$ by notic-
 1059 ing that $\hat{\mathbf{L}}_{0k}^T \mathbf{S}_{\mathbf{X}|\mathbf{Y}} \hat{\mathbf{L}}_{0k}$ and $\hat{\mathbf{L}}_{0k}^T \mathbf{S}_{\mathbf{X}}^{-1} \hat{\mathbf{L}}_{0k}$ are $\sqrt{n}$ -consistent estimators for $\mathbf{L}_{0k}^T \boldsymbol{\Sigma}_{\mathbf{X}|\mathbf{Y}} \mathbf{L}_{0k}$ and
 1060 $\mathbf{L}_{0k}^T \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \mathbf{L}_{0k}$. Since all the d_X directions will be $\sqrt{n}$ -consistent, the projection onto $\hat{\mathbf{L}}_{1D} =$
 1061 $(\hat{\ell}_1, \dots, \hat{\ell}_{d_X})$ will be a $\sqrt{n}$ -consistent estimator for the projection onto the envelope
 1062 $\mathcal{E}_{\boldsymbol{\Sigma}_{\mathbf{X}}}(\mathcal{L})$.

G Proof for Proposition 4

The Jacobian matrix $\mathbf{J}_h$ can be obtained by following Cook et al. (2010), and we make use of the following lemma from Cook et al. (2013).

Lemma 6. *Suppose that $\mathbf{A} \in \mathbb{S}^{t \times t}$ is non-singular and that the column-partition $(\mathbf{B}, \mathbf{B}_0) \in \mathbb{R}^{t \times t}$ is orthogonal. Then (1) $|\mathbf{B}^T \mathbf{A} \mathbf{B}| = |\mathbf{A}| \cdot |\mathbf{B}_0^T \mathbf{A}^{-1} \mathbf{B}_0|$, and (2) $|\mathbf{A}| \leq |\mathbf{B}^T \mathbf{A} \mathbf{B}| \cdot |\mathbf{B}_0^T \mathbf{A} \mathbf{B}_0|$ with equality if and only if $\text{span}(\mathbf{B})$ reduces $\mathbf{A}$.*

In the population, the objective function $F(\mathbf{L} \oplus \mathbf{R})$ from (4.4) can be written as,

$$\begin{aligned} F(\mathbf{L} \oplus \mathbf{R}) &= \log |\mathbf{L}^T \Sigma_{\mathbf{X}} \mathbf{L}| + \log |\mathbf{L}_0^T \Sigma_{\mathbf{X}} \mathbf{L}_0| + \log |\mathbf{R}_0^T \Sigma_{\mathbf{Y}} \mathbf{R}_0| \\ &\quad + \log |\mathbf{R}^T (\Sigma_{\mathbf{Y}} - \Sigma_{\mathbf{X}\mathbf{Y}}^T \mathbf{L} (\mathbf{L}^T \Sigma_{\mathbf{X}} \mathbf{L})^{-1} \mathbf{L}^T \Sigma_{\mathbf{X}\mathbf{Y}}) \mathbf{R}|. \\ &\geq \log |\Sigma_{\mathbf{X}}| + \log |\mathbf{R}_0^T \Sigma_{\mathbf{Y}} \mathbf{R}_0| + \log |\mathbf{R}^T \Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}} \mathbf{R}| \\ &\geq \log |\Sigma_{\mathbf{X}}| + \log |\mathbf{R}_0^T \Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}} \mathbf{R}_0| + \log |\mathbf{R}^T \Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}} \mathbf{R}| \\ &\geq \log |\Sigma_{\mathbf{X}}| + \log |\Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}}| \\ &\geq \log |\Sigma_{\mathbf{X}}| + \log |\Sigma_{\mathbf{Y}|\mathbf{X}}|, \end{aligned}$$

where $\Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}} = \Sigma_{\mathbf{Y}} - \Sigma_{\mathbf{X}\mathbf{Y}}^T \mathbf{L} (\mathbf{L}^T \Sigma_{\mathbf{X}} \mathbf{L})^{-1} \mathbf{L}^T \Sigma_{\mathbf{X}\mathbf{Y}}$. The first and third inequalities came directly from Lemma 6, which also says that the first inequality turns to equality if and only if (C1) $\text{span}(\mathbf{L})$ reduces $\Sigma_{\mathbf{X}}$ and that the third inequality becomes equality if and only if (C3) $\text{span}(\mathbf{R})$ reduces $\Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}}$. The second and the last inequalities are obtained from the fact that $\Sigma_{\mathbf{Y}} \geq \Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}}$ and $\Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}} \geq \Sigma_{\mathbf{Y}|\mathbf{X}}$ for any $\mathbf{L}$. Hence, these two inequalities become equalities if and only if (C2) $\text{span}(\Sigma_{\mathbf{Y}} - \Sigma_{\mathbf{Y}|\mathbf{L}^T \mathbf{X}}) \subseteq \text{span}(\mathbf{R})$ and (C4) $\text{span}(\Sigma_{\mathbf{X}\mathbf{Y}}^T) \subseteq \text{span}(\mathbf{L})$. Since d_X is the dimension of the smallest subspace satisfying condition (C1) and (C4), $\text{span}(\mathbf{L}) = \mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{L})$. Therefore, condition (C2) and (C3) will be equivalent to (C2') $\text{span}(\Sigma_{\mathbf{Y}} - \Sigma_{\mathbf{Y}|\mathbf{X}}) = \text{span}(\Sigma_{\mathbf{X}\mathbf{Y}}) \subseteq \text{span}(\mathbf{R})$, and (C3') $\text{span}(\mathbf{R})$ reduces $\Sigma_{\mathbf{Y}|\mathbf{X}}$. Again, because d_Y is the dimension of the smallest subspace satisfying (C2') and (C3'), we can conclude that $\text{span}(\mathbf{R}) = \mathcal{E}_{\Sigma_{\mathbf{Y}|\mathbf{X}}}(\mathcal{R})$.

The rest of the proof relies on Shapiro's (1986) results on the asymptotics of overparameterized structural models. In order to apply Shapiro's (1986) theory in our context, we define the minimum discrepancy function as

$$F(\hat{\mathbf{h}}_{\text{full}}, \mathbf{h}) = \log |\Sigma_{\mathbf{C}}| + \text{trace}(\Sigma_{\mathbf{C}}^{-1} \mathbf{S}_{\mathbf{C}}) - \log |\mathbf{S}_{\mathbf{C}}| - \text{trace}(\mathbf{S}_{\mathbf{C}}^{-1} \mathbf{S}_{\mathbf{C}}). \quad (\text{G.1})$$

1084 Then $F(\hat{\mathbf{h}}_{\text{full}}, \mathbf{h})$ satisfies: (1) $F(\hat{\mathbf{h}}_{\text{full}}, \mathbf{h}) \geq 0$ for all $\hat{\mathbf{h}}_{\text{full}}$ and $\mathbf{h}$; (2) $F(\hat{\mathbf{h}}_{\text{full}}, \mathbf{h}) = 0$
1085 if and only if $\hat{\mathbf{h}}_{\text{full}} = \mathbf{h}$; and (3) $F(\hat{\mathbf{h}}_{\text{full}}, \mathbf{h})$ is twice continuously differentiable in $\hat{\mathbf{h}}_{\text{full}}$
1086 and $\mathbf{h}$. Recall from Section 5.1 that we use the subscript 0 to emphasize the true param-
1087 eter: $\mathbf{h}_0$ and ϕ_0 correspond to the true distribution of $\mathbf{C}$. Then $\hat{\mathbf{h}}_{\text{full}}$ is $\sqrt{n}$ -consistent
1088 for $\mathbf{h}_0$. Notice that $\hat{\mathbf{h}}_{\text{full}}$ is a smooth function of the sample covariance matrices which
1089 converges in distribution to the population covariance matrices, then by the delta method
1090 we know $\sqrt{n}(\hat{\mathbf{h}}_{\text{full}} - \mathbf{h}_0) \rightarrow N(0, \Gamma)$, for some positive definite covariance Γ . Using
1091 Shapiro's (1986) Proposition 3.1 and Proposition 4.1, we will have $\sqrt{n}$ -consistency re-
1092 sults for $\text{vec}(\hat{\mathbf{h}})$. We next need to compare $\Gamma = \text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{full}})$ with $\mathbf{W} = \text{avar}(\sqrt{n}\hat{\mathbf{h}})$.

1093 By noticing that $\mathbf{P}_{\mathbf{J}_h^{1/2}\Delta} = \mathbf{J}_h^{1/2}\Delta(\Delta^T\mathbf{J}_h\Delta)^\dagger\Delta^T\mathbf{J}_h^{1/2}$, we have

$$\begin{aligned}\Gamma - \mathbf{W} &= \Gamma - \mathbf{J}_h^{-1/2}\mathbf{P}_{\mathbf{J}_h^{1/2}\Delta}\mathbf{J}_h^{1/2}\Gamma\mathbf{J}_h^{1/2}\mathbf{P}_{\mathbf{J}_h^{1/2}\Delta}\mathbf{J}_h^{-1/2} \\ &= \mathbf{J}_h^{-1/2}\left(\mathbf{J}_h^{1/2}\Gamma\mathbf{J}_h^{1/2} - \mathbf{P}_{\mathbf{J}_h^{1/2}\Delta}\mathbf{J}_h^{1/2}\Gamma\mathbf{J}_h^{1/2}\mathbf{P}_{\mathbf{J}_h^{1/2}\Delta}\right)\mathbf{J}_h^{-1/2}.\end{aligned}$$

1094 Then the conclusion follows from the above expression.

1095 H Proof for Proposition 5

1096 Proposition 5 can be seen from Proposition 4 by letting $\Gamma = \mathbf{J}_h^{-1}$ under the normality
1097 assumption.

1098 I Proof for Proposition 6

1099 I.1 The gradient matrix $\Delta = \frac{\partial \mathbf{h}(\phi)}{\partial \phi}$

1100 The first three columns of the gradient matrix $\Delta \equiv (\delta_1, \dots, \delta_7)$ are

$$(\delta_1, \dots, \delta_3) = \begin{pmatrix} \mathbf{R} \otimes \mathbf{L} & \mathbf{R}\boldsymbol{\eta}^T \otimes \mathbf{I}_p & \mathbf{K}_{r,p}(\mathbf{L}\boldsymbol{\eta} \otimes \mathbf{I}_r) \\ 0 & 2\mathbf{C}_p(\mathbf{L}\boldsymbol{\Omega} \otimes \mathbf{I}_p - \mathbf{L} \otimes \mathbf{L}_0\boldsymbol{\Omega}_0\mathbf{L}_0^T) & 0 \\ 0 & 0 & 2\mathbf{C}_r(\mathbf{R}\boldsymbol{\Phi} \otimes \mathbf{I}_r - \mathbf{R} \otimes \mathbf{R}_0\boldsymbol{\Phi}_0\mathbf{R}_0^T) \end{pmatrix},$$

1101 and the last four columns are

$$(\delta_4, \dots, \delta_7) = \begin{pmatrix} 0 & 0 & 0 & 0 \\ \mathbf{C}_p(\mathbf{L} \otimes \mathbf{L})\mathbf{E}_{d_X} & \mathbf{C}_p(\mathbf{L}_0 \otimes \mathbf{L}_0)\mathbf{E}_{p-d_X} & 0 & 0 \\ 0 & 0 & \mathbf{C}_r(\mathbf{R} \otimes \mathbf{R})\mathbf{E}_{d_Y} & \mathbf{C}_r(\mathbf{R}_0 \otimes \mathbf{R}_0)\mathbf{E}_{r-d_Y} \end{pmatrix},$$

1102 where $\mathbf{C}_p \in \mathbb{R}^{p(p+1)/2 \times p^2}$ is called the contraction matrix and $\mathbf{E}_p \in \mathbb{R}^{p^2 \times p(p+1)/2}$ is called
 1103 the extraction matrix (Henderson and Searle, 1979). They relate the vec operation and
 1104 vech operation for any symmetric matrix $\mathbf{M} \in \mathbb{S}^{p \times p}$ as: $\text{vech}(\mathbf{M}) = \mathbf{C}_p \text{vec}(\mathbf{M})$ and
 1105 $\text{vec}(\mathbf{M}) = \mathbf{E}_p \text{vech}(\mathbf{M})$. More properties about the contraction and extraction matrices
 1106 can be found from Henderson and Searle (1979).

1107 *Proof.* We can write $\Delta = \frac{\partial \mathbf{h}(\phi)}{\partial \phi}$ as $\begin{pmatrix} \partial \mathbf{h}_1 / \partial \phi \\ \partial \mathbf{h}_2 / \partial \phi \\ \partial \mathbf{h}_3 / \partial \phi \end{pmatrix}$ and compute it by rows.

1108 (1) Computation of $\frac{\partial \mathbf{h}_1}{\partial \phi}$.

1109 First, $\mathbf{h}_1(\phi) = \text{vec}(\beta) = \text{vec}(\mathbf{L}\eta\mathbf{R}^T)$ and thus, $\partial \mathbf{h}_1(\phi) / \partial \phi_k = 0$, $k = 4, 5, 6, 7$.

1110 The rest of the terms are:

$$\begin{aligned} \frac{\partial \mathbf{h}_1(\phi)}{\partial \phi_1} &= \frac{\partial [\text{vec}(\mathbf{L}\eta\mathbf{R}^T)]}{\partial \text{vec}(\eta)} = \frac{\partial [\{\mathbf{R} \otimes \mathbf{L}\} \text{vec}(\eta)]}{\partial \text{vec}(\eta)} = \mathbf{R} \otimes \mathbf{L} \\ \frac{\partial \mathbf{h}_1(\phi)}{\partial \phi_2} &= \frac{\partial [\text{vec}(\mathbf{L}\eta\mathbf{R}^T)]}{\partial \text{vec}(\mathbf{L})} = \frac{\partial [\{\mathbf{R}\eta^T \otimes \mathbf{I}_p\} \text{vec}(\mathbf{L})]}{\partial \text{vec}(\mathbf{L})} = \mathbf{R}\eta^T \otimes \mathbf{I}_p \\ \frac{\partial \mathbf{h}_1(\phi)}{\partial \phi_3} &= \frac{\partial [\text{vec}(\mathbf{L}\eta\mathbf{R}^T)]}{\partial \text{vec}(\mathbf{R})} = \frac{\partial [\{\mathbf{I}_r \otimes \mathbf{L}\eta\} \text{vec}(\mathbf{R}^T)]}{\partial \text{vec}(\mathbf{R})} = \{\mathbf{I}_r \otimes \mathbf{L}\Omega^{-1}\eta\} \mathbf{K}_{r,m} \\ &= \mathbf{K}_{r,p} \{\mathbf{L}\eta \otimes \mathbf{I}_r\}. \end{aligned}$$

1111 (2) Computation of $\partial \mathbf{h}_2 / \partial \phi$.

1112 Since $\mathbf{h}_2(\phi) = \text{vech}(\Sigma_{\mathbf{X}}) = \text{vech}(\mathbf{L}\Omega\mathbf{L}^T + \mathbf{L}_0\Omega_0\mathbf{L}_0^T)$, it follows that $\partial \mathbf{h}_2(\phi) / \partial \phi_k =$
 1113 0 , $k = 1, 3, 6, 7$. The rest terms can be found in Cook, Li and Chiaromonte (2010). In
 1114 our notation, they are

$$\begin{aligned} \frac{\partial \mathbf{h}_2(\phi)}{\partial \phi_2} &= 2\mathbf{C}_p(\mathbf{L}\Omega \otimes \mathbf{I}_p - \mathbf{L} \otimes \mathbf{L}_0\Omega_0\mathbf{L}_0^T), \\ \frac{\partial \mathbf{h}_2(\phi)}{\partial \phi_4} &= \mathbf{C}_p(\mathbf{L} \otimes \mathbf{L})\mathbf{E}_{d_X}, \\ \frac{\partial \mathbf{h}_2(\phi)}{\partial \phi_5} &= \mathbf{C}_p(\mathbf{L}_0 \otimes \mathbf{L}_0)\mathbf{E}_{p-d_X}. \end{aligned}$$

1115 (3) Computation of $\partial \mathbf{h}_3 / \partial \phi$.

1116 Since $\mathbf{h}_3(\phi) = \text{vech}(\Sigma_{\mathbf{Y}|\mathbf{X}}) = \text{vech}(\mathbf{R}\Phi\mathbf{R}^T + \mathbf{R}_0\Phi_0\mathbf{R}_0^T)$, we have

$$\begin{aligned}\frac{\partial \mathbf{h}_3(\phi)}{\partial \phi_k} &= 0, \quad k = 1, 2, 4, 5, \\ \frac{\partial \mathbf{h}_3(\phi)}{\partial \phi_3} &= 2\mathbf{C}_r(\mathbf{R}\Phi \otimes \mathbf{I}_r - \mathbf{R} \otimes \mathbf{R}_0\Phi_0\mathbf{R}_0^T), \\ \frac{\partial \mathbf{h}_3(\phi)}{\partial \phi_6} &= \mathbf{C}_r(\mathbf{R} \otimes \mathbf{R})\mathbf{E}_{d_Y}, \\ \frac{\partial \mathbf{h}_3(\phi)}{\partial \phi_7} &= \mathbf{C}_r(\mathbf{R}_0 \otimes \mathbf{R}_0)\mathbf{E}_{r-d_Y}.\end{aligned}$$

1117

□

1118 I.2 Transformation $\Delta = \tilde{\Delta}\mathbf{T}$

1119 We want to find a transformation such that $\text{span}(\Delta) = \text{span}(\tilde{\Delta})$ and $\tilde{\Delta}^T \mathbf{J}_h \tilde{\Delta}$ is block
1120 diagonal. Partition Δ into its first three columns and its last four column, and denote them
1121 as $\Delta \equiv (\Delta_1, \Delta_2)$. It is easy to check that $\Delta_2^T \mathbf{J}_h \Delta_2$ is already block diagonal. So we
1122 need to find a transformed $\tilde{\Delta}_1$ such that $\tilde{\Delta}_1^T \mathbf{J}_h \tilde{\Delta}_1$ is block diagonal and $\tilde{\Delta}_1^T \mathbf{J}_h \Delta_2 = 0$.

1123 By direct computation, we have

$$\begin{aligned}\tilde{\Delta}_1 &= \begin{pmatrix} \mathbf{R} \otimes \mathbf{L} & \mathbf{R}\boldsymbol{\eta}^T \otimes \mathbf{L}_0 & \mathbf{K}_{r,p}(\mathbf{L}\boldsymbol{\eta} \otimes \mathbf{R}_0) \\ 0 & 2\mathbf{C}_p(\mathbf{L}\boldsymbol{\Omega} \otimes \mathbf{L}_0 - \mathbf{L} \otimes \mathbf{L}_0\boldsymbol{\Omega}_0) & 0 \\ 0 & 0 & 2\mathbf{C}_r(\mathbf{R}\Phi \otimes \mathbf{R}_0 - \mathbf{R} \otimes \mathbf{R}_0\Phi_0) \end{pmatrix}, \\ &\equiv (\tilde{\boldsymbol{\delta}}_1, \tilde{\boldsymbol{\delta}}_2, \tilde{\boldsymbol{\delta}}_3) \\ \tilde{\Delta}_2 &= \Delta_2.\end{aligned}$$

1124 And the 7×7 blocks transformation matrix $\mathbf{T}$ can be written as a 7×3 blocks matrix $\mathbf{T}_1$
1125 and a 7×4 block matrix $\mathbf{T}_2 = \begin{pmatrix} \mathbf{0}_{3 \times 4} \\ \mathbf{I}_{4 \times 4} \end{pmatrix}$ whose blocks conform to those of $\tilde{\Delta}$. Also,
1126 $\tilde{\Delta}\mathbf{T}_1 = \Delta_1$ and $\mathbf{T}_1$ is given by

$$\mathbf{T}_1 = \begin{pmatrix} \mathbf{I}_{d_X d_Y} & \boldsymbol{\eta}^T \otimes \mathbf{L}^T & \mathbf{K}_{d_Y, d_X}(\boldsymbol{\eta} \otimes \mathbf{R}^T) \\ 0 & \mathbf{I}_{d_X} \otimes \mathbf{L}_0^T & 0 \\ 0 & 0 & \mathbf{I}_{d_Y} \otimes \mathbf{R}_0^T \\ 0 & 2\mathbf{C}_{d_X}(\boldsymbol{\Omega} \otimes \mathbf{L}^T) & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 2\mathbf{C}_{d_Y}(\Phi \otimes \mathbf{R}^T) \\ 0 & 0 & 0 \end{pmatrix}.$$

1127 It can be easily noticed that $\mathbf{T}$ has full row rank, hence, $\text{span}(\mathbf{\Delta}) = \text{span}(\tilde{\mathbf{\Delta}})$. There-
 1128 fore, the asymptotic covariance for $\hat{\mathbf{h}}$ under the envelope model is given by

$$\begin{aligned} \text{avar}(\sqrt{n}\hat{\mathbf{h}}) &= \mathbf{\Delta} (\mathbf{\Delta}^T \mathbf{J}_h \mathbf{\Delta})^\dagger \mathbf{\Delta}^T = \tilde{\mathbf{\Delta}} \left(\tilde{\mathbf{\Delta}}^T \mathbf{J}_h \tilde{\mathbf{\Delta}} \right)^\dagger \tilde{\mathbf{\Delta}}^T \\ &= \sum_{i=1}^3 \tilde{\boldsymbol{\delta}}_i (\tilde{\boldsymbol{\delta}}_i^T \mathbf{J}_h \tilde{\boldsymbol{\delta}}_i)^\dagger \tilde{\boldsymbol{\delta}}_i^T + \sum_{j=4}^7 \boldsymbol{\delta}_j (\boldsymbol{\delta}_j^T \mathbf{J}_h \boldsymbol{\delta}_j)^\dagger \boldsymbol{\delta}_j^T. \end{aligned}$$

1129 **I.3 Asymptotic covariance matrix for $\text{vec}(\hat{\boldsymbol{\beta}})$**

1130 The asymptotic covariance matrix for $\text{vec}(\hat{\boldsymbol{\beta}})$ is the upper left block of the full covariance
 1131 matrix for $\hat{\mathbf{h}}$. Under the standard model, it is just $\boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}$. From Proposition 2, we
 1132 recognize this as the asymptotic variance of $\text{vec}(\hat{\boldsymbol{\beta}}_{\text{OLS}})$.

$$\text{avar}(\sqrt{n}\text{vec}(\hat{\boldsymbol{\beta}}_{\text{OLS}})) = \boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}. \quad (\text{I.1})$$

1133 For the envelope model, only $\tilde{\boldsymbol{\delta}}_i$, $i = 1, 2, 3$, have contributions to this block. We use
 1134 the subscript $\beta\beta$ to denote the upper left block of dimension $pr \times pr$ for a matrix and
 1135 compute the corresponding three terms.

1136 (1) Contribution from $\tilde{\boldsymbol{\delta}}_1$.

$$\begin{aligned} [\tilde{\boldsymbol{\delta}}_1 (\tilde{\boldsymbol{\delta}}_1^T \mathbf{J}_h \tilde{\boldsymbol{\delta}}_1)^\dagger \tilde{\boldsymbol{\delta}}_1^T]_{\beta\beta} &= (\mathbf{R} \otimes \mathbf{L}) \left((\mathbf{R} \otimes \mathbf{L})^T (\boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}) (\mathbf{R} \otimes \mathbf{L}) \right)^\dagger (\mathbf{R} \otimes \mathbf{L})^T \\ &= (\mathbf{R} \otimes \mathbf{L}) (\boldsymbol{\Phi}^{-1} \otimes \boldsymbol{\Omega})^\dagger (\mathbf{R} \otimes \mathbf{L})^T \\ &= \mathbf{R} \boldsymbol{\Phi} \mathbf{R}^T \otimes \mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T = \text{avar}(\sqrt{n}\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}})). \end{aligned}$$

1137 (2) Contribution from $\tilde{\boldsymbol{\delta}}_2$. Let $\mathbf{A}_1 = \mathbf{E}_p(\mathbf{C}_p(\mathbf{L}\boldsymbol{\Omega} \otimes \mathbf{L}_0 - \mathbf{L} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0))$ then,

$$\begin{aligned} \tilde{\boldsymbol{\delta}}_2^T \mathbf{J}_h \tilde{\boldsymbol{\delta}}_2 &= (\mathbf{R} \boldsymbol{\eta}^T \otimes \mathbf{L}_0)^T (\boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}) (\mathbf{R} \boldsymbol{\eta}^T \otimes \mathbf{L}_0) + 2\mathbf{A}_1^T (\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}) \mathbf{A}_1 \\ &= \boldsymbol{\eta} \boldsymbol{\Phi}^{-1} \boldsymbol{\eta}^T \otimes \boldsymbol{\Omega}_0 + \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1} + \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}_0 - 2\mathbf{I}_{d_X} \otimes \mathbf{I}_{p-d_X} \\ &\equiv \mathbf{M}_2. \end{aligned} \quad (\text{I.2})$$

1138 Thus for the covariance,

$$[\tilde{\boldsymbol{\delta}}_2 (\tilde{\boldsymbol{\delta}}_2^T \mathbf{J}_h \tilde{\boldsymbol{\delta}}_2)^\dagger \tilde{\boldsymbol{\delta}}_2^T]_{\beta\beta} = (\mathbf{R} \boldsymbol{\eta}^T \otimes \mathbf{L}_0) (\mathbf{M}_2)^\dagger (\boldsymbol{\eta} \mathbf{R}^T \otimes \mathbf{L}_0^T). \quad (\text{I.3})$$

1139 (3) Contribution from $\tilde{\delta}_3$. Let $\mathbf{A}_2 = \mathbf{E}_r \mathbf{C}_r (\mathbf{R} \Phi \otimes \mathbf{R}_0 - \mathbf{R} \otimes \mathbf{R}_0 \Phi_0)$ then,

$$\begin{aligned}
\tilde{\delta}_3^T \mathbf{J}_h \tilde{\delta}_3 &= (\mathbf{K}_{r,p} (\mathbf{L} \eta \otimes \mathbf{R}_0))^T (\Sigma_{\mathbf{Y}|\mathbf{X}}^{-1} \otimes \Sigma_{\mathbf{X}}) (\mathbf{K}_{r,p} (\mathbf{L} \eta \otimes \mathbf{R}_0)) \\
&+ 2\mathbf{A}^T (\Sigma_{\mathbf{Y}|\mathbf{X}}^{-1} \otimes \Sigma_{\mathbf{Y}|\mathbf{X}}) \mathbf{A} \\
&= \eta^T \Omega \eta \otimes \Phi_0^{-1} + \Phi \otimes \Phi_0^{-1} + \Phi^{-1} \otimes \Phi_0 - 2\mathbf{I}_{d_Y} \otimes \mathbf{I}_{r-d_Y} \\
&\equiv \mathbf{M}_3.
\end{aligned}$$

1140 Thus for the covariance,

$$\begin{aligned}
[\tilde{\delta}_3(\tilde{\delta}_3^T \mathbf{J}_h \tilde{\delta}_3)^\dagger \tilde{\delta}_3^T]_{\beta\beta} &= \mathbf{K}_{r,p} (\mathbf{L} \eta \otimes \mathbf{R}_0) (\mathbf{M}_3)^\dagger (\mathbf{K}_{r,p} (\mathbf{L} \eta \otimes \mathbf{R}_0))^T \\
&= (\mathbf{R}_0 \otimes \mathbf{L} \eta) \mathbf{K}_{r-d_Y, d_Y} (\mathbf{M}_3)^\dagger \mathbf{K}_{r-d_Y, d_Y}^T (\eta^T \mathbf{L}^T \otimes \mathbf{R}_0^T) \\
&= (\mathbf{R}_0 \otimes \mathbf{L} \eta) (\mathbf{K}_{r-d_Y, d_Y}^T \mathbf{M}_3 \mathbf{K}_{r-d_Y, d_Y})^\dagger (\eta^T \mathbf{L}^T \otimes \mathbf{R}_0^T)
\end{aligned}$$

1141 the last equality is because $\mathbf{K}_{r-d_Y, d_Y}$ is non-singular and $\mathbf{K}_{r-d_Y, d_Y} \mathbf{K}_{r-d_Y, d_Y}^T = \mathbf{I}_{(r-d_Y)d_Y}$.

1142 Therefore,

$$\begin{aligned}
[\tilde{\delta}_3(\tilde{\delta}_3^T \mathbf{J}_h \tilde{\delta}_3)^\dagger \tilde{\delta}_3^T]_{\beta\beta} &= (\mathbf{R}_0 \otimes \mathbf{L} \eta) (\mathbf{K}_{r-d_Y, d_Y}^T \mathbf{M}_3 \mathbf{K}_{r-d_Y, d_Y})^\dagger (\eta^T \mathbf{L}^T \otimes \mathbf{R}_0^T) \\
&= (\mathbf{R}_0 \otimes \mathbf{L} \eta) (\mathbf{M}_4)^\dagger (\eta^T \mathbf{L}^T \otimes \mathbf{R}_0^T),
\end{aligned}$$

1143 where

$$\mathbf{M}_4 = \Phi_0^{-1} \otimes \eta^T \Omega \eta + \Phi_0^{-1} \otimes \Phi + \Phi_0 \otimes \Phi^{-1} - 2\mathbf{I}_{r-d_Y} \otimes \mathbf{I}_{d_Y}. \quad (\text{I.4})$$

1144 I.4 Interpretations

1145 The Fisher information matrix for $\phi = (\phi_1^T, \dots, \phi_7^T)^T$, $\Delta^T \mathbf{J}_h \Delta$, has the form

$$\begin{pmatrix}
\mathbf{J}_\eta & \mathbf{J}_{\eta\mathbf{L}} & \mathbf{J}_{\eta\mathbf{R}} & 0 & 0 & 0 & 0 \\
\mathbf{J}_{\eta\mathbf{L}}^T & \mathbf{J}_{\mathbf{L}} & \mathbf{J}_{\mathbf{LR}} & \mathbf{J}_{\mathbf{L}\Omega} & 0 & 0 & 0 \\
\mathbf{J}_{\eta\mathbf{R}}^T & \mathbf{J}_{\mathbf{LR}}^T & \mathbf{J}_{\mathbf{R}} & 0 & 0 & \mathbf{J}_{\mathbf{R}\Phi} & 0 \\
0 & \mathbf{J}_{\mathbf{L}\Omega}^T & 0 & \mathbf{J}_{\Omega} & 0 & 0 & 0 \\
0 & 0 & 0 & 0 & \mathbf{J}_{\Omega_0} & 0 & 0 \\
0 & 0 & \mathbf{J}_{\mathbf{R}\Phi}^T & 0 & 0 & \mathbf{J}_{\Phi} & 0 \\
0 & 0 & 0 & 0 & 0 & 0 & \mathbf{J}_{\Phi_0}
\end{pmatrix}.$$

1146 If $\mathbf{L}$ and $\mathbf{R}$ are known, then the asymptotic variance of the MLE of $\text{vec}(\eta)$, denoted

1147 by $\text{vec}(\hat{\boldsymbol{\eta}}_{\mathbf{L},\mathbf{R}})$, is simply $\mathbf{J}_{\boldsymbol{\eta}}^{-1} = \boldsymbol{\Phi} \otimes \boldsymbol{\Omega}^{-1}$. As we discussed before,

$$\begin{aligned} \text{avar}(\sqrt{n}\text{vec}(\hat{\boldsymbol{\beta}}_{\mathbf{L},\mathbf{R}})) &= \text{avar}(\sqrt{n}\text{vec}(\mathbf{L}\hat{\boldsymbol{\eta}}_{\mathbf{L},\mathbf{R}}\mathbf{R}^T)) \\ &= (\mathbf{R} \otimes \mathbf{L})\text{avar}(\sqrt{n}\text{vec}(\hat{\boldsymbol{\eta}}_{\mathbf{L},\mathbf{R}}))(\mathbf{R}^T \otimes \mathbf{L}^T) \\ &= \mathbf{R}\boldsymbol{\Phi}\mathbf{R}^T \otimes \mathbf{L}\boldsymbol{\Omega}^{-1}\mathbf{L}^T. \end{aligned}$$

1148 **I.5 If $\boldsymbol{\eta}$ and $\mathbf{R}$ are known,**

1149 then the asymptotic variance of the MLE of $\text{vec}(\mathbf{L})$, denoted by $\text{vec}(\hat{\mathbf{L}}_{\boldsymbol{\eta},\mathbf{R}})$, is $(\mathbf{J}_{\mathbf{L}} - \mathbf{J}_{\mathbf{L}\boldsymbol{\Omega}}\mathbf{J}_{\boldsymbol{\Omega}}^{-1}\mathbf{J}_{\mathbf{L}\boldsymbol{\Omega}}^T)^{-1}$.

1150 Notice that

$$\begin{aligned} \mathbf{J}_{\mathbf{L}} - \mathbf{J}_{\mathbf{L}\boldsymbol{\Omega}}\mathbf{J}_{\boldsymbol{\Omega}}^{-1}\mathbf{J}_{\mathbf{L}\boldsymbol{\Omega}}^T &= \boldsymbol{\delta}_2^T \mathbf{J}_h \boldsymbol{\delta}_2 - \boldsymbol{\delta}_2^T \mathbf{J}_h \boldsymbol{\delta}_4 (\boldsymbol{\delta}_4^T \mathbf{J}_h \boldsymbol{\delta}_4)^\dagger \boldsymbol{\delta}_4^T \mathbf{J}_h \boldsymbol{\delta}_2 \\ &= \boldsymbol{\delta}_2^T \mathbf{J}_h^{1/2} \left(\mathbf{I} - \mathbf{P}_{\mathbf{J}_h^{1/2} \boldsymbol{\delta}_4} \right) \mathbf{J}_h^{1/2} \boldsymbol{\delta}_2 \\ &= \boldsymbol{\delta}_2^T \left(\mathbf{J}_h - \mathbf{J}_h \boldsymbol{\delta}_4 (\boldsymbol{\delta}_4^T \mathbf{J}_h \boldsymbol{\delta}_4)^\dagger \boldsymbol{\delta}_4^T \mathbf{J}_h \right) \boldsymbol{\delta}_2 \\ &\equiv \boldsymbol{\delta}_2^T (\mathbf{J}_h - \mathbf{J}_0) \boldsymbol{\delta}_2. \end{aligned}$$

1151 In the second equality above, $\mathbf{P}_{\mathbf{J}_h^{1/2} \boldsymbol{\delta}_4}$ is the projection matrix onto $\mathbf{J}_h^{1/2} \boldsymbol{\delta}_4$. But it

1152 not obvious how the projection might facilitate the calculation here. We look into the

1153 last equality, recall that $\boldsymbol{\delta}_4^T = (0, (\mathbf{C}_p(\mathbf{L} \otimes \mathbf{L})\mathbf{E}_{d_X})^T, 0)$ and that $\mathbf{J}_h$ was given in (5.1).

1154 Therefore,

$$\begin{aligned} \mathbf{J}_0 &= \mathbf{J}_h \boldsymbol{\delta}_4 (\boldsymbol{\delta}_4^T \mathbf{J}_h \boldsymbol{\delta}_4)^\dagger \boldsymbol{\delta}_4^T \mathbf{J}_h \\ &= \mathbf{J}_h \boldsymbol{\delta}_4 \left(\frac{1}{2} \mathbf{E}_{d_X}^T (\mathbf{L}^T \otimes \mathbf{L}^T) \mathbf{P}_{\mathbf{E}_p} (\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}) \mathbf{P}_{\mathbf{E}_{d_X}} (\mathbf{L} \otimes \mathbf{L}) \mathbf{E}_p \right)^\dagger \boldsymbol{\delta}_4^T \mathbf{J}_h \\ &= \mathbf{J}_h \boldsymbol{\delta}_4 \left(\frac{1}{2} \mathbf{E}_{d_X}^T (\mathbf{L}^T \otimes \mathbf{L}^T) (\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}) (\mathbf{L} \otimes \mathbf{L}) \mathbf{E}_{d_X} \right)^\dagger \boldsymbol{\delta}_4^T \mathbf{J}_h \\ &= \mathbf{J}_h \boldsymbol{\delta}_4 \left(\frac{1}{2} \mathbf{E}_{d_X}^T (\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}) \mathbf{E}_{d_X} \right)^\dagger \boldsymbol{\delta}_4^T \mathbf{J}_h. \end{aligned}$$

1155 The matrix $\boldsymbol{\delta}_4 \left(\frac{1}{2} \mathbf{E}_{d_X}^T (\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}) \mathbf{E}_{d_X} \right)^\dagger \boldsymbol{\delta}_4^T$ has 3×3 blocks. Only its middle $[\cdot]_{22}$ block

1156 is nonzero, other 8 blocks are all zeros. Then we compute the middle block as following,

1157 where we let $\mathbf{A}_3 = \mathbf{C}_p(\mathbf{L} \otimes \mathbf{L})\mathbf{E}_{d_X}$.

$$\left[\delta_4 \left(\frac{1}{2} \mathbf{E}_{d_X}^T (\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}) \mathbf{E}_{d_X} \right)^\dagger \delta_4^T \right]_{22} \quad (\text{I.5})$$

$$\begin{aligned} &= \mathbf{A}_3 \left(\frac{1}{2} \mathbf{E}_{d_X}^T (\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}) \mathbf{E}_{d_X} \right)^\dagger \mathbf{E}_{d_X}^T \mathbf{A}_3^T \\ &= 2\mathbf{A}_3 \mathbf{E}_{d_X} (\mathbf{E}_{d_X}^T (\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}) \mathbf{E}_{d_X})^\dagger \mathbf{E}_{d_X}^T \mathbf{A}_3^T \\ &= 2\mathbf{A}_3 \mathbf{P}_{\mathbf{E}_{d_X}} (\boldsymbol{\Omega} \otimes \boldsymbol{\Omega}) \mathbf{P}_{\mathbf{E}_{d_X}} \mathbf{A}_3^T \\ &= 2\mathbf{C}_p (\mathbf{L} \boldsymbol{\Omega} \mathbf{L}^T \otimes \mathbf{L} \boldsymbol{\Omega} \mathbf{L}^T) \mathbf{C}_p^T. \end{aligned} \quad (\text{I.6})$$

1158 Where the equation (I.6) is obtained using Corollary E.1. in CLC (2010; supplemen-
1159 tary material): $\mathbf{E}_{d_X} \in \mathbb{R}^{d_X^2 \times d_X(d_X+1)/2}$ is nonsingular and $\mathbf{P}_{\mathbf{E}_l}$ commute with $(\boldsymbol{\Omega} \otimes \boldsymbol{\Omega})$,
1160 then $\mathbf{E}_{d_X} (\mathbf{E}_{d_X}^T (\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}) \mathbf{E}_{d_X})^{-1} \mathbf{E}_{d_X}^T = \mathbf{P}_{\mathbf{E}_{d_X}} [(\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1})]^{-1} \mathbf{P}_{\mathbf{E}_{d_X}}$.

1161 Now, $\mathbf{J}_0 = \mathbf{J}_h \delta_4 (\delta_4^T \mathbf{J}_h \delta_4)^\dagger \delta_4^T \mathbf{J}_h$ only has its $[\cdot]_{22}$ block nonzero. For notational
1162 convenience, let $\mathbf{A} = \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}$ and $\mathbf{B} = \mathbf{L} \boldsymbol{\Omega} \mathbf{L}^T \otimes \mathbf{L} \boldsymbol{\Omega} \mathbf{L}^T$. Then

$$\begin{aligned} [\mathbf{J}_0]_{22} &= \left[\mathbf{J}_h \delta_4 (\delta_4^T \mathbf{J}_h \delta_4)^\dagger \delta_4^T \mathbf{J}_h \right]_{22} = 2^{-1} \mathbf{E}_p^T \mathbf{A} \mathbf{E}_p \times 2\mathbf{C}_p \mathbf{B} \mathbf{C}_p^T \times 2^{-1} \mathbf{E}_p^T \mathbf{A} \mathbf{E}_p \\ &= 2^{-1} \mathbf{E}_p^T \mathbf{A} \mathbf{P}_{\mathbf{E}_p} \mathbf{B} \mathbf{P}_{\mathbf{E}_p} \mathbf{A} \mathbf{E}_p = 2^{-1} \mathbf{E}_p^T \mathbf{A} \mathbf{B} \mathbf{A} \mathbf{E}_p = 2^{-1} \mathbf{E}_p^T \mathbf{B} \mathbf{E}_p. \end{aligned}$$

1163 Then, recall that $\delta_2 = \begin{pmatrix} \mathbf{R} \boldsymbol{\eta}^T \otimes \mathbf{I}_p \\ 2\mathbf{C}_p (\mathbf{L} \boldsymbol{\Omega} \otimes \mathbf{I}_p - \mathbf{L} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T) \\ 0 \end{pmatrix}$ and,

$$\begin{aligned} \mathbf{J}_L - \mathbf{J}_{L\boldsymbol{\Omega}} \mathbf{J}_{\boldsymbol{\Omega}}^{-1} \mathbf{J}_{L\boldsymbol{\Omega}}^T &= \delta_2^T \mathbf{J}_h \delta_2 - \delta_2^T \mathbf{J}_0 \delta_2 \\ &= (\boldsymbol{\eta} \mathbf{R}^T \otimes \mathbf{I}_p) (\boldsymbol{\Sigma}_{\mathbf{Y}|\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}) (\mathbf{R} \boldsymbol{\eta}^T \otimes \mathbf{I}_p) \\ &\quad + 2 (\boldsymbol{\Omega} \mathbf{L}^T \otimes \mathbf{I}_p - \mathbf{L}^T \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T) \mathbf{P}_{\mathbf{E}_p} \{ (\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}) - \\ &\quad (\mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T \otimes \mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T) \} \mathbf{P}_{\mathbf{E}_p} (\mathbf{L} \boldsymbol{\Omega} \otimes \mathbf{I}_p - \mathbf{L} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T) \\ &= \boldsymbol{\eta} \boldsymbol{\Phi}^{-1} \boldsymbol{\eta}^T \otimes \boldsymbol{\Sigma}_{\mathbf{X}} + 2\mathbf{A}^T \mathbf{P}_{\mathbf{E}_p} \mathbf{B} \mathbf{P}_{\mathbf{E}_p} \mathbf{A}. \end{aligned}$$

1164 The second term $2\mathbf{A}^T \mathbf{P}_{\mathbf{E}_p} \mathbf{B} \mathbf{P}_{\mathbf{E}_p} \mathbf{A}$ is defined by $\mathbf{A} = \mathbf{L} \boldsymbol{\Omega} \otimes \mathbf{I}_p - \mathbf{L} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T$, and

$$\begin{aligned} \mathbf{B} &= (\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}_{\mathbf{X}}^{-1}) - (\mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T \otimes \mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T) \\ &= \mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T + \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T \otimes \mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T + \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T. \end{aligned}$$

1165 Then we compute this term directly by noticing $\mathbf{P}_{\mathbf{E}_p} = \frac{1}{2}(\mathbf{I}_{p^2} + \mathbf{K}_{p,p})$ and $\mathbf{B}$ commutes
1166 with $\mathbf{P}_{\mathbf{E}_p}$, $2\mathbf{A}^T \mathbf{P}_{\mathbf{E}_p} \mathbf{B} \mathbf{P}_{\mathbf{E}_p} \mathbf{A} = 2\mathbf{A}^T \mathbf{B} \mathbf{P}_{\mathbf{E}_p} \mathbf{A} = \mathbf{A}^T \mathbf{B} (\mathbf{I}_{p^2} + \mathbf{K}_{p,p}) \mathbf{A}$. Then, notice that in

1167 $\mathbf{A}^T \mathbf{B}$ only the first term in $\mathbf{B}$ will not vanish. Hence,

$$\begin{aligned}\mathbf{A}^T \mathbf{B} &= (\mathbf{L} \boldsymbol{\Omega} \otimes \mathbf{I}_p - \mathbf{L} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T)^T \mathbf{L} \boldsymbol{\Omega}^{-1} \mathbf{L}^T \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T \\ &= \mathbf{L}^T \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T - \boldsymbol{\Omega}^{-1} \mathbf{L}^T \otimes \mathbf{L}_0 \mathbf{L}_0^T.\end{aligned}$$

1168 And,

$$\begin{aligned}\mathbf{A}^T \mathbf{B} (\mathbf{I}_{p^2} + \mathbf{K}_{p,p}) \mathbf{A} &= \mathbf{A}^T \mathbf{B} (\mathbf{L} \boldsymbol{\Omega} \otimes \mathbf{I}_p - \mathbf{L} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T) \\ &\quad + \mathbf{A}^T \mathbf{B} (\mathbf{I}_p \otimes \mathbf{L} \boldsymbol{\Omega} - \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T \otimes \mathbf{L}) \mathbf{K}_{l,p} \\ &= \mathbf{A}^T \mathbf{B} (\mathbf{L} \boldsymbol{\Omega} \otimes \mathbf{I}_p - \mathbf{L} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T) \\ &= \boldsymbol{\Omega} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T - \mathbf{I}_{d_X} \otimes \mathbf{L}_0 \mathbf{L}_0^T - \mathbf{I}_{d_X} \otimes \mathbf{L}_0 \mathbf{L}_0^T + \boldsymbol{\Omega}^{-1} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T \\ &= \boldsymbol{\Omega} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T - 2\mathbf{I}_{d_X} \otimes \mathbf{L}_0 \mathbf{L}_0^T + \boldsymbol{\Omega}^{-1} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T.\end{aligned}$$

1169 So far, we have found,

$$\begin{aligned}\mathbf{J}_L - \mathbf{J}_{L\Omega} \mathbf{J}_\Omega^{-1} \mathbf{J}_{L\Omega}^T &= \boldsymbol{\eta} \boldsymbol{\Phi}^{-1} \boldsymbol{\eta}^T \otimes \boldsymbol{\Sigma}_X + 2\mathbf{A}^T \mathbf{P}_{\mathbf{E}_p} \mathbf{B} \mathbf{P}_{\mathbf{E}_p} \mathbf{A} \\ &= \boldsymbol{\eta} \boldsymbol{\Phi}^{-1} \boldsymbol{\eta}^T \otimes \boldsymbol{\Sigma}_X + \mathbf{A}^T \mathbf{B} (\mathbf{I}_{p^2} + \mathbf{K}_{p,p}) \mathbf{A} \\ &= \boldsymbol{\eta} \boldsymbol{\Phi}^{-1} \boldsymbol{\eta}^T \otimes \boldsymbol{\Sigma}_X + \boldsymbol{\Omega} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0^{-1} \mathbf{L}_0^T - 2\mathbf{I}_{d_X} \otimes \mathbf{L}_0 \mathbf{L}_0^T + \boldsymbol{\Omega}^{-1} \otimes \mathbf{L}_0 \boldsymbol{\Omega}_0 \mathbf{L}_0^T.\end{aligned}$$

1170 Comparing this expression with (I.2) and (I.3), we have

$$\begin{aligned}\tilde{\boldsymbol{\delta}}_2^T \mathbf{J}_h \tilde{\boldsymbol{\delta}}_2 &= (\mathbf{I}_{d_X} \otimes \mathbf{L}_0^T) (\mathbf{J}_L - \mathbf{J}_{L\Omega} \mathbf{J}_\Omega^{-1} \mathbf{J}_{L\Omega}^T) (\mathbf{I}_{d_X} \otimes \mathbf{L}_0) \\ &= (\mathbf{I}_{d_X} \otimes \mathbf{L}_0^T) \left(\text{avar}(\sqrt{n} \text{vec}(\hat{\mathbf{L}}_{\boldsymbol{\eta}, \mathbf{R}})) \right)^\dagger (\mathbf{I}_{d_X} \otimes \mathbf{L}_0).\end{aligned}$$

1171 Consequently,

$$\begin{aligned}& [\tilde{\boldsymbol{\delta}}_2^T (\tilde{\boldsymbol{\delta}}_2^T \mathbf{J}_h \tilde{\boldsymbol{\delta}}_2)^\dagger \tilde{\boldsymbol{\delta}}_2^T]_{\beta\beta} \\ &= (\mathbf{R} \boldsymbol{\eta}^T \otimes \mathbf{L}_0) \left((\mathbf{I}_{d_X} \otimes \mathbf{L}_0^T) \left(\text{avar}(\sqrt{n} \text{vec}(\hat{\mathbf{L}}_{\boldsymbol{\eta}, \mathbf{R}})) \right)^\dagger (\mathbf{I}_{d_X} \otimes \mathbf{L}_0) \right)^\dagger (\boldsymbol{\eta} \mathbf{R}^T \otimes \mathbf{L}_0^T) \\ &= (\mathbf{R} \boldsymbol{\eta}^T \otimes \mathbf{L}_0 \mathbf{L}_0^T) \left(\text{avar}(\sqrt{n} \text{vec}(\hat{\mathbf{L}}_{\boldsymbol{\eta}, \mathbf{R}})) \right) (\boldsymbol{\eta} \mathbf{R}^T \otimes \mathbf{L}_0 \mathbf{L}_0^T) \quad (\text{I.7}) \\ &= \text{avar} \left(\sqrt{n} \text{vec}(\mathbf{Q}_L \hat{\mathbf{L}}_{\boldsymbol{\eta}, \mathbf{R}} \boldsymbol{\eta} \mathbf{R}^T) \right) = \text{avar} \left(\sqrt{n} \text{vec}(\mathbf{Q}_L \hat{\boldsymbol{\beta}}_{\boldsymbol{\eta}, \mathbf{R}}) \right).\end{aligned}$$

1172 where (I.7) can be obtained using Corollary E.1. in CLC (2010; supplementary ma-

1173 terial): $\mathbf{I}_{d_X} \otimes \mathbf{L}_0$ has full column rank, and its projection $\mathbf{I}_{d_X} \otimes \mathbf{L}_0 \mathbf{L}_0^T$ commute with

1174 $\mathbf{C} \equiv \left(\text{avar}(\sqrt{n} \text{vec}(\hat{\mathbf{L}}_{\boldsymbol{\eta}, \mathbf{R}})) \right)^\dagger = (\mathbf{J}_L - \mathbf{J}_{L\Omega} \mathbf{J}_\Omega^{-1} \mathbf{J}_{L\Omega}^T)$, then

$$\mathbf{I}_{d_X} \otimes \mathbf{L}_0 ((\mathbf{I}_{d_X} \otimes \mathbf{L}_0^T) \mathbf{C} (\mathbf{I}_{d_X} \otimes \mathbf{L}_0))^{-1} \mathbf{I}_{d_X} \otimes \mathbf{L}_0^T = (\mathbf{I}_{d_X} \otimes \mathbf{L}_0 \mathbf{L}_0^T) \mathbf{C}^{-1} (\mathbf{I}_{d_X} \otimes \mathbf{L}_0 \mathbf{L}_0^T).$$

1175 I.6 If η and $\mathbf{L}$ are known,

1176 then the asymptotic variance of the MLE of $\text{vec}(\mathbf{R})$, denoted by $\text{vec}(\hat{\mathbf{R}}_{\eta, \mathbf{L}})$, is $(\mathbf{J}_{\mathbf{R}} - \mathbf{J}_{\mathbf{R}\Phi} \mathbf{J}_{\Phi}^{-1} \mathbf{J}_{\mathbf{R}\Phi}^T)^{-1}$.

1177 Either by similar computation or by symmetry, we have

$$\begin{aligned} [\tilde{\delta}_3(\tilde{\delta}_3^T \mathbf{J}_h \tilde{\delta}_3)^\dagger \tilde{\delta}_3^T]_{\beta\beta} &= \text{avar} \left(\sqrt{n} \text{vec}(\hat{\beta}_{\eta, \mathbf{L}} \mathbf{Q}_{\mathbf{R}}) \right) = \text{avar} \left(\sqrt{n} \text{vec}(\mathbf{L} \eta \hat{\mathbf{R}}_{\eta, \mathbf{L}}^T \mathbf{Q}_{\mathbf{R}}) \right) \\ &= (\mathbf{R}_0 \mathbf{R}_0^T \otimes \mathbf{L} \eta) \text{avar} \left(\sqrt{n} \text{vec}(\hat{\mathbf{R}}_{\eta, \mathbf{L}}^T) \right) (\mathbf{R}_0 \mathbf{R}_0^T \otimes \eta^T \mathbf{L}^T). \end{aligned}$$

1178 J Additional simulations

1179 We used the same models as in Section 7.1. We simulated 1000 data sets for each setting
1180 and used 500 for training and the others for testing. We used the known dimension for
1181 each method: the true ranks of β were used for CCA and RRR, the $\mathbf{X}$ -envelope dimen-
1182 sions d_X were used for PLS, and the true envelope dimensions were used for envelope
1183 methods. The Frobenius norm $\|\alpha\|_F$ was used as a measure of overall prediction er-
1184 ror. And we also used the Frobenius norm $\|\beta - \hat{\beta}\|_F$ as a measure of overall parameter
1185 estimation accuracy.

1186 Simulation results were summarized in Table 4. We tried different dimensions and
1187 different sample sizes for all these models. As expected, the simultaneous envelope es-
1188 timators always had the best performance among all estimators and always performed
1189 substantially better than OLS. For each of the other methods, it could perform as well
1190 as the simultaneous envelope in some cases but could have much worse performances in
1191 other cases. We did not include the model (R2) because the true envelope dimensions
1192 were p and r , which gave no reductions. Interpretations of the results were similar as in
1193 Section 7.

1194 References

- 1195 [1] AMEMIYA, T. (1985), *Advanced Econometrics*, *Harvard University Press*.

Prediction Criterion $\ \alpha\ _F$									
(d_X, d_Y, N) or (d, N)	OLS	PLS	CCA	RRR	X-env.	Y-env.	S. env.	S. E.	$\leq$
E1	(3, 3, 100)	3.2	1.3	0.75	2.3	1.3	0.57	0.48	0.02
	(3, 3, 400)	0.70	0.32	0.14	0.37	0.31	0.11	0.10	0.004
	(3, 3, 900)	0.30	0.14	0.06	0.15	0.13	0.05	0.04	0.002
E2	(2, 5, 200)	1.6	0.55	2.8	1.6	0.51	0.76	0.28	0.03
	(5, 3, 200)	1.9	1.2	11	1.9	1.2	0.12	0.10	0.04*
	(7, 4, 200)	1.9	1.8	44	1.9	1.5	0.29	0.23	0.04*
R1	(1, 100)	3.8	1.5	1.83	2.3	2.2	1.7	1.4	0.03
	(3, 100)	3.8	2.5	3.9	2.7	2.8	2.7	2.6	0.03
	(5, 100)	3.8	3.0	4.7	3.3	3.3	3.2	3.1	0.04
Estimation Criterion $\ \beta - \hat{\beta}\ _F$									
(d_X, d_Y, N) or (d, N)	OLS	PLS	CCA	RRR	X-env.	Y-env.	S. env.	S. E.	$\leq$
E1	(3, 3, 100)	253	28	31	142	28	25	11	2.6
	(3, 3, 400)	14	2	1	3	2	1	1	0.1
	(3, 3, 900)	2.6	0.3	0.2	0.4	0.3	0.2	0.1	0.02
E2	(2, 5, 200)	451	23	116	425	127	197	59	11
	(5, 3, 200)	994	94	80	946	493	45	26	26
	(7, 4, 200)	1152	599	225	1112	876	154	104	36
R1	(1, 100)	73	30	36	46	42	34	27	0.7
	(3, 100)	73	49	76	53	54	52	50	0.8
	(5, 100)	73	58	90	64	64	62	60	0.8

Table 4: Prediction performances measured by $\|\alpha\|_F$ and estimation performances measured by $\|\beta - \hat{\beta}\|_F$, where * in the last column means the corresponding S.E. upper bound were computed without the S.E. of CCA. The best two methods for each setting were emphasized with boldface.