

Envelopes and reduced-rank regression

R. Dennis Cook^{*}, Liliana Forzani[†] and Xin Zhang[‡]

October 1, 2013

Abstract

We incorporate the nascent idea of envelopes (Cook, Li and Chiaromonte 2010) into reduced-rank regression, which is a popular techniques for dimension reduction and estimation in the multivariate linear model. We propose a reduced-rank envelope model, which is a hybrid of reduced-rank regression and envelope models. The resulting estimator is at least as efficient as both existing estimators and has total number of parameters no more than either of them. The methodology of this paper can be adapted easily to other envelope models such as partial envelopes (Su and Cook 2011) and envelopes in predictor space (Cook, Helland and Su 2013).

Key Words: Envelope model; Grassmannians; reduced-rank regression.

1 Introduction

The multivariate linear regression model for $p \times 1$ non-stochastic predictor $\mathbf{X}$ and $r \times 1$ stochastic response $\mathbf{Y}$ can be written as

$$\mathbf{Y} = \boldsymbol{\alpha} + \boldsymbol{\beta}\mathbf{X} + \boldsymbol{\epsilon}, \quad (1.1)$$

where the error vector $\boldsymbol{\epsilon}$ has mean zero and covariance matrix $\boldsymbol{\Sigma} > 0$ and is independent of $\mathbf{X}$. This model is a foundation of multivariate statistics where interest lies in prediction

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

[†]Liliana Forzani is Facultad de Ingeniera Quimica and Instituto de Matematica Aplicada (UNL-CONICET) Guemes 3450 - 3000 Santa Fe, Argentina (E-mail: liliana.forzani@gmail.com).

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

and in studying the interrelation between $\mathbf{X}$ and $\mathbf{Y}$ through the regression coefficient matrix $\beta \in \mathbb{R}^{r \times p}$. There is a general awareness that the estimation of β may often be improved by reducing the dimensionalities of $\mathbf{X}$ and $\mathbf{Y}$, and reduced-rank regression is popular method for doing so. We propose a reduced-rank envelope model that extends the nascent idea of envelopes to reduced-rank regression. The purpose of this paper is to integrate reduced-rank regression and envelopes, resulting in an overarching method that can choose the better of the two methods when appropriate and that has the potential to perform better than either of them.

Reduced-rank regression (Anderson 1951; Izenman 1975; Reinsel and Velu 1998) arises frequently in multivariate statistical analysis, and has been applied widely across the applied sciences. By restricting the rank of the regression coefficient matrix $\text{rank}(\beta) = d < \min(r, p)$, the total number of parameters is reduced and efficiency in estimation is improved. The analysis of reduced-rank regression (Izenman 1975; Tso 1981; Reinsel and Velu 1998; Anderson 2002) connects with many important multivariate methods such as principal components analysis, canonical correlation analysis and multiple time series modeling. The asymptotic advantages of the reduced-rank regression estimator over the standard ordinary least squares estimator were studied by Stoica and Viberg (1996) and Anderson (1999). Chen et al. (2012) and Chen and Huang (2012) extended reduced-rank regression to high-dimensional settings and demonstrated the advantages of parsimoniously reducing model parameters and interrelating response variables.

Envelope regression, which was first proposed by Cook et al. (2010), is another way of parsimoniously reducing the total number of parameters from the standard model (1.1) and gaining both efficiency in estimation and accuracy in prediction. The key idea of envelopes is to identify and eliminate information in the responses and the predictors that is immaterial to the estimation of β but still introduces unnecessary variation into estimation. Envelope reduction can be effective even when $d = \min(p, r)$, which is the case where reduced-rank regression has no reduction.

Envelope and reduced-rank regressions have different perspectives on dimension reduction.

It may take considerable effort to find which method is more efficient for a problem in practice. The proposed reduced-rank envelope model combines the strengths of envelopes and reduced-rank regression, which mitigates the burden of selecting among the two methods. When one of the two methods behaves poorly, the reduced-rank envelope model automatically degenerates towards the other one; when both methods show efficiency gains, the reduced-rank envelope estimator will enjoy a synergy from combining the two approaches and may improve over both estimators.

The rest of this paper is organized as follows. In Section 2, we review and summarize some fundamental results for reduced-rank regression and envelopes that are relevant to our development. We set up our reduced-rank envelope model in Section 2.3, where we also give intuitive connections to reduced-rank regression and envelope models. In Section 3.1, we summarize parameterizations for each model and show that the total number of parameters in the reduced-rank envelope model is fewer than that of the other models. Likelihood-based estimators for the reduced-rank envelope model are derived in Section 3.2. Asymptotic properties are studied in Section 4. We show that the reduced-rank envelope estimator is asymptotically more efficient than ordinary least squares, reduced-rank regression and envelope estimators under normal errors, and is still $\sqrt{n}$ -consistent without the normality assumption. Section 5 discusses procedures for selecting the rank of the coefficient matrix and the dimension of the envelope. Encouraging simulation results and real data examples are presented in Section 6 and 7. Proofs and other technical details are included in a Supplement to this article.

The following notations and definitions will be used in our exposition. Let $\mathbb{R}^{m \times n}$ be the set of all real $m \times n$ matrices. The Grassmannian consisting of the set of all u dimensional subspaces of $\mathbb{R}^r$, $u \leq r$, is denoted as $\mathcal{G}_{r,u}$. If $\mathbf{M} \in \mathbb{R}^{m \times n}$, then $\text{span}(\mathbf{M}) \subseteq \mathbb{R}^m$ is the subspace spanned by columns of $\mathbf{M}$. If $\sqrt{n}(\hat{\boldsymbol{\theta}} - \boldsymbol{\theta})$ converges to a normal random vector with mean 0 and covariance matrix $\mathbf{V}$ we write its asymptotic covariance matrix as $\text{avar}(\sqrt{n}\hat{\boldsymbol{\theta}}) = \mathbf{V}$. We use $\mathbf{P}_{\mathbf{A}(\mathbf{V})} = \mathbf{A}(\mathbf{A}^T \mathbf{V} \mathbf{A})^{-1} \mathbf{A}^T \mathbf{V}$ to denote the projection onto $\text{span}(\mathbf{A})$ with the $\mathbf{V}$ inner product and use $\mathbf{P}_{\mathbf{A}}$ to denote projection onto $\text{span}(\mathbf{A})$ with the identity inner product. Let

74 $\mathbf{Q}_{\mathbf{A}(\mathbf{V})} = \mathbf{I} - \mathbf{P}_{\mathbf{A}(\mathbf{V})}$. We will use operators $\text{vec} : \mathbb{R}^{a \times b} \rightarrow \mathbb{R}^{ab}$, which vectorizes an arbitrary
75 matrix by stacking its columns, and $\text{vech} : \mathbb{R}^{a \times a} \rightarrow \mathbb{R}^{a(a+1)/2}$, which vectorizes a symmetric
76 matrix by extracting its columns of elements below or on the diagonal. Let $\mathbf{A} \otimes \mathbf{B}$ denote
77 the Kronecker product of two matrices $\mathbf{A}$ and $\mathbf{B}$. We use $\hat{\boldsymbol{\theta}}_{\boldsymbol{\xi}}$ to denote estimator of $\boldsymbol{\theta}$ with
78 known true parameter value of $\boldsymbol{\xi}$. For a common parameter $\boldsymbol{\theta}$ in different models, we will use
79 subscripts to distinguish the estimators according to different models: $\hat{\boldsymbol{\theta}}_{\text{RR}}$ for the reduced-
80 rank regression estimator, $\hat{\boldsymbol{\theta}}_{\text{Env}}$ for the envelope estimator, $\hat{\boldsymbol{\theta}}_{\text{RE}}$ for the reduced-rank envelope
81 estimator and $\hat{\boldsymbol{\theta}}_{\text{OLS}}$ for the ordinary least square estimator.

82 2 Reduced-rank envelope model

83 2.1 Reduced-rank regression

84 Reduced-rank regression allows that $\text{rank}(\boldsymbol{\beta}) = d < \min(p, r)$ so that we can write the model
85 parameterization as

$$\boldsymbol{\beta} = \mathbf{A}\mathbf{B}, \mathbf{A} \in \mathbb{R}^{r \times d}, \mathbf{B} \in \mathbb{R}^{d \times p}, \text{rank}(\mathbf{A}) = \text{rank}(\mathbf{B}) = d, \quad (2.1)$$

86 where no additional constraints are imposed on $\mathbf{A}$ or $\mathbf{B}$. The maximum likelihood estimators
87 for the reduced-rank regression parameters were derived by Anderson (1999), Reinsel and Velu
88 (1998) and Stoica and Viberg (1996), under various constraints on $\mathbf{A}$ and $\mathbf{B}$ for identifiabil-
89 ity, such as $\mathbf{B}\boldsymbol{\Sigma}_{\mathbf{X}}\mathbf{B}^T = \mathbf{I}_d$ or $\mathbf{A}^T\mathbf{A} = \mathbf{I}_d$. The decomposition $\boldsymbol{\beta} = \mathbf{A}\mathbf{B}$ is still non-unique
90 even with those identifiable constraints: for any orthogonal matrix $\mathbf{O} \in \mathbb{R}^{d \times d}$, $\mathbf{A}_1 = \mathbf{A}\mathbf{O}$
91 and $\mathbf{B}_1 = \mathbf{O}^T\mathbf{B}$ offer another valid decomposition that satisfies the constraints. The param-
92 eters of interests, $\boldsymbol{\beta}$ and $\boldsymbol{\Sigma}$, are nevertheless identifiable as well as $\text{span}(\mathbf{A}) = \text{span}(\boldsymbol{\beta})$ and
93 $\text{span}(\mathbf{B}^T) = \text{span}(\boldsymbol{\beta}^T)$. We present this article in an apparently novel unified framework so
94 that every statement involving $\mathbf{A}$ or $\mathbf{B}$ holds universally for any decomposition $\boldsymbol{\beta} = \mathbf{A}\mathbf{B}$ satis-
95 fying (2.1).

The log-likelihood of model (1.1) under normality of ϵ can be written as,

$$L_n(\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\Sigma}) \simeq -\frac{n}{2} \left\{ \log |\boldsymbol{\Sigma}| + \frac{1}{n} \sum_{i=1}^n (\mathbf{Y}_i - \boldsymbol{\alpha} - \boldsymbol{\beta} \mathbf{X}_i)^T \boldsymbol{\Sigma}^{-1} (\mathbf{Y}_i - \boldsymbol{\alpha} - \boldsymbol{\beta} \mathbf{X}_i) \right\}, \quad (2.2)$$

which is to be maximized under the constraint that $\text{rank}(\boldsymbol{\beta}) = d$, or equivalently under the parameterization $\boldsymbol{\beta} = \mathbf{A}\mathbf{B}$. The symbol $\simeq$ denotes an equality from which any unimportant additive constant has been eliminated. We treat $L_n(\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\Sigma})$ as a general purpose objective function, which will be maximized under (2.1). The following lemma summarizes the reduced-rank regression estimator that maximizes (2.2). Rigorous derivation can be found in Anderson (1999).

Sample covariance matrices in this article are represented as $\mathbf{S}_{(\cdot)}$ and defined with the divisor n . For instance, $\mathbf{S}_{\mathbf{X}} = \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^T / n$, $\mathbf{S}_{\mathbf{XY}} = \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{Y}_i - \bar{\mathbf{Y}})^T / n$, $\mathbf{S}_{\mathbf{Y}|\mathbf{X}} = \mathbf{S}_{\mathbf{Y}} - \mathbf{S}_{\mathbf{YX}}\mathbf{S}_{\mathbf{X}}^{-1}\mathbf{S}_{\mathbf{XY}}$ denotes the sample covariance matrix of the residuals from the linear fit of $\mathbf{Y}$ on $\mathbf{X}$, and $\mathbf{S}_{\mathbf{Y} \circ \mathbf{X}} = \mathbf{S}_{\mathbf{YX}}\mathbf{S}_{\mathbf{X}}^{-1}\mathbf{S}_{\mathbf{XY}}$ denotes the sample covariance matrix of the fitted vectors from the linear fit of $\mathbf{Y}$ on $\mathbf{X}$. We define the sample canonical correlation matrix between $\mathbf{Y}$ and $\mathbf{X}$ as $\mathbf{C}_{\mathbf{YX}} = \mathbf{S}_{\mathbf{Y}}^{-1/2}\mathbf{S}_{\mathbf{YX}}\mathbf{S}_{\mathbf{X}}^{-1/2}$ and $\mathbf{C}_{\mathbf{XY}} = \mathbf{C}_{\mathbf{YX}}^T$. Truncated matrices are represented with superscripts. For example, $\mathbf{C}_{\mathbf{YX}}^{(d)}$ and $\mathbf{S}_{\mathbf{YX}}^{(d)}$ are constructed by truncated singular value decompositions of $\mathbf{C}_{\mathbf{YX}}$ and $\mathbf{S}_{\mathbf{YX}}$ with only the largest d singular values being kept.

Lemma 1. *Under the reduced-rank regression parameterization (2.1), the likelihood-based objective function from (2.2) is maximized at $\hat{\boldsymbol{\alpha}}_{\text{RR}} = \bar{\mathbf{Y}} - \hat{\boldsymbol{\beta}}_{\text{RR}}\bar{\mathbf{X}}$ and*

$$\begin{aligned} \hat{\boldsymbol{\beta}}_{\text{RR}} &= \mathbf{S}_{\mathbf{Y}}^{1/2} \mathbf{C}_{\mathbf{YX}}^{(d)} \mathbf{S}_{\mathbf{X}}^{-1/2} \\ \hat{\boldsymbol{\Sigma}}_{\text{RR}} &= \mathbf{S}_{\mathbf{Y}} - \hat{\boldsymbol{\beta}}_{\text{RR}} \mathbf{S}_{\mathbf{XY}} = \mathbf{S}_{\mathbf{Y}}^{1/2} \left\{ \mathbf{I}_r - \mathbf{C}_{\mathbf{YX}}^{(d)} \mathbf{C}_{\mathbf{XY}}^{(d)} \right\} \mathbf{S}_{\mathbf{Y}}^{1/2}. \end{aligned}$$

There are a variety forms of maximizers $\hat{\mathbf{A}}$ and $\hat{\mathbf{B}}$ in the literature under different constraints on $\mathbf{A}$ and $\mathbf{B}$. They could all be reproduced by decomposing the rank- d estimator $\hat{\boldsymbol{\beta}}_{\text{RR}}$ in Lemma 1. The ordinary least squares estimators for $\boldsymbol{\beta}$ and $\boldsymbol{\Sigma}$ can be written as $\hat{\boldsymbol{\beta}}_{\text{OLS}} = \mathbf{S}_{\mathbf{Y}}^{1/2} \mathbf{C}_{\mathbf{YX}} \mathbf{S}_{\mathbf{X}}^{-1/2}$ and $\hat{\boldsymbol{\Sigma}}_{\text{OLS}} = \mathbf{S}_{\mathbf{Y}|\mathbf{X}} = \mathbf{S}_{\mathbf{Y}}^{1/2} \{ \mathbf{I}_r - \mathbf{C}_{\mathbf{YX}} \mathbf{C}_{\mathbf{XY}} \} \mathbf{S}_{\mathbf{Y}}^{1/2}$ by replacing the

truncated sample canonical correlation matrices $\mathbf{C}_{(\cdot)}^{(d)}$ with the untruncated ones $\mathbf{C}_{(\cdot)}$. This Lemma also reveals the scale equivariant property of both reduced-rank regression and ordinary least squares estimators since the truncated sample canonical correlation matrices are scale invariant.

2.2 Review of Envelopes

The envelope model (Cook et al. 2010) seeks the smallest subspace $\mathcal{E} \subseteq \mathbb{R}^r$ such that

$$\mathbf{Q}_{\mathcal{E}}\mathbf{Y}|\mathbf{X} \sim \mathbf{Q}_{\mathcal{E}}\mathbf{Y} \text{ and } \text{cov}(\mathbf{Q}_{\mathcal{E}}\mathbf{Y}, \mathbf{P}_{\mathcal{E}}\mathbf{Y}|\mathbf{X}) = 0. \quad (2.3)$$

For any $\mathcal{E}$ with those properties, $\mathbf{Q}_{\mathcal{E}}\mathbf{Y}$ carries only information that is irrelevant to the linear regression. The projected response $\mathbf{Q}_{\mathcal{E}}\mathbf{Y}$ is linearly immaterial to the estimation of β in the sense that it responds to neither the predictor nor the rest of response $\mathbf{P}_{\mathcal{E}}\mathbf{Y}$, which presents material information in the response. When the conditional distribution of $\mathbf{Y}|\mathbf{X}$ is normal, the second statement in (2.3) implies that $\mathbf{P}_{\mathcal{E}}\mathbf{Y}$ is independent of $\mathbf{Q}_{\mathcal{E}}\mathbf{Y}$ given $\mathbf{X}$. The smallest subspace satisfying (2.3) always uniquely exists and is denoted by $\mathcal{E}_{\Sigma}(\beta)$ as defined formally in the following definitions.

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

The definition of a reducing subspace is basic in functional analysis (Conway 1990) but the notion of reduction is different from the common statistical meaning. Reducing subspaces are central to the study of envelope models and methods.

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

Definition 2 guarantees the existence and the uniqueness of envelopes by noticing that the intersection of any two reducing subspaces of $\mathbf{M}$ is still a reducing subspace of $\mathbf{M}$. To avoid proliferation of notation, we may use a matrix in the argument of an envelope as $\mathcal{E}_{\mathbf{M}}(\mathbf{B}) :=$

141 $\mathcal{E}_{\mathbf{M}}(\text{span}(\mathbf{B}))$. Under the reduced-rank regression model (2.1), $\mathcal{E}_{\Sigma}(\beta) = \mathcal{E}_{\Sigma}(\mathbf{A})$ and the di-
 142 mension of the envelope denoted by u is always no less than d since $\dim(\mathcal{E}_{\Sigma}(\beta)) \geq \dim(\text{span}(\beta)) =$
 143 $\text{rank}(\beta) = d$. The following proposition from Cook, et al. (2010) gives a characterization of
 144 envelopes.

145 **Proposition 1.** *If $\mathbf{M} \in \mathbb{R}^{r \times r}$ has $s < r$ eigenspaces, then the $\mathbf{M}$ -envelope of $\mathcal{S} \subseteq \text{span}(\mathbf{M})$
 146 can be constructed as $\mathcal{E}_{\mathbf{M}}(\mathcal{S}) = \sum_{i=1}^s \mathbf{P}_i \mathcal{S}$, where $\mathbf{P}_i$ is the projection onto the i -th eigenspace
 147 of $\mathbf{M}$.*

148 From this proposition, we see that the $\mathbf{M}$ -envelope of $\mathcal{S}$ is the sum of the eigenspaces of $\mathbf{M}$
 149 that are not orthogonal to $\mathcal{S}$. This implies that the envelope is the span of some subset of the
 150 eigenvectors of $\mathbf{M}$.

151 2.3 Reduced-rank envelope model

152 Let (Γ, Γ_0) be an orthogonal basis for $\mathbb{R}^r$ so that $\text{span}(\Gamma) = \mathcal{E}_{\Sigma}(\beta)$ and $\Gamma \in \mathbb{R}^{r \times u}$. Then
 153 $\dim(\mathcal{E}_{\Sigma}(\beta)) = u$ and

$$\beta = \mathbf{A}\mathbf{B} = \Gamma\xi = \Gamma\eta\mathbf{B}, \quad \Sigma = \Gamma\Omega\Gamma^T + \Gamma_0\Omega_0\Gamma_0^T, \quad (2.4)$$

154 where Ω and Ω_0 are symmetric positive definite matrices in $\mathbb{R}^{u \times u}$ and $\mathbb{R}^{(r-u) \times (r-u)}$ respectively
 155 and $\eta \in \mathbb{R}^{u \times d}$, $u \geq d$, are the coordinates of $\mathbf{A}$ with respect to Γ . The parameterization $\beta = \Gamma\xi$
 156 with $\xi \in \mathbb{R}^{u \times p}$ occurs in the envelope model of Cook et al. (2010). We still impose no addi-
 157 tional constraint on $\mathbf{A}$, $\mathbf{B}$ or η other than requiring them all to have rank d . The decompositions
 158 of β and Σ in (2.4) are not unique but β and Σ are unique.

159 To see the connections between the reduced-rank envelope model and reduced-rank regres-
 160 sion, we next consider the situation in which Γ is known. Notice that $\text{span}(\Gamma)$ is uniquely
 161 defined while Γ is unique up to an orthogonal transformation in $\mathbb{R}^u$. Although expressions
 162 in Lemma 2 are given in terms of Γ , the final estimators $\hat{\beta}_{\Gamma}$ and $\hat{\Sigma}_{\Gamma}$ depend on Γ only via
 163 $\text{span}(\Gamma)$: for any orthogonal transformation $\mathbf{O} \in \mathbb{R}^{u \times u}$, we have $\hat{\beta}_{\Gamma} = \hat{\beta}_{\Gamma\mathbf{O}}$ and $\hat{\Sigma}_{\Gamma} = \hat{\Sigma}_{\Gamma\mathbf{O}}$.

Lemma 2. Under the reduced-rank envelope model (2.4), the likelihood-based objective function from (2.2) with given Γ is maximized at $\hat{\alpha}_\Gamma = \bar{\mathbf{Y}} - \hat{\beta}_\Gamma \bar{\mathbf{X}}$ and

$$\begin{aligned}\hat{\beta}_\Gamma &= \Gamma \hat{\eta}_\Gamma \hat{\mathbf{B}}_\Gamma = \Gamma \mathbf{S}_{\Gamma^T \mathbf{Y}}^{1/2} \mathbf{C}_{\Gamma^T \mathbf{Y}, \mathbf{X}}^{(d)} \mathbf{S}_{\mathbf{X}}^{-1/2} \\ \hat{\Sigma}_\Gamma &= \Gamma \mathbf{S}_{\Gamma^T \mathbf{Y}}^{1/2} \left\{ \mathbf{I}_u - \mathbf{C}_{\Gamma^T \mathbf{Y}, \mathbf{X}}^{(d)} \mathbf{C}_{\mathbf{X}, \Gamma^T \mathbf{Y}}^{(d)} \right\} \mathbf{S}_{\Gamma^T \mathbf{Y}}^{1/2} \Gamma^T + \mathbf{Q}_\Gamma \mathbf{S}_{\mathbf{Y}} \mathbf{Q}_\Gamma.\end{aligned}$$

The implication of Lemma 2 is clear: once we know the envelope, we can focus our attention on the reduced response $\Gamma^T \mathbf{Y}$ and find $\hat{\eta}_\Gamma \hat{\mathbf{B}}_\Gamma$, which is the rank- d reduced-rank regression estimator of $\Gamma^T \mathbf{Y}$ on $\mathbf{X}$. By Definition 1, the covariance estimator $\hat{\Sigma}_\Gamma$ is now reduced by $\text{span}(\Gamma)$ since $\hat{\Sigma}_\Gamma = \mathbf{P}_\Gamma \hat{\Sigma}_\Gamma \mathbf{P}_\Gamma + \mathbf{Q}_\Gamma \hat{\Sigma}_\Gamma \mathbf{Q}_\Gamma$. Hence $\text{span}(\Gamma)$ is a reducing subspace of $\hat{\Sigma}_\Gamma$ that also contains $\text{span}(\hat{\beta}_\Gamma)$, and the envelope structure is preserved by the construction of these estimators. In Section 3.2, we derive the likelihood-based estimator $\hat{\Gamma}$ and demonstrate that the reduced-rank envelope estimators for β and Σ coincide with the estimators in Lemma 2 by replacing Γ with $\hat{\Gamma}$.

When the envelope dimension $u = r$, there is no immaterial information to be reduced by the envelope method. Then the reduced-rank envelope model degenerates to the reduced-rank regression (2.1), $\Gamma = \mathbf{I}_r$. When the regression coefficient matrix is full rank $\text{rank}(\beta) = p \leq r$, reduced-rank regression is equivalent to ordinary least squares and the reduced-rank envelope model degenerates to the ordinary envelope model. Two extreme situations are then: (a) if $p > r = 1$ then both methods degenerate to the standard method, which produces no reduction; (b) if $r > p = 1$ then reduced-rank regression can not provide any response reduction while reduced-rank envelopes can still gain efficiency by projecting the response onto the envelope $\mathcal{E}_\Sigma(\beta)$. The reduced-rank envelope model can be extended to the predictor envelopes by Cook et al. (2013), so that it can resolve the problem in (a) and provide potential gain by enveloping in the predictor space.

3 Likelihood-based estimation for reduced-rank envelope

3.1 Parameters in different models

Following Cook, Li and Chiaromonte (2010), we define the following estimable functions $\mathbf{h}$ for the standard model (1.1), parameters $\boldsymbol{\psi}$ for the reduced-rank model, parameters $\boldsymbol{\delta}$ for the envelope model and parameter ϕ for the reduced-rank envelope model. The common parameter $\boldsymbol{\alpha}$ is omitted because its estimator takes the following form for all methods: $\hat{\boldsymbol{\alpha}} = \bar{\mathbf{Y}} - \hat{\boldsymbol{\beta}}\bar{\mathbf{X}}$, while $\bar{\mathbf{Y}}$ and $\bar{\mathbf{X}}$ are asymptotically independent of the other estimators.

$$\mathbf{h} = \begin{pmatrix} \text{vec}(\boldsymbol{\beta}) \\ \text{vech}(\boldsymbol{\Sigma}) \end{pmatrix}, \boldsymbol{\psi} = \begin{pmatrix} \text{vec}(\mathbf{A}) \\ \text{vec}(\mathbf{B}) \\ \text{vech}(\boldsymbol{\Sigma}) \end{pmatrix}, \boldsymbol{\delta} = \begin{pmatrix} \text{vec}(\boldsymbol{\Gamma}) \\ \text{vec}(\boldsymbol{\xi}) \\ \text{vech}(\boldsymbol{\Omega}) \\ \text{vech}(\boldsymbol{\Omega}_0) \end{pmatrix}, \phi = \begin{pmatrix} \text{vec}(\boldsymbol{\Gamma}) \\ \text{vec}(\boldsymbol{\eta}) \\ \text{vec}(\mathbf{B}) \\ \text{vech}(\boldsymbol{\Omega}) \\ \text{vech}(\boldsymbol{\Omega}_0) \end{pmatrix}, \quad (3.1)$$

where we define $\mathbf{h} = (\mathbf{h}_1^T, \mathbf{h}_2^T)^T$, $\boldsymbol{\psi} = (\boldsymbol{\psi}_1^T, \boldsymbol{\psi}_2^T, \boldsymbol{\psi}_3^T)^T$, $\boldsymbol{\delta} = (\boldsymbol{\delta}_1^T, \dots, \boldsymbol{\delta}_4^T)^T$ and $\phi = (\phi_1, \dots, \phi_5)^T$ correspondingly. We have $\mathbf{h} = \mathbf{h}(\boldsymbol{\psi})$ under the reduced-rank model, $\mathbf{h} = \mathbf{h}(\boldsymbol{\delta})$ under the envelope model and $\mathbf{h} = \mathbf{h}(\phi)$ under the reduced-rank envelope model.

We use $\mathcal{N}(\cdot)$ to denote the total number of unique real parameters in a vector of model parameters. We have the following summary for each method:

(i) standard linear model, $\mathcal{N}_{\text{OLS}} := \mathcal{N}(\mathbf{h}) = pr + r(r + 1)/2$;

(ii) reduced-rank model, $\mathcal{N}_{\text{RR}} := \mathcal{N}(\boldsymbol{\psi}) = (p + r - d)d + r(r + 1)/2$;

(iii) envelope model, $\mathcal{N}_{\text{Env}} := \mathcal{N}(\boldsymbol{\delta}) = pu + r(r + 1)/2$;

(iv) reduced-rank envelope model, $\mathcal{N}_{\text{RE}} := \mathcal{N}(\phi) = (p + u - d)d + r(r + 1)/2$.

By straightforward calculation we observe that the total number of unique parameters is reduced by $(p - d)(r - d) \geq 0$ from standard model to reduced-rank regression, and is further reduced by $(r - u)d \geq 0$ from reduced-rank regression to reduced-rank envelopes. Similarly, the total number of unique parameters is reduced by $p(r - u) \geq 0$ from the standard model to envelopes,

and is further reduced by $(p - d)(u - d) \geq 0$ from the envelope model to the reduced-rank envelope model.

3.2 Estimators for the reduced-rank envelope model parameters

The goal of this section is to derive the reduced-rank envelope estimators for given d and u . Procedures for selecting d and u are discussed in Section 5. The likelihood-based reduced-rank envelope estimators is obtained by substituting $\mathbf{h} = \mathbf{h}(\phi)$ into (2.2) and maximizing $L_n(\boldsymbol{\alpha}, \boldsymbol{\beta}(\phi), \boldsymbol{\Sigma}(\phi)) \equiv L_n(\boldsymbol{\alpha}, \boldsymbol{\eta}, \mathbf{B}, \boldsymbol{\Omega}, \boldsymbol{\Omega}_0, \boldsymbol{\Gamma}|d, u)$ over all parameters except $\boldsymbol{\Gamma}$ because they live on a product space and the optimizing value of $\boldsymbol{\Gamma}$ cannot be found analytically. We then arrive at the estimator $\hat{\boldsymbol{\Gamma}}$ from optimization over a Grassmannian as described in the following Proposition. For any semi-orthogonal $r \times u$ matrix $\mathbf{G}$, we define $\mathbf{Z}_{\mathbf{G}} = (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}} \mathbf{G})^{-1/2} \mathbf{G}^T \mathbf{Y}$ to be the standardized version of $\mathbf{G}^T \mathbf{Y} \in \mathbb{R}^u$ with sample covariance $\mathbf{I}_u$, and let $\hat{\omega}_i(\mathbf{G})$, $i = 1, \dots, u$, be the i -th eigenvalue of $\mathbf{S}_{\mathbf{Z}_{\mathbf{G}}|\mathbf{X}}^{-1} = (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}|\mathbf{X}} \mathbf{G})^{-1/2} (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}} \mathbf{G}) (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}|\mathbf{X}} \mathbf{G})^{-1/2}$.

Proposition 2. *The estimator $\hat{\boldsymbol{\Gamma}} = \arg \min_{\mathbf{G}} F_n(\mathbf{G}|d, u)$ is the maximizer of $L_n(\boldsymbol{\alpha}, \boldsymbol{\eta}, \mathbf{B}, \boldsymbol{\Omega}, \boldsymbol{\Omega}_0, \boldsymbol{\Gamma}|d, u)$, where the optimization is over $\mathcal{G}_{r,u}$ and*

$$F_n(\mathbf{G}|d, u) = \log |\mathbf{G}^T \mathbf{S}_{\mathbf{Y}} \mathbf{G}| + \log |\mathbf{G}^T \mathbf{S}_{\mathbf{Y}}^{-1} \mathbf{G}| + \log |\mathbf{I}_u - \mathbf{S}_{\mathbf{Z}_{\mathbf{G}} \circ \mathbf{X}}^{(d)}| \quad (3.2)$$

$$= \log |\mathbf{G}^T \mathbf{S}_{\mathbf{Y}|\mathbf{X}} \mathbf{G}| + \log |\mathbf{G}^T \mathbf{S}_{\mathbf{Y}}^{-1} \mathbf{G}| + \sum_{i=d+1}^u \log [\hat{\omega}_i(\mathbf{G})]. \quad (3.3)$$

We find in practice that the form of objective function (3.3) can be more easily and stably evaluated than (3.2). The analytical expression of $\partial F_n(\mathbf{G}|d, u)/\partial \mathbf{G}$ based on (3.3) is used to facilitate the Newton-Raphson or conjugate gradient iterations. The formulation in (3.2) describes some operating characteristics of the reduced-rank envelope objective function. Lemma 1 and the relationship $\mathbf{S}_{\mathbf{Z}_{\mathbf{G}} \circ \mathbf{X}}^{(d)} = \mathbf{C}_{\mathbf{Z}_{\mathbf{G}} \mathbf{X}}^{(d)} \mathbf{C}_{\mathbf{X} \mathbf{Z}_{\mathbf{G}}}^{(d)}$ implies that the term $\mathbf{I}_u - \mathbf{S}_{\mathbf{Z}_{\mathbf{G}} \circ \mathbf{X}}^{(d)}$ equals the sample covariance of the residuals from reduced-rank regression fit of $\mathbf{Z}_{\mathbf{G}}$ on $\mathbf{X}$ with rank d . Let $F_{n,1}(\mathbf{G}|u) = \log |\mathbf{G}^T \mathbf{S}_{\mathbf{Y}} \mathbf{G}| + \log |\mathbf{G}^T \mathbf{S}_{\mathbf{Y}}^{-1} \mathbf{G}|$ and $F_{n,2}(\mathbf{G}|d, u) = \log |\mathbf{I}_u - \mathbf{S}_{\mathbf{Z}_{\mathbf{G}} \circ \mathbf{X}}^{(d)}|$ so that $F_n(\mathbf{G}|d, u) = F_{n,1}(\mathbf{G}|u) + F_{n,2}(\mathbf{G}|d, u)$. The first part $F_{n,1}(\mathbf{G}|u) \geq 0$ for all $\mathbf{G} \in \mathcal{G}_{r,u}$ and equals zero when $\mathbf{G}$ is a u -dimensional reducing subspace of $\mathbf{S}_{\mathbf{Y}}$. The effect of $F_{n,1}(\mathbf{G}|u)$ is

then to pull the solution towards eigenvectors of $\mathbf{S}_Y$. The second part $F_{n,2}(\mathbf{G}|d, u)$ represents the magnitude of the sample covariance of the residual from reduced-rank regression fit of the standardized variable $\mathbf{Z}_G$ on $\mathbf{X}$ with given rank d . Simply put, this part is a scale-invariant measure for the lack-of-fit of the rank- d reduced-rank regression of $\mathbf{G}^T \mathbf{Y}$ on $\mathbf{X}$.

Our formulation and decomposition based on (3.2) offer a generic way of interpreting the likelihood-based objective functions for envelope methods. For example, the objective function for the standard envelope model in Cook et al. (2010) can be expressed as

$$\log |\mathbf{G}^T \mathbf{S}_Y \mathbf{G}| + \log |\mathbf{G}^T \mathbf{S}_Y^{-1} \mathbf{G}| + \log |\mathbf{I}_u - \mathbf{S}_{\mathbf{Z}_G \circ \mathbf{X}}|, \quad (3.4)$$

which can be interpreted similar to (3.2) except the lack-of-fit term is now based on ordinary least squares fit rather than reduced-rank regression fit. The above objective function is the same as (3.2) when $d = p$ or $d = u$.

Additional properties of the objective function are given in the following Proposition.

Proposition 3. *The objective function $F_n(\mathbf{G}|d, u)$ in (3.3) converges in probability as $n \rightarrow \infty$ to the population objective function $F(\mathbf{G}|u) = \log |\mathbf{G}^T \Sigma \mathbf{G}| + \log |\mathbf{G}^T \Sigma_Y^{-1} \mathbf{G}|$ uniformly in $\mathbf{G}$. The estimator $\hat{\Gamma} = \arg \min_{\mathbf{G}} F_n(\mathbf{G}|d, u)$ is Fisher consistent, $\mathcal{E}_\Sigma(\beta) = \text{span}\{\arg \min_{\mathbf{G}} F(\mathbf{G}|u)\}$.*

The population objective function $F(\mathbf{G}|u)$, which does not depend explicitly on the given rank d , is exactly the same one as in Cook et al. (2010) for estimating an u -dimensional envelope $\mathcal{E}_\Sigma(\beta)$. In the proof of Proposition 3, we show that $\log[\hat{\omega}_i(\mathbf{G})]$, for any $i > d$, converges in probability to zero uniformly in $\mathbf{G}$. Therefore, we could view $F_n(\mathbf{G}|d, u)$ in (3.3) as a sample version of $F(\mathbf{G}|u)$, $F_n(\mathbf{G}|u) := \log |\mathbf{G}^T \mathbf{S}_{Y|\mathbf{X}} \mathbf{G}| + \log |\mathbf{G}^T \mathbf{S}_Y^{-1} \mathbf{G}|$, plus a finite sample adjustment for the rank deficiency, $\sum_{i=d+1}^u \log[\hat{\omega}_i(\mathbf{G})]$, which goes to zero as $n \rightarrow \infty$. Minimizing $F_n(\mathbf{G}|u)$ leads to another $\sqrt{n}$ -consistent envelope estimator but it will not be optimal since it does not account for the rank deficiency. The impact of the rank $d < p$ on the envelope estimation diminishes as sample size increases and reduced-rank envelope estimation moves towards a two-stage estimation procedure: first estimate the envelope from $F_n(\mathbf{G}|u)$ ignoring the rank, then obtain a rank- d estimator within the estimated envelope. The effects of rank deficiency

and envelope interdigitate at finite samples and there is a noticeable synergy when sample size is not large.

Finally, we summarize estimators for the parameters in the reduced-rank envelope model as follows. The results come naturally from Lemma 2.

Proposition 4. *The estimators for the reduced-rank envelope model (2.4) that minimize (2.2) are $\hat{\alpha}_{\text{RE}} = \bar{\mathbf{Y}} - \hat{\beta}_{\text{RE}} \bar{\mathbf{X}}$, $\hat{\Gamma} = \arg \max_{\mathbf{G} \in \mathcal{G}_{r,u}} F_n(\mathbf{G}|d, u)$, $\hat{\Omega}_0 = \hat{\Gamma}_0^T \mathbf{S}_{\mathbf{Y}} \hat{\Gamma}_0$ and*

$$\begin{aligned}\hat{\Omega} &= \mathbf{S}_{\hat{\Gamma}^T \mathbf{Y}}^{1/2} \left\{ \mathbf{I}_u - \mathbf{C}_{\hat{\Gamma}^T \mathbf{Y}, \mathbf{X}}^{(d)} \mathbf{C}_{\mathbf{X}, \hat{\Gamma}^T \mathbf{Y}}^{(d)} \right\} \mathbf{S}_{\hat{\Gamma}^T \mathbf{Y}}^{1/2} \\ \hat{\Sigma}_{\text{RE}} &= \hat{\Gamma} \hat{\Omega} \hat{\Gamma}^T + \hat{\Gamma}_0 \hat{\Omega}_0 \hat{\Gamma}_0^T \\ \hat{\beta}_{\text{RE}} &= \hat{\Gamma} \hat{\eta} \hat{\mathbf{B}}_{\text{RE}} = \hat{\Gamma} \mathbf{S}_{\hat{\Gamma}^T \mathbf{Y}}^{1/2} \mathbf{C}_{\hat{\Gamma}^T \mathbf{Y}, \mathbf{X}}^{(d)} \mathbf{S}_{\mathbf{X}}^{-1/2}.\end{aligned}$$

The rank of $\hat{\beta}_{\text{RE}}$ is d and the span of $\hat{\beta}_{\text{RE}}$ is a subset of the entire u -dimensional envelope. In contrast to reduced-rank regression, the estimator for $\hat{\Sigma}_{\text{RE}}$ now has an envelope structure:

$$\hat{\Sigma}_{\text{RE}} = \mathbf{P}_{\hat{\Gamma}} \hat{\Sigma}_{\text{RE}} \mathbf{P}_{\hat{\Gamma}} + \mathbf{Q}_{\hat{\Gamma}} \hat{\Sigma}_{\text{RE}} \mathbf{Q}_{\hat{\Gamma}}.$$

If we let $u = r$, which is equivalent to setting $\Gamma = \mathbf{I}_r$ in Proposition 4, then there is no envelope reduction and the estimator $\hat{\beta}_{\text{RE}}$ is the same as the estimator $\hat{\beta}_{\text{RR}}$ in Lemma 1. If we let $d = p$, then the estimators in Proposition 4 is the same as the envelope estimators in Cook et al. (2010). The estimators for the reduced-rank envelope model parameters coincide with those estimators in Lemma 2 by replacing Γ by its estimator $\hat{\Gamma}$.

4 Asymptotics

4.1 Asymptotic properties under normality

In this section, we present asymptotic results assuming that the error term is normal, $\epsilon \sim N(0, \Sigma)$, so that the estimators derived in Section 3 are all maximum likelihood estimators. We focus attention on the comparison between $\hat{\beta}_{\text{RE}}$ and $\hat{\beta}_{\text{RR}}$ because (1) comparisons between $\hat{\beta}_{\text{Env}}$ and $\hat{\beta}_{\text{OLS}}$ can be found in Cook et al. (2010); and (2) the advantage of $\hat{\beta}_{\text{RE}}$ over $\hat{\beta}_{\text{Env}}$ is

similar to the advantage of $\hat{\beta}_{\text{RR}}$ over $\hat{\beta}_{\text{OLS}}$, which is due to the rank reduction in the material response $\Gamma^T \mathbf{Y}$. We then relax the normality assumption in Section 4.2 and show the $\sqrt{n}$ -consistency of the reduced-rank envelope estimator and its asymptotic distribution.

From Cook et al. (2010) we know that the Fisher information for $\mathbf{h}$ is

$$\mathbf{J}_{\mathbf{h}} = \begin{pmatrix} \mathbf{J}_{\beta} & 0 \\ 0 & \mathbf{J}_{\Sigma} \end{pmatrix} = \begin{pmatrix} \Sigma_{\mathbf{X}} \otimes \Sigma^{-1} & 0 \\ 0 & \frac{1}{2} \mathbf{E}_r^T (\Sigma^{-1} \otimes \Sigma^{-1}) \mathbf{E}_r \end{pmatrix}, \quad (4.1)$$

where $\Sigma_{\mathbf{X}} = \lim_{n \rightarrow \infty} \mathbf{S}_{\mathbf{X}}$ and $\mathbf{E}_r$ is the expansion matrix, $\mathbf{E}_r \text{vec}(\mathbf{S}) = \text{vech}(\mathbf{S})$ for any $r \times r$ symmetric matrix $\mathbf{S}$. The asymptotic covariance for the ordinary least squares estimator $\hat{\mathbf{h}}_{\text{OLS}}$ is $\mathbf{J}_{\mathbf{h}}^{-1}$, which is the asymptotic covariance of the unrestricted maximum likelihood estimator.

Define the gradient matrices

$$\mathbf{H} = \frac{\partial \mathbf{h}(\psi)}{\partial \psi} \text{ and } \mathbf{R} = \frac{\partial \mathbf{h}(\phi)}{\partial \phi}. \quad (4.2)$$

Then the asymptotic covariance for the reduced-rank regression estimator $\hat{\mathbf{h}}_{\text{RR}} = \mathbf{h}(\hat{\psi})$ and for the reduced-rank envelope estimator $\hat{\mathbf{h}}_{\text{RE}} = \mathbf{h}(\hat{\phi})$ are summarized in the following Proposition.

Proposition 5. *Assuming that $\epsilon \sim N(0, \Sigma)$, then $\text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{OLS}}) = \mathbf{J}_{\mathbf{h}}^{-1}$, $\text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{RR}}) = \mathbf{H}(\mathbf{H}^T \mathbf{J}_{\mathbf{h}} \mathbf{H})^{\dagger} \mathbf{H}^T$ and $\text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{RE}}) = \mathbf{R}(\mathbf{R}^T \mathbf{J}_{\mathbf{h}} \mathbf{R})^{\dagger} \mathbf{R}^T$. Moreover,*

$$\begin{aligned} \text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{OLS}}) - \text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{RR}}) &= \mathbf{J}_{\mathbf{h}}^{-1/2} \mathbf{Q}_{\mathbf{J}_{\mathbf{h}}^{1/2} \mathbf{H}} \mathbf{J}_{\mathbf{h}}^{-1/2} \geq 0, \\ \text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{RR}}) - \text{avar}(\sqrt{n}\hat{\mathbf{h}}_{\text{RE}}) &= \mathbf{J}_{\mathbf{h}}^{-1/2} \left(\mathbf{P}_{\mathbf{J}_{\mathbf{h}}^{1/2} \mathbf{H}} - \mathbf{P}_{\mathbf{J}_{\mathbf{h}}^{1/2} \mathbf{R}} \right) \mathbf{J}_{\mathbf{h}}^{-1/2} \\ &= \mathbf{J}_{\mathbf{h}}^{-1/2} \mathbf{P}_{\mathbf{J}_{\mathbf{h}}^{1/2} \mathbf{H}} \mathbf{Q}_{\mathbf{J}_{\mathbf{h}}^{1/2} \mathbf{R}} \mathbf{J}_{\mathbf{h}}^{-1/2} \geq 0, \end{aligned}$$

where $\dagger$ indicates the Moore-Penrose inverse. In particular, $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{OLS}})] \geq \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RR}})] \geq \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}})]$.

Proposition 5 follows directly from $\psi = \psi(\phi)$. Therefore, we have $\mathbf{R} = \mathbf{H} \partial \psi(\phi) / \partial \phi$ and $\text{span}(\mathbf{J}_{\mathbf{h}}^{1/2} \mathbf{R}) \subseteq \text{span}(\mathbf{J}_{\mathbf{h}}^{1/2} \mathbf{H})$. Similarly, it can be shown that $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{OLS}})] \geq \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{Env}})] \geq \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}})]$.

Since we are particularly interested in the asymptotic covariance of $\hat{\mathbf{h}}_1 = \text{vec}(\hat{\beta})$ from different estimators, we summarize some of the results in the following Propositions.

291 **Proposition 6.** Assume that $\epsilon \sim N(0, \Sigma)$ and that $\text{rank}(\beta) = d$. Then $\sqrt{n}\text{vec}(\hat{\beta}_{\text{OLS}} - \beta)$ and
 292 $\sqrt{n}\text{vec}(\hat{\beta}_{\text{RR}} - \beta)$ are both asymptotically normal with mean zero and covariances as follows.

$$\begin{aligned} \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{OLS}})] &= \Sigma_{\mathbf{X}}^{-1} \otimes \Sigma \\ \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RR}})] &= (\mathbf{I}_{pr} - \mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})} \otimes \mathbf{Q}_{\mathbf{A}(\Sigma^{-1})}) \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{OLS}})] \end{aligned} \quad (4.3)$$

$$= \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\mathbf{A}} \mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})}^T)] + \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\mathbf{B}})] \quad (4.4)$$

293 where $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\mathbf{A}})] = \Sigma_{\mathbf{X}}^{-1} \otimes (\mathbf{P}_{\mathbf{A}(\Sigma^{-1})} \Sigma)$ and $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\mathbf{B}})] = (\mathbf{P}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})} \Sigma_{\mathbf{X}}^{-1}) \otimes \Sigma$.

294 The asymptotic result in (4.3) follows from Anderson (1999; equation (3.20)). The results
 295 in Proposition 6 rely on $\mathbf{A}$ and $\mathbf{B}$ only through their projections $\mathbf{Q}_{\mathbf{A}(\Sigma^{-1})}$ and $\mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})}$, which
 296 serve to orthogonalize the parameters in the asymptotic variance decompositions. This implies
 297 that all the equalities in Proposition 6 hold for any decomposition $\beta = \mathbf{A}\mathbf{B}$, with $\mathbf{A} \in \mathbb{R}^{r \times d}$
 298 and $\mathbf{B} \in \mathbb{R}^{d \times p}$. Hence, Proposition 6 is a unification for all the asymptotic studies of reduced-
 299 rank regression in the literature such as Anderson (1999), Reinsel and Velu (1998), Stoica and
 300 Viberg (1996) and so on.

301 For the reduced-rank envelope model (2.4), we have the following results on asymptotic
 302 distributions.

303 **Proposition 7.** Under the reduced-rank envelope model with normal error $\epsilon \sim N(0, \Sigma)$,
 304 $\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}} - \beta)$ is asymptotically normal with mean zero and covariance

$$\begin{aligned} \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}})] &= \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma})] + \text{avar}[\sqrt{n}\text{vec}(\mathbf{Q}_{\Gamma} \hat{\beta}_{\eta, \mathbf{B}})] \\ &= \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma, \eta} \mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})}^T)] + \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma, \mathbf{B}})] \\ &\quad + \text{avar}[\sqrt{n}\text{vec}(\mathbf{Q}_{\Gamma} \hat{\beta}_{\eta, \mathbf{B}})], \end{aligned} \quad (4.5)$$

305 where $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma, \eta} \mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})}^T)] = \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\mathbf{A}} \mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})}^T)]$ from (4.4). Explicit expres-
 306 sions for $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}})]$ can be found in the Supplemental material (D.7). The above equal-
 307 ities hold for any decomposition $\beta = \Gamma \eta \mathbf{B}$, where Γ is semi-orthogonal and the dimensions of
 308 Γ , η and $\mathbf{B}$ are $r \times u$, $u \times d$ and $d \times p$.

We view the asymptotic advantages of reduced-rank envelopes over reduced-rank regression by contrasting (4.4) with (4.5). From Propositions 6 and 7, we can write $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RR}})] - \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}})]$ as

$$\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{B}})] - \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma, \text{B}})] - \text{avar}[\sqrt{n}\text{vec}(\mathbf{Q}_{\Gamma}\hat{\beta}_{\eta, \text{B}})] \geq 0, \quad (4.6)$$

where $\hat{\beta}_{\text{B}} = \hat{\mathbf{A}}_{\text{B}}\mathbf{B}$, $\hat{\beta}_{\Gamma, \text{B}} = \Gamma\hat{\eta}_{\Gamma, \text{B}}\mathbf{B} = \hat{\mathbf{A}}_{\Gamma, \text{B}}\mathbf{B}$ and $\hat{\beta}_{\eta, \text{B}} = \hat{\Gamma}_{\eta, \text{B}}\eta\mathbf{B} = \hat{\mathbf{A}}_{\eta, \text{B}}\mathbf{B}$ are estimators with given $\mathbf{B}$. When $\mathbf{B}$ is known, the original regression problem simplifies to the regression of $\mathbf{Y}$ on $\mathbf{BX}$ and $\mathbf{A}$ is the new regression coefficient matrix. The estimator $\hat{\mathbf{A}}_{\text{B}}$ is the ordinary least squares estimator of $\mathbf{Y}$ on $\mathbf{BX}$ and the estimators $\hat{\mathbf{A}}_{\Gamma, \text{B}}$ and $\hat{\mathbf{A}}_{\eta, \text{B}}$ correspond to the usual envelope estimators for $\mathbf{A} = \Gamma\eta$. The difference in asymptotic covariances $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RR}})] - \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}})]$ from (4.6) equals the asymptotic efficiency gain of envelope estimator over the ordinary least squares estimator for regression of $\mathbf{Y}$ on $\mathbf{BX}$ and is consistent with the results presented in Cook et al. (2010).

Two special situations where the inequality in (4.6) becomes equality are: $\Gamma = \mathbf{I}_r$ and $\Sigma = \sigma^2\mathbf{I}_r$, while the envelope estimator is asymptotically equivalent to the ordinary least squares estimator in these two cases.

To see the potential gain of the reduced-rank envelope estimator, we have the following Corollary, where we have ignored the cost of estimating an envelope.

Corollary 1. *Under the reduced-rank envelope model with normal error $\epsilon \sim N(0, \Sigma)$,*

$$\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma})] = \mathbf{F}_1\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{Env}, \Gamma})] = \mathbf{F}_2\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RR}})] = \mathbf{F}_1\mathbf{F}_2\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{OLS}})],$$

where $\mathbf{F}_1 = \mathbf{I}_{pr} - \mathbf{Q}_{\text{B}^T(\Sigma_{\text{X}})} \otimes \mathbf{Q}_{\mathbf{A}(\Sigma^{-1})}$ and $\mathbf{F}_2 = \mathbf{I}_p \otimes \mathbf{P}_{\Gamma}$ are two positive semi-definite matrices with eigenvalues between 0 and 1.

The two matrices $\mathbf{F}_1$ and $\mathbf{F}_2$ represent fractions of asymptotic covariance reduction from the ordinary least squares estimator to the reduced-rank regression estimator and to the envelope estimator with given Γ . Then the efficiency gain of reduced-rank envelope with known Γ over ordinary least squares is the superimposition of the efficiency gain of the reduced-rank regression and the envelope regression with known Γ .

4.2 Consistency without the normality assumption

Let $\hat{\mathbf{h}}_{\text{OLS}} = \left(\text{vec}^T(\hat{\boldsymbol{\beta}}_{\text{OLS}}), \text{vech}^T(\mathbf{S}_{\mathbf{Y}|\mathbf{X}}) \right)^T$ denote the ordinary least squares estimator of $\mathbf{h}$ under the standard linear regression model, and let $\hat{\mathbf{h}}_{\text{RE}} = \mathbf{h}(\hat{\boldsymbol{\phi}})$ denote the reduced-rank envelope estimator. The true values of $\mathbf{h}$ and $\boldsymbol{\phi}$ are denoted as $\mathbf{h}_0$ and $\boldsymbol{\phi}_0$. The objective function $L_n(\boldsymbol{\alpha}, \boldsymbol{\beta}, \boldsymbol{\Sigma})$ in (2.2) can be written as, after partially maximized over $\boldsymbol{\alpha}$,

$$L_n(\boldsymbol{\beta}, \boldsymbol{\Sigma}) \simeq -\frac{n}{2} \log |\boldsymbol{\Sigma}| - \frac{n}{2} \text{trace} \left\{ \boldsymbol{\Sigma}^{-1} [\mathbf{S}_{\mathbf{Y}|\mathbf{X}} + (\hat{\boldsymbol{\beta}}_{\text{OLS}} - \boldsymbol{\beta}) \mathbf{S}_{\mathbf{X}} (\hat{\boldsymbol{\beta}}_{\text{OLS}} - \boldsymbol{\beta})^T] \right\}. \quad (4.7)$$

We treat the objective function $L_n(\boldsymbol{\beta}, \boldsymbol{\Sigma})$ as a function of $\mathbf{h}$ and $\hat{\mathbf{h}}_{\text{OLS}}$ and define $\mathcal{F}(\mathbf{h}, \hat{\mathbf{h}}_{\text{OLS}}) = 2/n \left\{ L_n(\hat{\boldsymbol{\beta}}_{\text{OLS}}, \mathbf{S}_{\mathbf{Y}|\mathbf{X}}) - L_n(\boldsymbol{\beta}, \boldsymbol{\Sigma}) \right\}$, which satisfies the conditions of Shapiro's (1986) minimum discrepancy function (see Supplement Section F). Hence $\mathbf{J}_{\mathbf{h}} = 1/2 \times \partial^2 \mathcal{F}(\mathbf{h}, \hat{\mathbf{h}}_{\text{OLS}}) / \partial \mathbf{h} \partial \mathbf{h}^T$ evaluated at $\hat{\mathbf{h}}_{\text{OLS}} = \mathbf{h} = \mathbf{h}_0$ is the Fisher information matrix for $\mathbf{h}$ when ϵ is normal. The following proposition formally states the asymptotic distribution of $\hat{\mathbf{h}}_{\text{RE}}$ without normality of ϵ .

Proposition 8. *Assume that the reduced-rank envelope model (2.4) holds and that ϵ_i 's are independent and identically distributed with finite fourth moments. Then $\sqrt{n}(\hat{\mathbf{h}}_{\text{OLS}} - \mathbf{h}_0) \rightarrow N(0, \mathbf{K})$, for some positive definite covariance matrix $\mathbf{K}$. And $\sqrt{n}(\hat{\mathbf{h}}_{\text{RE}} - \mathbf{h}_0)$ converges in distribution to a normal random variable with mean $\mathbf{0}$ and covariance matrix*

$$\mathbf{W} = \mathbf{R} (\mathbf{R}^T \mathbf{J}_{\mathbf{h}} \mathbf{R})^\dagger \mathbf{R}^T \mathbf{J}_{\mathbf{h}} \mathbf{K} \mathbf{J}_{\mathbf{h}} \mathbf{R} (\mathbf{R}^T \mathbf{J}_{\mathbf{h}} \mathbf{R})^\dagger \mathbf{R}^T.$$

In particular, $\sqrt{n}(\text{vec}(\hat{\boldsymbol{\beta}}_{\text{RE}}) - \text{vec}(\boldsymbol{\beta}))$ 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 $\mathbf{W}$. The explicit expression for the gradient matrix $\mathbf{R} = \partial \mathbf{h}(\boldsymbol{\phi}) / \partial \boldsymbol{\phi}$ is given in the Supplement equation (D.1).

The $\sqrt{n}$ -consistency of the reduced-rank envelope estimator $\hat{\boldsymbol{\beta}}_{\text{RE}}$ is essentially because that $\hat{\boldsymbol{\beta}}_{\text{OLS}}$ and $\mathbf{S}_{\mathbf{Y}|\mathbf{X}}$ are $\sqrt{n}$ -consistent regardless of normality assumption and also because of the properties of $\mathcal{F}(\mathbf{h}, \hat{\mathbf{h}}_{\text{OLS}})$. The asymptotic covariance matrix $\mathbf{W}_{11}$ can be estimated straightforwardly using the plug-in method once $\mathbf{K}$ is estimated, but its accuracy for any fixed sample size will depend on the distribution of ϵ , which is usually unknown in practice. Fortunately, bootstrap methods can provide good estimates of $\mathbf{W}_{11}$, as illustrated in Section 6.3.

5 Selections of rank and envelope dimension

5.1 Rank: $d = \text{rank}(\beta)$

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

$$\hat{\beta}_{\text{std}} = \{(n - p - 1)\} / n \}^{1/2} \mathbf{S}_{\mathbf{X}}^{1/2} \hat{\beta}_{\text{OLS}} \mathbf{S}_{\mathbf{Y}|\mathbf{X}}^{-1/2}. \quad (5.1)$$

Under the null hypothesis that $H_0 : d = d_0$, Bura and Cook (2003) showed that Λ_{d_0} is asymptotically distributed as a $\chi_{(p-d_0)(r-d_0)}^2$ random variable. The rank d is then determined by comparing a sequence of test statistics Λ_{d_0} , $d_0 = 0, \dots, \min(p, r) - 1$, to the percentiles of their null distribution $\chi_{(p-d_0)(r-d_0)}^2$. The sequence of tests terminates at the first non-significant test of $H_0 : d = d_0$ and then d_0 serves as an estimate of the rank of β .

5.2 Envelope dimension: $u = \dim(\mathcal{E}_{\Sigma}(\beta))$

Since the envelope dimension satisfies $d \leq u \leq r$, standard techniques such as sequential likelihood-ratio tests, AIC and BIC can be applied to select u , as in Cook et al. (2010).

For any possible combination (d, u) with $0 \leq d \leq u \leq r$, let $\hat{L}_{d,u}$ denote the maximized log-likelihood function (c.f. (A.3)), which is evaluated at the maximum likelihood estimators in Proposition 4. Assuming d is known, then $\Lambda_{d,u_0} = 2(\hat{L}_{d,r} - \hat{L}_{d,u_0})$ is asymptotically distributed as a $\chi_{(r-u_0)d}^2$ random variable under the null hypothesis $H_0 : u = u_0$. Thus, a sequence of likelihood ratio tests of $u_0 = d, \dots, r - 1$ can be used to determine u after d is determined by the method described in Section 5.1. The first non-significant value of u_0 will serve as the envelope dimension.

Information criteria such as AIC and BIC can be used to select (d, u) simultaneously. We write AIC as $\mathcal{A}_{d,u} = 2K_{d,u} - 2\hat{L}_{d,u}$, where $K_{d,u} = (p+u-d)d + r(r+1)/2$ is the total number of parameters in the reduced-rank envelope model, and write BIC as $\mathcal{B}_{d,u} = \log(n)K_{d,u} - 2\hat{L}_{d,u}$. We search (d, u) from $(0, 0)$ to (r, r) with constraint $d \leq u$ and choose the pair that has the

smallest AIC or BIC. Alternatively, we can first determine d from the asymptotic chi-squared tests in Section 5.1 and then search u from $d, \dots, r$ with the smallest AIC or BIC, which could save a lot of computation. The computation cost for determining d by the sequential chi-squared tests in Section 5.1 is substantially cheaper than the computation cost in calculating AIC and BIC, which involves sequence of Grassmannian optimizations.

When sample size is not too small, our experience suggests that the most favorable procedure is BIC selection for $u = d, \dots, r$ where d is guided by the sequential chi-squared tests. Since the true envelope dimension always exist, BIC is consistent in the sense that the probability of selecting the correct u approaches 1, given the correct d . There are many articles comparing AIC and BIC, from both theoretical and practical points of view, for example Shao (1997) and Yang (2005).

The rank d and envelope dimension u can also be determined by cross-validation or by using hold-out samples. These approaches are especially appropriate when prediction is the primary goal of the study rather than correctness of the selected model.

6 Simulations

6.1 Rank and dimension

In all the simulations, we first filled in Γ , $\boldsymbol{\eta}$ and $\mathbf{B}$ with random uniform (0,1) numbers, and then Γ was standardized so that $\Gamma^T \Gamma = \mathbf{I}_u$ and $\boldsymbol{\beta} = \Gamma \boldsymbol{\eta} \mathbf{B}$ was standardized so that $\|\boldsymbol{\beta}\|_F = 1$. Estimation errors were defined as $\|\boldsymbol{\beta} - \hat{\boldsymbol{\beta}}\|_F$. Unless otherwise specified, the predictors and errors were simulated independently from $N(0, \mathbf{I}_p)$ and $N(0, \boldsymbol{\Sigma})$ distributions. All figures were generated based on averaging over 200 independent replicate data sets.

In this section, we present simulation results to demonstrate the behavior of the proposed method using various sample sizes, ranks d , dimensions u . We simulated data from model (2.4), where $[\boldsymbol{\Omega}]_{ij} = (-0.9)^{|i-j|}$ and $[\boldsymbol{\Omega}_0]_{ij} = 5 \cdot (-0.5)^{|i-j|}$. Figure 6.1 summarizes the effect of dimension and rank on the relative performances of each methods. In the left plot $(d, u, p, r) = (1, 10, 10, 20)$. Since the rank was only one but the envelope dimension was ten,

reduced-rank regression had a dramatic improvement over ordinary least squares, while the ordinary envelope method had a relatively modest gain over ordinary least squares. The reduced-rank envelope had a relatively small edge over reduced-rank regression. The second case was $(d, u, p, r) = (4, 5, 6, 20)$, where β had nearly full column rank and the envelope dimension was much smaller than the number of response variables. Not surprisingly, reduced-rank regression had modest gain over ordinary least squares while the envelope estimator and the reduced-rank envelope estimator had similar behavior and significantly improved over ordinary least squares and reduced-rank regression. The last case was chosen as $(d, u, p, r) = (5, 10, 15, 20)$ so that there was no particular favor towards either the envelope method or reduced-rank regression. We found good improvement over ordinary least squares by both reduced-rank regression and envelopes. However, reduced-rank envelopes combined both of their strengths and resulted in a bigger gain.

We found in practice that reduced-rank envelopes typically have improved performance over reduced-rank regression and envelope estimators, and it has similar behavior to one of the two estimators if the other one performed poorly. Even in the extreme cases where $d = p$ or $u = r$, reduced-rank envelopes can still gain drastically over ordinary least squares similar to the results in Figure 6.1.

We next illustrate the asymptotic chi-squared test for rank detection combined with BIC selection for envelope dimension, as discussed in Section 5. Using the same simulation model, we took $(d, u, p, r) = (3, 5, 8, 12)$, where the total number of parameters in the reduced-rank envelope model was 108. The percentages of correct detections for d and u were plotted in Figure 6.2 versus sample size. The BIC selection of u was based on the correct rank d . The significance level of the chi-squared tests was 0.05. As seen from the figure, the probability of selecting the correct d was about 0.9 at $n = 400$ samples and the probability of correct detection settled at 95% for larger n , as predicted by the hypothesis testing theory. BIC selection for the envelope dimension u seemed to be very accurate even with small samples. The likelihood-ratio tests and AIC selection for u were not nearly as effective as BIC and thus were omitted from the plot.

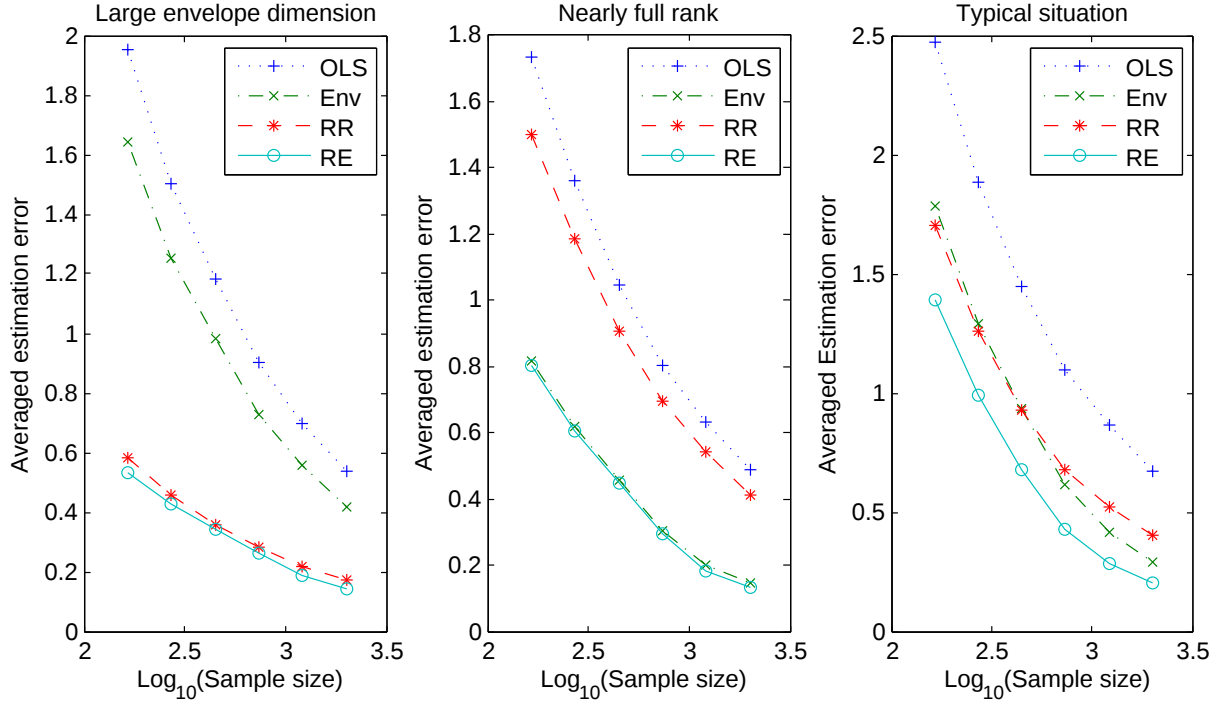


Figure 6.1: Effect of rank and dimension. Averaged estimation error on the vertical axis is defined as averaged $\|\beta - \hat{\beta}\|_F$ over 200 independent data sets. The dimensions of the three plots were: (1) large envelope dimension case $(d, u, p, r) = (1, 10, 10, 20)$; (2) nearly full rank case $(d, u, p, r) = (4, 5, 6, 20)$ and (3) a typical situation $(d, u, p, r) = (5, 10, 15, 20)$. The sample sizes varied from 160 to 2000 and were shown in a logarithmic scale.

We also considered BIC selection for u and d simultaneously. The probability of simultaneous correctness was less than 70% for $n \leq 600$ but reached more than 95% correctness for $n \geq 900$. In our experience the best method for determining dimensions is to use the chi-squared test for d and BIC selection on u based on the selected d . Overestimation of d and u usually is not a serious issue but underestimation of d and u will certainly cause bias in estimation.

6.2 Signal-versus-noise and material-versus-immaterial

In this section, we describe the behavior of each method with varying signal-to-noise ratios and ratios of immaterial variation to material variation. We fixed the sample size at 400 and the dimensions were $(d, u, p, r) = (3, 7, 10, 20)$. The covariances had the forms of $[\Omega]_{ij} = \sigma^2 \cdot (-0.9)^{|i-j|}$ and $[\Omega_0]_{ij} = \sigma_0^2 \cdot (-0.9)^{|i-j|}$ with varying constants $\sigma^2, \sigma_0^2 > 0$.

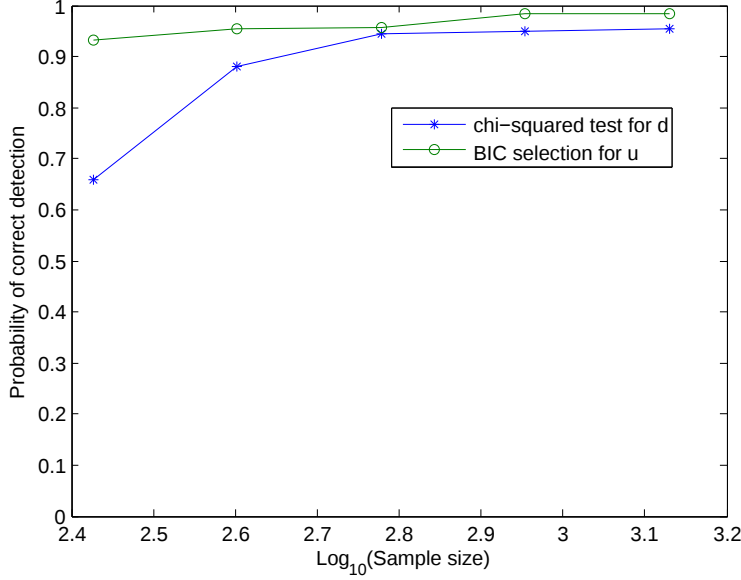


Figure 6.2: The empirical probability of correct detection versus sample sizes. BIC selection on u was based on true rank d .

In the study of varying signal-to-noise ratio, we kept $\sigma^2 = \sigma_0^2$. And because $\|\beta\|_F = 1$, the signal-to-noise ratio was simply $1/\sigma^2$ which varied from 0.1 to 10. Figure 6.3 summarizes the results of two numerical experiments. All the four lines in this log-log scale signal-to-noise ratio plot are roughly parallel, which implies that the four methods are exponentially more distinguishable in weaker signal. Comparing reduced-rank regression to envelopes, the reduced-rank regression seemed to perform better in stronger signals (signal-to-noise ratio ≥ 1), but the envelope estimator was less vulnerable to weaker signals (signal-to-noise ratio ≤ 1). This was because the envelope method can gain information from the error term $\Sigma = \Gamma\Omega\Gamma^T + \Gamma_0\Omega_0\Gamma_0^T$ while reduced-rank regression and the standard method cannot. Reduced-rank envelope estimators combined the strengths of reduced-rank regression and envelopes, and hence outperformed both estimators in strong and weak signals.

In the study of varying immaterial-to-material variance ratio, we kept $\sigma^2 = 1$ and changed σ_0^2 . The ratio is then defined as σ_0^2 and the horizontal axis in the plot is $\log_{10}(\sigma_0^2)$, which varied from -0.5 to 2. Not surprisingly, reduced-rank regression and ordinary least squares behaved similarly because they did not gain information from the covariance structure of Σ .

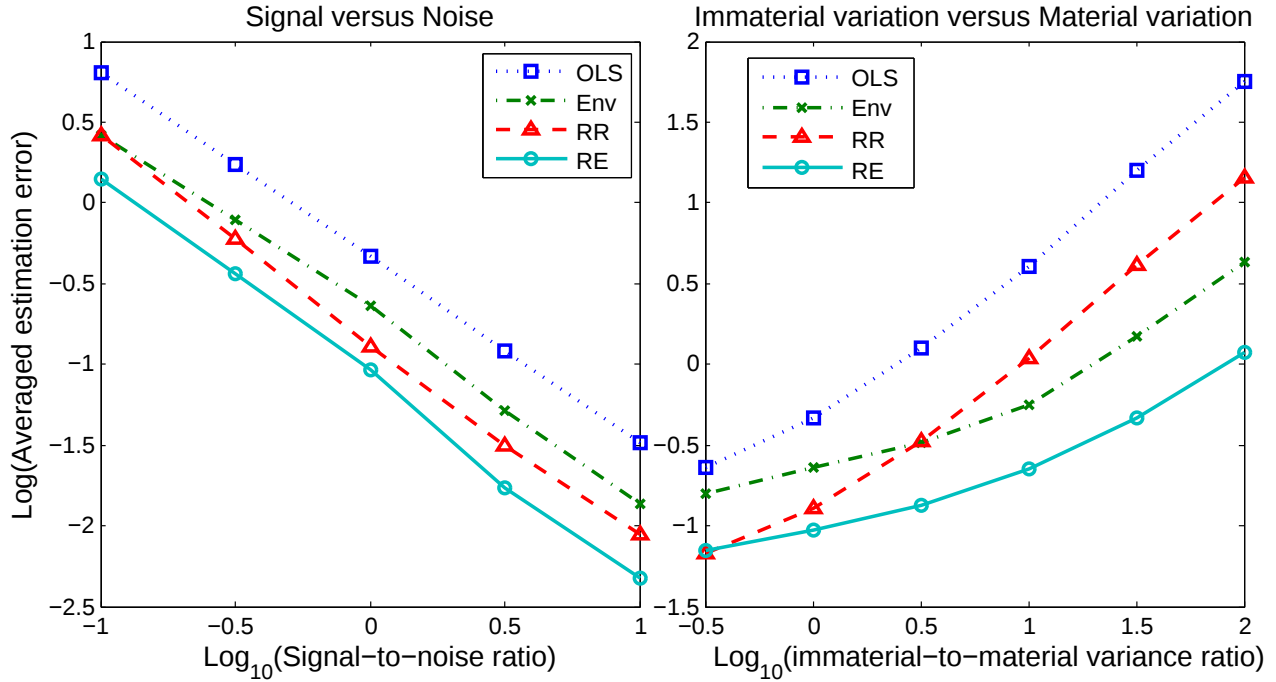


Figure 6.3: Varying the signal-to-noise ratio and the immaterial-to-material variance ratio.

The envelope estimator and the reduced-rank envelope estimator had similar behavior, and they had much better performances over ordinary least squares and reduced-rank regression when the immaterial variation was large. This is due to the fact that envelope methods can efficiently eliminate the immaterial information. In this example, the averaged estimation errors for ordinary least squares, reduced-rank regression and envelope were 7.2, 3.9 and 1.8 times of that of the envelope reduced-rank regression when $\sigma_0^2 = 100$.

6.3 Bootstrap standard errors

To illustrate the application of the bootstrap for estimating the standard errors of regression coefficients, we considered a model with $(d, u, p, r) = (2, 4, 6, 8)$. Residual bootstrap samples were used since we considered $\mathbf{X}$ as a non-stochastic predictor. Both $\mathbf{\Omega}$ and $\mathbf{\Omega}_0$ were randomly generated as $\mathbf{M}\mathbf{M}^T$, where $\mathbf{M} \in \mathbb{R}^{4 \times 4}$ was filled with uniform (0,1) numbers. The error term ϵ_i was simulated as $\epsilon_i = \Sigma^{1/2}\mathbf{U}_i$, where $\mathbf{U}_i$ was a vector of i.i.d. random variable with mean 0 standard deviation 1. We simulated both normal and uniform $\mathbf{U}_i$. The standard errors of a se-

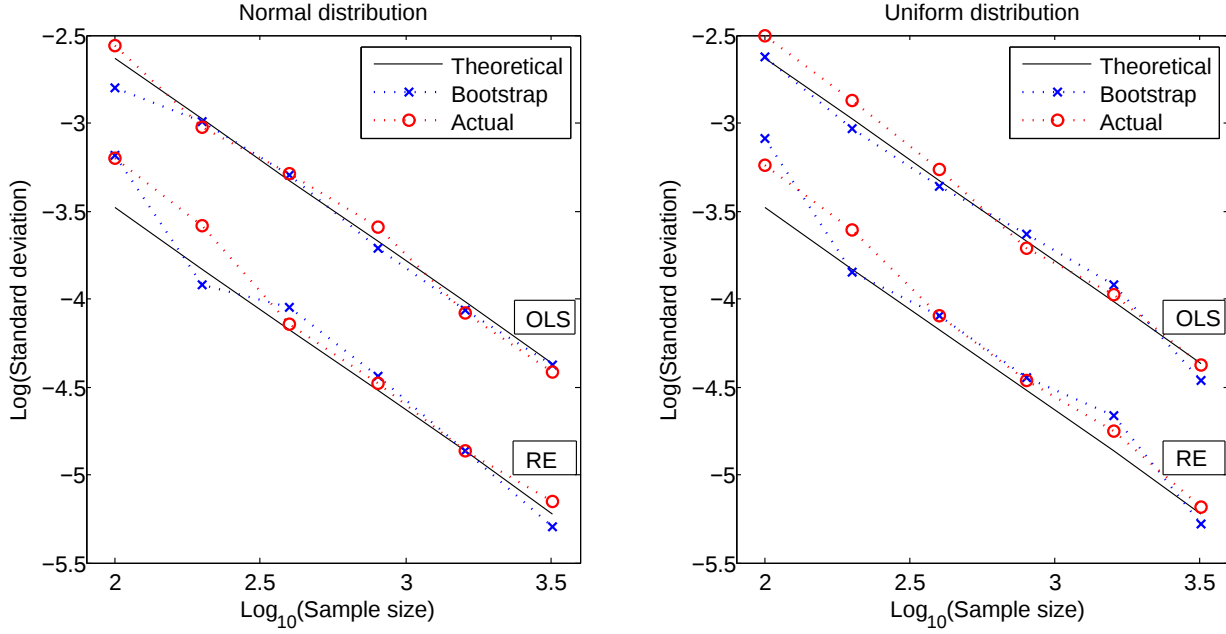


Figure 6.4: Theoretical, bootstrap and actual standard errors with normal and uniform errors ϵ . The sample sizes were 100, 200, 400, ..., 3200. The standard errors for reduced-rank regression and envelope estimators were consistently between the ordinary least squares and the envelope reduced-rank regression standard errors, and were not included in these plots for better visualization.

lected element in $\hat{\beta}$ were plotted in Figure 6.4. For both normal and non-normal data, the three types of standard error estimates agreed well: the theoretical standard errors were the squared roots of the diagonal elements in the asymptotic covariances of each estimators divided by $\sqrt{n}$; the actual standard errors were based on 200 independent realizations; and the bootstrap standard errors were based on 200 bootstrap replicate data sets. Moreover, the bootstrap standard errors were close to the theoretical standard errors of the maximum likelihood estimators even when the normality assumption was violated. As expected, the reduced-rank envelope estimator had much smaller standard errors than those of the ordinary least squares estimator. We also simulated non-normal errors from t -distribution and χ^2 -distribution, and obtained results similar to Figure 6.4.

7 Sales people test scores data

This data set consisted of 50 sales people from a firm. Three performances variables were used as predictors: growth of sales (X_1), profitability of sales (X_2) and new account sales (X_3). And four response variables were test scores on creativity (Y_1), mechanical reasoning (Y_2), abstract reasoning (Y_3) and mathematical ability (Y_4). This data set can be found in Johnson and Wichern (2007).

The chi-squared rank test in Section 5.1 suggested that $d = 2$ at level 0.01. Then based on BIC we selected the envelope dimension to be $u = 3$. We computed the fractions $f_{ij} := 1 - \widehat{\text{avar}}^{1/2}(\sqrt{n}\widehat{\beta}_{ij,\text{RE}})/\widehat{\text{avar}}^{1/2}(\sqrt{n}\widetilde{\beta}_{ij})$ for all i and j , where $\widetilde{\beta}$ denotes one of the estimators to be compared: $\widehat{\beta}_{\text{RR}}$, $\widehat{\beta}_{\text{Env}}$ and $\widehat{\beta}_{\text{OLS}}$. Comparing to ordinary least squares, the standard deviations of the elements in the reduced-rank envelope estimator were 5% to 60% smaller, $0.05 \leq f_{ij} \leq 0.60$. Hypothetically, a sample size of more than 300 observations, in contrast to the original 50 observations, would be needed to achieve a 60% smaller standard deviation in ordinary least squares. The fractions for comparing with the reduced-rank regression estimator were $0.01 \leq f_{ij} \leq 0.24$, where 24% smaller standard deviation than reduced-rank regression implies a doubling of the observations for reduced-rank regression to achieve the same performance as reduced-rank envelope estimator. At last, the reduced-rank envelope estimator compared to the ordinary envelope estimator, had 3% to 51% smaller standard deviations, where 51% smaller standard deviation meant four times the sample size, $n = 200$, for the ordinary envelope estimator.

Supplementary Materials

Proofs and Technical Details: Detailed proofs for all Lemmas and Propositions are provided in the online supplement to this article. (PDF file)

References

- [1] ANDERSON, T. W. (1951), Estimating linear restrictions on regression coefficients for multivariate normal distributions. *The Annals of Mathematical Statistics*, **22**, 327–351.
- [2] ANDERSON, T. W. (1999), Asymptotic distribution of the reduced rank regression estimator under general conditions. *The Annals of Statistics*, **27**, 1141–1154.
- [3] ANDERSON, T. W. (2002), Canonical correlation analysis and reduced rank regression in autoregressive models. *The Annals of Statistics*, **30**, 1134–1154.
- [4] ADRAGNI, K., COOK, R.D. AND WU, S. (2012), GrassmannOptim: An R Package for Grassmann Manifold Optimization. *Journal of Statistical Software*, **50**, 1–18.
- [5] BURA, E. AND COOK, R.D. (2003), Rank estimation in reduced-rank regression. *Journal of Multivariate Analysis*, **87**, 159–176.
- [6] CHEN, L. AND HUANG, J.Z. (2012), Sparse reduced-rank regression for simultaneous dimension reduction and variable selection. *Journal of the American Statistical Association*, **107**, 1533–1545.
- [7] CHEN, K., CHAN, K.S. AND STENSETH, N.C. (2012), Reduced rank stochastic regression with a sparse singular value decomposition. *JRSS-B*, **74**, 203–221.
- [8] CONWAY, J. (1990). A course in functional analysis. Second edition. *Springer, New York*.
- [9] COOK, R. D., HELLAND, I. S. AND SU, Z. (2013), Envelopes and partial least squared regression. *To appear in JRSS-B*.

- [10] COOK, R. D., LI, B. AND CHIAROMONTE, F. (2010). Envelope models for parsimonious and efficient multivariate linear regression (with discussion). *Statistica Sinica*, **20**, 927–1010.
- [11] EDELMAN, A., TOMAS, A.A. AND SMITH, S.T. (1998), The geometry of algorithms with orthogonality constraints, *SIAM Journal of Matrix Analysis and Applications*, **20**, 303–353.
- [12] HENDERSON, H. V. AND SEARLE, S. R. (1979), Vec and Vech operators for matrices, with some uses in Jacobians and multivariate statistics. *Canadian Journal of Statistics*, **7**, 65–81.
- [13] IZENMAN, A. J. (1975), Reduced-rank regression for the multivariate linear model. *Journal of Multivariate Analysis*. **5**, 248–264.
- [14] REINSEL, G. C. AND VELU, R. P. (1998), Multivariate reduced-rank regression: theory and applications. *Springer, New York*.
- [15] SHAO, J. (1997), An asymptotic theory for linear model selection (with discussion). *Statistica sinica*, **7**, 221–264.
- [16] SHAPIRO, A. (1986), Asymptotic theory of overparameterized structural models, *Journal of the American Statistical Association*, **81**, 142–149.
- [17] STOICA, P. AND VIBERG, M. (1996), Maximum likelihood parameter and rank estimation in reduced-rank multivariate linear regressions. *IEEE Transactions on Signal Processing*, **44**, 3069–3079.
- [18] SU, Z. AND COOK, R.D. (2011), Partial envelopes for efficient estimation in multivariate linear regression. *Biometrika*, **98**, 133–146.
- [19] TSO, M.K.S. (1981), Reduced-rank regression and canonical analysis. *JRSS-B*, **43**, 183–189.

- 550 [20] TYLER, D.E. (1981). Asymptotic inference for eigenvectors. *Annals of Statis-*
551 *tics*, **9**, 725–736.
- 552 [21] YANG, Y. (2005), Can the strengths of AIC and BIC be shared? A conflict be-
553 tween model identification and regression estimation. *Biometrika*. **92**, 934–950.

Supplement: Proofs and Technical Details for “Reduced-rank Envelope Model”

R. Dennis Cook, Liliana Forzani and Xin Zhang

A Maximizing the likelihood-based objective function (2.2)

In this Section, we consider maximizing $L_n(\alpha, \beta, \Sigma)$ from (2.2) under different model parameterizations regarding standard, reduced-rank, envelope and reduced-rank envelope models. Maximizing L_n from (2.2) is equivalent to deriving maximum likelihood estimators with normally distributed error $\epsilon \sim N(0, \Sigma)$ as follows. Lemmas 1 and 2 and Propositions 2 and 4 are proved directly in the derivation of estimators.

A.1 Standard regression and envelope regression

Maximum likelihood estimators for the standard regression model is the ordinary least squares estimator, $\hat{\beta}_{OLS} = S_{YX}S_X^{-1}$ and $\hat{\Sigma}_{OLS} = S_{Y|X}$. From Cook et al. (2010), we have the maximum likelihood estimators for the envelope model as

$$\begin{aligned}\hat{\Gamma}_{Env} &= \arg \min_{G \in \mathcal{G}_{r,u}} \{ \log |G^T S_{Y|X} G| + \log |G^T S_Y^{-1} G| \} \\ \hat{\beta}_{Env} &= \hat{\Gamma}_{Env} S_{\hat{\Gamma}_{Env}^T Y, X} S_X^{-1} = P_{\hat{\Gamma}_{Env}} \hat{\beta}_{OLS} \\ \hat{\Sigma}_{Env} &= P_{\hat{\Gamma}_{Env}} S_{Y|X} P_{\hat{\Gamma}_{Env}} + Q_{\hat{\Gamma}_{Env}} S_Y Q_{\hat{\Gamma}_{Env}}.\end{aligned}$$

A.2 Reduced-rank regression (proof of Lemma 1)

Following Anderson (1999) equation (2.13), we let $\hat{L} \in \mathbb{R}^{p \times d}$ denote $S_X^{-1/2}[\mathbf{v}_1, \dots, \mathbf{v}_d]$, where $\mathbf{v}_i$ is the i -th eigenvector of $S_X^{-1/2} S_{X \circ Y} S_X^{-1/2}$. Then the estimators can be written as $\hat{\alpha}_{RR} = \bar{Y} - \hat{\beta}_{RR} \bar{X}$, $\hat{\beta}_{RR} = S_{YX} \hat{L} \hat{L}^T$ and $\hat{\Sigma}_{RR} = S_Y - \hat{\beta}_{RR} S_{XY}$. We then use the sample canonical

571 correlation matrix notation to get the results in Lemma 1: $\mathbf{S}_\mathbf{X}^{-1/2} \mathbf{S}_{\mathbf{X} \circ \mathbf{Y}} \mathbf{S}_\mathbf{X}^{-1/2} = \mathbf{C}_{\mathbf{X}\mathbf{Y}} \mathbf{C}_{\mathbf{Y}\mathbf{X}}$ and

$$\begin{aligned}\hat{\beta}_{\text{RR}} &= \mathbf{S}_{\mathbf{Y}\mathbf{X}} \hat{\mathbf{L}} \hat{\mathbf{L}}^T = \mathbf{S}_{\mathbf{Y}\mathbf{X}} \mathbf{S}_\mathbf{X}^{-1/2} \mathbf{P}_{\mathbf{C}_{\mathbf{X}\mathbf{Y}}^{(d)}} \mathbf{S}_\mathbf{X}^{-1/2} \\ &= \mathbf{S}_\mathbf{Y}^{1/2} \mathbf{C}_{\mathbf{Y}\mathbf{X}} \mathbf{P}_{\mathbf{C}_{\mathbf{X}\mathbf{Y}}^{(d)}} \mathbf{S}_\mathbf{X}^{-1/2} = \mathbf{S}_\mathbf{Y}^{1/2} \mathbf{C}_{\mathbf{Y}\mathbf{X}}^{(d)} \mathbf{S}_\mathbf{X}^{-1/2}.\end{aligned}$$

572 **A.3 Reduced-rank envelope regression**

573 **A.3.1 Proof of Lemma 2**

574 Estimation for the envelope model is facilitated by the following consideration which is straight-
575 forward from (2.4).

$$\mathbf{\Gamma}^T \mathbf{Y}_i = \mathbf{\Gamma}^T \boldsymbol{\alpha} + \boldsymbol{\eta} \mathbf{B} \mathbf{X}_i + \mathbf{\Gamma}^T \boldsymbol{\epsilon}_i, \quad (\text{A.1})$$

$$\mathbf{\Gamma}_0^T \mathbf{Y}_i = \mathbf{\Gamma}_0^T \boldsymbol{\alpha} + \mathbf{\Gamma}_0^T \boldsymbol{\epsilon}_i, \quad (\text{A.2})$$

576 where $\mathbf{\Gamma}^T \boldsymbol{\epsilon} \sim N(0, \boldsymbol{\Omega})$, $\mathbf{\Gamma}_0^T \boldsymbol{\epsilon} \sim N(0, \boldsymbol{\Omega}_0)$, $\mathbf{\Gamma}^T \boldsymbol{\epsilon} \perp \mathbf{\Gamma}_0^T \boldsymbol{\epsilon}$.

577 The maximum likelihood estimator of $\boldsymbol{\alpha}$ is $\hat{\boldsymbol{\alpha}}_{\text{RE}} = \bar{\mathbf{Y}} - \hat{\beta}_{\text{RE}} \bar{\mathbf{X}}$ and effectively we could use
578 centered response $\mathbf{Y}_{ci} := \mathbf{Y}_i - \bar{\mathbf{Y}}$ and centered predictors $\mathbf{X}_{ci} = \mathbf{X}_i - \bar{\mathbf{X}}$ to omit the analysis
579 on $\boldsymbol{\alpha}$ and $\hat{\boldsymbol{\alpha}}_{\text{RE}}$. Then the partially maximized log-likelihood with known dimensions u and d
580 can be decomposed into the following two additive parts since $\mathbf{\Gamma}^T \boldsymbol{\epsilon}$ is independent of $\mathbf{\Gamma}_0^T \boldsymbol{\epsilon}$.

$$L_n(\mathbf{\Gamma}, \boldsymbol{\eta}, \mathbf{B}, \boldsymbol{\Omega}_0, \boldsymbol{\Omega} | d, u) \simeq L_{1,n}(\mathbf{\Gamma}, \boldsymbol{\eta}, \mathbf{B}, \boldsymbol{\Omega} | d, u) + L_{2,n}(\mathbf{\Gamma}_0, \boldsymbol{\Omega}_0 | u) \quad (\text{A.3})$$

581 where $L_{1,n}(\mathbf{\Gamma}, \boldsymbol{\eta}, \mathbf{B}, \boldsymbol{\Omega} | d, u)$ corresponds to the likelihood from (A.1) and is given by

$$-\frac{n}{2} \left\{ \log |\boldsymbol{\Omega}| + \text{trace} \left[\boldsymbol{\Omega}^{-1} \frac{1}{n} \sum_{i=1}^n (\mathbf{\Gamma}^T \mathbf{Y}_{ci} - \boldsymbol{\eta} \mathbf{B} \mathbf{X}_{ci}) (\mathbf{\Gamma}^T \mathbf{Y}_{ci} - \boldsymbol{\eta} \mathbf{B} \mathbf{X}_{ci})^T \right] \right\}, \quad (\text{A.4})$$

582 and $L_{2,n}(\mathbf{\Gamma}_0, \boldsymbol{\Omega}_0 | u)$ corresponds to the likelihood from (A.2) and is equal to

$$-\frac{n}{2} \left\{ \log |\boldsymbol{\Omega}_0| + \text{trace} \left[\boldsymbol{\Omega}_0^{-1} \frac{1}{n} \sum_{i=1}^n \mathbf{\Gamma}_0^T \mathbf{Y}_{ci} \mathbf{Y}_{ci}^T \mathbf{\Gamma}_0 \right] \right\}.$$

583 It follows that $L_{2,n}$ is maximized over $\boldsymbol{\Omega}_0$ by $\sum_{i=1}^n \mathbf{\Gamma}_0^T \mathbf{Y}_{ci} \mathbf{Y}_{ci}^T \mathbf{\Gamma}_0 / n = \mathbf{\Gamma}_0^T \mathbf{S}_\mathbf{Y} \mathbf{\Gamma}_0$. Substituting
584 back, we find the following partially maximized form for $L_{2,n}$:

$$L_{2,n}(\mathbf{\Gamma}_0 | u) \simeq -(n/2) \log |\mathbf{\Gamma}_0^T \mathbf{S}_\mathbf{Y} \mathbf{\Gamma}_0|. \quad (\text{A.5})$$

585 Holding Γ fixed, the log-likelihood $L_{1,n}$ is same as the log-likelihood for reduced rank regres-
 586 sion of $\Gamma^T \mathbf{Y}$ on $\mathbf{X}$. Therefore, by replacing $r \rightarrow u$, $\mathbf{Y} \rightarrow \Gamma^T \mathbf{Y}$, $\mathbf{A} \rightarrow \boldsymbol{\eta}$, $\mathbf{B} \rightarrow \mathbf{B}$ and $\boldsymbol{\Sigma} \rightarrow \boldsymbol{\Omega}$
 587 in (2.2) and in Lemma 1, we partially maximize $L_{1,n}(\Gamma, \boldsymbol{\eta}, \mathbf{B}, \boldsymbol{\Omega}|d, u)$ over $\boldsymbol{\eta}$, $\mathbf{B}$ and $\boldsymbol{\Omega}$ and
 588 obtain the maximum likelihood estimators as

$$\begin{aligned}\widehat{\boldsymbol{\eta}}_\Gamma \widehat{\mathbf{B}}_\Gamma &= \mathbf{S}_{\Gamma^T \mathbf{Y}}^{1/2} \mathbf{C}_{\Gamma^T \mathbf{Y}, \mathbf{X}}^{(d)} \mathbf{S}_{\mathbf{X}}^{-1/2} \\ \widehat{\boldsymbol{\Omega}}_\Gamma &= \mathbf{S}_{\Gamma^T \mathbf{Y}}^{1/2} \left\{ \mathbf{I}_u - \mathbf{C}_{\Gamma^T \mathbf{Y}, \mathbf{X}}^{(d)} \mathbf{C}_{\mathbf{X}, \Gamma^T \mathbf{Y}}^{(d)} \right\} \mathbf{S}_{\Gamma^T \mathbf{Y}}^{1/2},\end{aligned}$$

589 from which Lemma 2 follows.

590 A.3.2 Proof of Proposition 2

591 The log-likelihood function in (A.3) after partial maximization becomes

$$L_n(\Gamma|d, u) \simeq -(n/2) \left\{ \log |\Gamma_0^T \mathbf{S}_\mathbf{Y} \Gamma_0| + \log |\widehat{\boldsymbol{\Omega}}_\Gamma| \right\}, \quad (\text{A.6})$$

592 which lead us to the objective function $F_n(\mathbf{G}|d, u) := (-2/n)L_n(\mathbf{G}|d, u)$ for numerical opti-
 593 mization over $\text{span}(\mathbf{G}) \in \mathcal{G}_{r,u}$. We next simplify the expression of $\log |\widehat{\boldsymbol{\Omega}}_\mathbf{G}|$ as

$$\begin{aligned}\log |\widehat{\boldsymbol{\Omega}}_\mathbf{G}| &= \log |\mathbf{S}_{\mathbf{G}^T \mathbf{Y}}^{1/2} \left\{ \mathbf{I}_u - \mathbf{C}_{\mathbf{G}^T \mathbf{Y}, \mathbf{X}}^{(d)} \mathbf{C}_{\mathbf{X}, \mathbf{G}^T \mathbf{Y}}^{(d)} \right\} \mathbf{S}_{\mathbf{G}^T \mathbf{Y}}^{1/2}| \\ &= 2 \cdot \log |\mathbf{S}_{\mathbf{G}^T \mathbf{Y}}^{1/2}| + \log |\mathbf{I}_u - \mathbf{C}_{\mathbf{G}^T \mathbf{Y}, \mathbf{X}}^{(d)} \mathbf{C}_{\mathbf{X}, \mathbf{G}^T \mathbf{Y}}^{(d)}| \\ &= \log |\mathbf{S}_{\mathbf{G}^T \mathbf{Y}}| + \log |\mathbf{I}_u - \mathbf{S}_{\mathbf{Z}_\mathbf{G} \circ \mathbf{X}}^{(d)}|,\end{aligned}$$

594 where $\mathbf{S}_{\mathbf{G}^T \mathbf{Y}} = \mathbf{G}^T \mathbf{S}_\mathbf{Y} \mathbf{G}$ and $\mathbf{Z}_\mathbf{G} = (\mathbf{G}^T \mathbf{S}_\mathbf{Y} \mathbf{G})^{-1} \mathbf{G}^T \mathbf{Y}$ is the standardized random vector in
 595 $\mathbb{R}^u$. Equation (3.2) is then obtained by noticing $\log |\mathbf{G}_0^T \mathbf{S}_\mathbf{Y} \mathbf{G}_0^T| = \log |\mathbf{S}_\mathbf{Y}| + \log |\mathbf{G}^T \mathbf{S}_\mathbf{Y}^{-1} \mathbf{G}|$ in
 596 the objective function (A.6).

597 We next prove the equality in (3.3). The first term in (3.3) can be re-expressed as $\log |\mathbf{G}^T \mathbf{S}_\mathbf{Y} \mathbf{G}| =$
 598 $\log |\mathbf{G}^T \mathbf{S}_\mathbf{Y} \mathbf{G}| + \log |\mathbf{S}_{\mathbf{Z}_\mathbf{G} \circ \mathbf{X}}|$ according to the following.

$$\mathbf{G}^T \mathbf{S}_\mathbf{Y} \mathbf{G} = \mathbf{G}^T \mathbf{S}_\mathbf{Y} \mathbf{X} \mathbf{S}_\mathbf{X}^{-1} \mathbf{S}_\mathbf{X} \mathbf{Y} \mathbf{G} = \mathbf{S}_{\mathbf{G}^T \mathbf{Y}, \mathbf{X}} \mathbf{S}_\mathbf{X}^{-1} \mathbf{S}_{\mathbf{X}, \mathbf{G}^T \mathbf{Y}} = \mathbf{S}_{\mathbf{G}^T \mathbf{Y}}^{1/2} \mathbf{S}_{\mathbf{Z}_\mathbf{G} \circ \mathbf{X}} \mathbf{S}_{\mathbf{G}^T \mathbf{Y}}^{1/2}.$$

599 The objective function in (3.3) now become

$$\log |\mathbf{G}^T \mathbf{S}_\mathbf{Y} \mathbf{G}| + \log |\mathbf{S}_{\mathbf{Z}_\mathbf{G} \circ \mathbf{X}}| + \log |\mathbf{G}^T \mathbf{S}_\mathbf{Y}^{-1} \mathbf{G}| + \sum_{i=d+1}^u \log [\widehat{\omega}_i(\mathbf{G})],$$

600 where $\widehat{\omega}_i(\mathbf{G})$ is the i -th eigenvalue of $\mathbf{S}_{\mathbf{Z}_G \circ \mathbf{X}}$. The equality connecting (3.2) and (3.3) is proved
 601 by noticing that $\mathbf{S}_{\mathbf{Z}_G|\mathbf{X}} = \mathbf{S}_{\mathbf{Z}_G} - \mathbf{S}_{\mathbf{Z}_G \circ \mathbf{X}} = \mathbf{I}_u - \mathbf{S}_{\mathbf{Z}_G \circ \mathbf{X}}$ and that the log-determinant of a positive
 602 definite matrix is the sum of the logarithms of its eigenvalues.

603 A.3.3 Proof of Proposition 4

604 The proof follows trivially by combining the results in Lemma 2 and Proposition 2.

605 B Proof for Proposition 3

606 Recall that in (3.3), $\widehat{\omega}_i(\mathbf{G})$ is the i -th eigenvalues of the following matrix.

$$\begin{aligned} \mathbf{S}_{\mathbf{Z}_G|\mathbf{X}}^{-1} &= (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}|\mathbf{X}} \mathbf{G})^{-1/2} (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}} \mathbf{G}) (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}|\mathbf{X}} \mathbf{G})^{-1/2} \\ &= \mathbf{I}_u + (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}|\mathbf{X}} \mathbf{G})^{-1/2} (\mathbf{G}^T \mathbf{S}_{\mathbf{Y} \circ \mathbf{X}} \mathbf{G}) (\mathbf{G}^T \mathbf{S}_{\mathbf{Y}|\mathbf{X}} \mathbf{G})^{-1/2}, \end{aligned}$$

607 which relies on the two sample covariance matrices: $\mathbf{S}_{\mathbf{Y}|\mathbf{X}}$ and $\mathbf{S}_{\mathbf{Y} \circ \mathbf{X}}$. These two matrices are
 608 both positive semi-definite and converge to Σ and $\Sigma_{\mathbf{Y} \circ \mathbf{X}} = \Sigma_{\mathbf{Y}\mathbf{X}} \Sigma_{\mathbf{X}}^{-1} \Sigma_{\mathbf{X}\mathbf{Y}}$ with probability
 609 one as $n \rightarrow \infty$. Since $\text{rank}(\Sigma_{\mathbf{Y} \circ \mathbf{X}}) = \text{rank}(\Sigma_{\mathbf{Y}\mathbf{X}} \Sigma_{\mathbf{X}}^{-1} \Sigma_{\mathbf{X}\mathbf{Y}}) = \text{rank}(\beta \Sigma_{\mathbf{X}} \beta^T) = d$, the last
 610 $(u - d)$ eigenvalues $\widehat{\omega}_j(\mathbf{G})$, $j = d + 1, \dots, u$, will equal to one with probability one as $n \rightarrow \infty$
 611 for any value of $\mathbf{G}$. Therefore, as $n \rightarrow \infty$,

$$\sup_{\mathbf{G} \in \mathcal{G}_{r,u}} \left\{ \sum_{i=d+1}^u \log[\widehat{\omega}_i(\mathbf{G})] \right\} \xrightarrow{p} 0. \quad (\text{B.1})$$

612 We next show that $\log |\mathbf{G}^T \mathbf{S}_{\mathbf{Y}}^{-1} \mathbf{G}|$ converges in probability to $\log |\mathbf{G}^T \Sigma_{\mathbf{Y}}^{-1} \mathbf{G}|$ uniformly in
 613 $\mathbf{G}$ by the following argument.

$$\begin{aligned} \delta(\mathbf{G}) &:= \sup_{\mathbf{G} \in \mathcal{G}_{r,u}} \{ \log |\mathbf{G}^T \mathbf{S}_{\mathbf{Y}}^{-1} \mathbf{G}| - \log |\mathbf{G}^T \Sigma_{\mathbf{Y}}^{-1} \mathbf{G}| \} \\ &= \sup_{\mathbf{G} \in \mathcal{G}_{r,u}} \{ \log |(\mathbf{G}^T \mathbf{S}_{\mathbf{Y}}^{-1} \mathbf{G})(\mathbf{G}^T \Sigma_{\mathbf{Y}}^{-1} \mathbf{G})^{-1}| \} \\ &= \sup_{\mathbf{G} \in \mathcal{G}_{r,u}} \{ \log |\mathbf{S}_{\mathbf{Y}}^{-1} \mathbf{G} (\mathbf{G}^T \Sigma_{\mathbf{Y}}^{-1} \mathbf{G})^{-1} \mathbf{G}^T|_0 \} \\ &= \sup_{\mathbf{G} \in \mathcal{G}_{r,u}} \left\{ \log |\Sigma_{\mathbf{Y}}^{1/2} \mathbf{S}_{\mathbf{Y}}^{-1} \Sigma_{\mathbf{Y}}^{1/2} \cdot \Sigma^{-1/2} \mathbf{G} (\mathbf{G}^T \Sigma_{\mathbf{Y}}^{-1} \mathbf{G})^{-1} \mathbf{G}^T \Sigma_{\mathbf{Y}}^{-1/2}|_0 \right\} \\ &= \sup_{\mathbf{G} \in \mathcal{G}_{r,u}} \left\{ \log |\Sigma_{\mathbf{Y}}^{1/2} \mathbf{S}_{\mathbf{Y}}^{-1} \Sigma_{\mathbf{Y}}^{1/2} \mathbf{P}_{\Sigma_{\mathbf{Y}}^{-1/2} \mathbf{G}}|_0 \right\}, \end{aligned}$$

614 where we use $|\cdot|_0$ to denote the product of the non-zero eigenvalues of a positive semi-definite
 615 matrix. We then can derive that

$$\delta(\mathbf{G}) = \sup_{\mathbf{G} \in \mathcal{G}_{r,u}} \left\{ \log |\mathbf{P}_{\Sigma_Y^{-1/2} \mathbf{G}} \Sigma_Y^{1/2} \mathbf{S}_Y^{-1} \Sigma_Y^{1/2} \mathbf{P}_{\Sigma_Y^{-1/2} \mathbf{G}}|_0 \right\}, \quad (\text{B.2})$$

616 where $\Sigma_Y^{1/2} \mathbf{S}_Y^{-1} \Sigma_Y^{1/2}$ was projected onto an u -dimensional subspace $\text{span}(\Sigma_Y^{-1/2} \mathbf{G})$. The quan-
 617 tity within $|\cdot|_0$ then has at most u nonzero eigenvalues. Because the projection matrix can not
 618 inflate the eigenvalues,

$$\delta(\mathbf{G}) \leq \sup_{\mathbf{G} \in \mathcal{G}_{r,u}} \left\{ \log |\Sigma_Y^{1/2} \mathbf{S}_Y^{-1} \Sigma_Y^{1/2}|_0 \right\}, \quad (\text{B.3})$$

619 which converges to zero in probability. Similarly, we can show that $\log |\mathbf{G}^T \mathbf{S}_{Y|\mathbf{X}} \mathbf{G}|$ converges
 620 in probability to $\log |\mathbf{G}^T \Sigma \mathbf{G}|$ uniformly in $\mathbf{G}$. Hence we have proved that the objective function
 621 $F_n(\mathbf{G}|d, u)$ in (3.3) converges in probability to $F(\mathbf{G}|u)$ uniformly in $\mathbf{G}$. The rest of the proof
 622 is similar to the proof of Proposition 4.2 in Cook et al. (2013) that

$$\begin{aligned} \log |\mathbf{G}^T \Sigma \mathbf{G}| + \log |\mathbf{G}^T \Sigma_Y^{-1} \mathbf{G}| &= \log |\mathbf{G}^T \Sigma \mathbf{G}| + \log |\mathbf{G}_0^T \Sigma_Y \mathbf{G}_0| \\ &= \log |\mathbf{G}^T \Sigma \mathbf{G}| + \log |\mathbf{G}_0^T (\Sigma + \beta \Sigma_X \beta^T) \mathbf{G}_0| \\ &\geq \log |\mathbf{G}^T \Sigma \mathbf{G}| + \log |\mathbf{G}_0^T \Sigma \mathbf{G}_0| \\ &\geq \log |\Sigma|, \end{aligned}$$

623 where the first inequality achieves its lower bound if $\text{span}(\beta) \subseteq \text{span}(\mathbf{G})$; and the second
 624 inequality achieves its lower bound if $\text{span}(\mathbf{G})$ is a reducing subspace of Σ . The uniqueness
 625 of the minimizer $\text{span}(\hat{\Gamma}) = \text{span}(\arg \min_{\mathbf{G}} F(\mathbf{G}|u))$ is guaranteed by the uniqueness of the
 626 envelope, which has dimension u .

627 C Proof for Proposition 6

628 For notation convenience, we define two covariance matrices $\mathbf{M}_B := \mathbf{B}^T (\mathbf{B} \Sigma_X \mathbf{B}^T)^{-1} \mathbf{B} \leq$
 629 Σ_X^{-1} and $\mathbf{M}_A := \mathbf{A} (\mathbf{A}^T \Sigma^{-1} \mathbf{A})^{-1} \mathbf{A}^T \leq \Sigma$. For any full row rank transformation $\mathbf{O} \in \mathbb{R}^{d \times q}$
 630 we could replace $\mathbf{A}$ by $\mathbf{A}\mathbf{O}$ and replace $\mathbf{B}$ by $\mathbf{O}\mathbf{B}$ without changing the value of $\mathbf{M}_A$ or $\mathbf{M}_B$.
 631 Also the projection matrices $\mathbf{P}_{\mathbf{A}(\Sigma^{-1})} = \mathbf{M}_A \Sigma^{-1}$ and $\mathbf{P}_{\mathbf{B}^T(\Sigma_X)} = \mathbf{M}_B \Sigma_X$.

C.1 Obtaining equation (4.3)

This result can be found in Anderson (1999) using canonical variables. We replicate the computation in our framework with details. Recall that the Fisher information is

$$\mathbf{J}_h = \begin{pmatrix} \mathbf{J}_\beta & 0 \\ 0 & \mathbf{J}_\Sigma \end{pmatrix} = \begin{pmatrix} \Sigma_{\mathbf{X}} \otimes \Sigma^{-1} & 0 \\ 0 & \frac{1}{2} \mathbf{E}_r^T (\Sigma^{-1} \otimes \Sigma^{-1}) \mathbf{E}_r \end{pmatrix}, \quad (\text{C.1})$$

where $\text{avar}(\sqrt{n}\hat{\beta}_{\text{OLS}}) = \mathbf{J}_\beta^{-1} = \Sigma_{\mathbf{X}}^{-1} \otimes \Sigma$.

By noticing $\mathbf{h}_1 = \text{vec}(\beta) = \text{vec}(\mathbf{A}\mathbf{B}) = (\mathbf{B}^T \otimes \mathbf{I}_r) \text{vec}(\mathbf{A}) = (\mathbf{I}_p \otimes \mathbf{A}) \text{vec}(\mathbf{B})$, we have

$$\mathbf{H} = \begin{pmatrix} \mathbf{B}^T \otimes \mathbf{I}_r & \mathbf{I}_p \otimes \mathbf{A} & 0 \\ 0 & 0 & \mathbf{I}_{r(r+1)/2} \end{pmatrix} := \begin{pmatrix} \mathbf{H}_1 & 0 \\ 0 & \mathbf{I}_{r(r+1)/2} \end{pmatrix}. \quad (\text{C.2})$$

Because of the similar block-diagonal structure in $\mathbf{J}_h = \text{diag}(\mathbf{J}_\beta, \mathbf{J}_\Sigma)$, we can get

$$\mathbf{H}(\mathbf{H}^T \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^T = \begin{pmatrix} \mathbf{H}_1(\mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1)^\dagger \mathbf{H}_1^T & 0 \\ 0 & \mathbf{J}_\Sigma^{-1} \end{pmatrix},$$

which means that $\beta = \mathbf{A}\mathbf{B}$ and Σ are orthogonal parameters in reduced-rank regression and the asymptotic covariance for $\text{vec}(\hat{\beta}_{\text{RR}})$ is $\mathbf{H}_1(\mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1)^\dagger \mathbf{H}_1^T$. Because $\mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1$ is not full rank under the reduced rank regression model, we can not use the block-matrix inversion formula. However, notice that asymptotic covariance $\mathbf{H}_1(\mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1)^\dagger \mathbf{H}_1^T$ depends only on the column space of $\mathbf{H}_1$, we thus could use any full row rank matrix $\mathbf{T}_1$ to get

$$\mathbf{H}_1(\mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1)^\dagger \mathbf{H}_1^T = \mathbf{H}_1 \mathbf{T}_1 (\mathbf{T}_1^T \mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1 \mathbf{T}_1)^\dagger \mathbf{T}_1^T \mathbf{H}_1. \quad (\text{C.3})$$

More specifically, we have each part

$$\mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1 = \begin{pmatrix} \mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T \otimes \Sigma^{-1} & \mathbf{B} \Sigma_{\mathbf{X}} \otimes \Sigma^{-1} \mathbf{A} \\ \Sigma_{\mathbf{X}} \mathbf{B}^T \otimes \mathbf{A}^T \Sigma^{-1} & \Sigma_{\mathbf{X}} \otimes \mathbf{A}^T \Sigma^{-1} \mathbf{A} \end{pmatrix} \quad (\text{C.4})$$

$$\mathbf{T}_1 = \begin{pmatrix} \mathbf{I}_{rd} & -(\mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T)^{-1} \mathbf{B} \Sigma_{\mathbf{X}} \otimes \mathbf{A} \\ 0 & \mathbf{I}_{pd} \end{pmatrix}$$

$$\mathbf{H}_1 \mathbf{T}_1 = \begin{pmatrix} \mathbf{B}^T \otimes \mathbf{I}_r & (\mathbf{I}_p - \mathbf{M}_B \Sigma_{\mathbf{X}}) \otimes \mathbf{A} \end{pmatrix}. \quad (\text{C.5})$$

where we have used $\mathbf{M}_B = \mathbf{B}^T (\mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T)^{-1} \mathbf{B}$ for notation convenience. Then,

$$\mathbf{T}_1^T \mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1 \mathbf{T}_1 = \begin{pmatrix} \mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T \otimes \Sigma^{-1} & 0 \\ 0 & (\Sigma_{\mathbf{X}} - \Sigma_{\mathbf{X}} \mathbf{M}_B \Sigma_{\mathbf{X}}) \otimes \mathbf{A}^T \Sigma^{-1} \mathbf{A} \end{pmatrix}.$$

645 To get the Moore-Penrose inverse of $\mathbf{T}_1^T \mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1 \mathbf{T}_1$, we first notice that it has rank $(p+r)d-d^2$
 646 and the only non-invertable part is $(\Sigma_{\mathbf{X}} - \Sigma_{\mathbf{X}} \mathbf{M}_B \Sigma_{\mathbf{X}})$ which causes rank deficiency of d^2 . The
 647 Moore-Penrose inverse of $\Sigma_{\mathbf{X}} - \Sigma_{\mathbf{X}} \mathbf{M}_B \Sigma_{\mathbf{X}}$ is obtained as follows by noticing $\mathbf{M}_B \Sigma_{\mathbf{X}} \mathbf{M}_B =$
 648 $\mathbf{M}_B$.

$$(\Sigma_{\mathbf{X}} - \Sigma_{\mathbf{X}} \mathbf{M}_B \Sigma_{\mathbf{X}})^\dagger = \Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B. \quad (\text{C.6})$$

649 Therefore,

$$(\mathbf{T}_1^T \mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1 \mathbf{T}_1)^\dagger = \begin{pmatrix} (\mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T)^{-1} \otimes \Sigma & 0 \\ 0 & (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes (\mathbf{A}^T \Sigma^{-1} \mathbf{A})^{-1} \end{pmatrix}. \quad (\text{C.7})$$

650 The asymptotic covariance $\text{avar}(\sqrt{n} \text{vec}(\hat{\beta}_{\text{RR}})) = \mathbf{H}_1 \mathbf{T}_1 (\mathbf{T}_1^T \mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1 \mathbf{T}_1)^\dagger \mathbf{T}_1^T \mathbf{H}_1$ is computed
 651 with (C.5),

$$\begin{aligned} \text{avar}(\sqrt{n} \text{vec}(\hat{\beta}_{\text{RR}})) &= \mathbf{M}_B \otimes \Sigma \\ &+ (\mathbf{I}_p - \mathbf{M}_B \Sigma_{\mathbf{X}}) (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) (\mathbf{I}_p - \mathbf{M}_B \Sigma_{\mathbf{X}}) \otimes \mathbf{M}_A \\ &= \mathbf{M}_B \otimes \Sigma + (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes \mathbf{M}_A \\ &= [\Sigma_{\mathbf{X}}^{-1} - (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B)] \otimes \Sigma + (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes \mathbf{M}_A \\ &= \Sigma_{\mathbf{X}}^{-1} \otimes \Sigma - (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes (\Sigma - \mathbf{M}_A), \end{aligned} \quad (\text{C.8})$$

652 then the equation (4.3) is derived from the following arguments.

$$\begin{aligned} \text{avar}(\sqrt{n} \text{vec}(\hat{\beta}_{\text{RR}})) &= \Sigma_{\mathbf{X}}^{-1} \otimes \Sigma - (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes (\Sigma - \mathbf{M}_A) \\ &= \Sigma_{\mathbf{X}}^{-1} \otimes \Sigma - [(\mathbf{I}_p - \mathbf{M}_B \Sigma_{\mathbf{X}}) \Sigma_{\mathbf{X}}^{-1}] \otimes [(\mathbf{I}_r - \mathbf{M}_A \Sigma^{-1}) \Sigma] \\ &= \Sigma_{\mathbf{X}}^{-1} \otimes \Sigma - [\mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})} \Sigma_{\mathbf{X}}^{-1}] \otimes [\mathbf{Q}_{\mathbf{A}(\Sigma)} \Sigma] \\ &= (\mathbf{I}_{pr} - \mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})} \otimes \mathbf{Q}_{\mathbf{A}(\Sigma)}) \Sigma_{\mathbf{X}}^{-1} \otimes \Sigma \\ &= (\mathbf{I}_{pr} - \mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})} \otimes \mathbf{Q}_{\mathbf{A}(\Sigma)}) \text{avar}[\sqrt{n} \text{vec}(\hat{\beta}_{\text{OLS}})]. \end{aligned}$$

C.2 Obtaining equation (4.4)

The Fisher information for $(\psi_1^T, \psi_2^T)^T = [\text{vec}^T(\mathbf{A}), \text{vec}^T(\mathbf{B})]^T$ is given in (C.4) as

$$\mathbf{H}_1^T \mathbf{J}_\beta \mathbf{H}_1 := \begin{pmatrix} \mathbf{J}_A & \mathbf{J}_{AB} \\ \mathbf{J}_{BA} & \mathbf{J}_B \end{pmatrix} = \begin{pmatrix} \mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T \otimes \Sigma^{-1} & \mathbf{B} \Sigma_{\mathbf{X}} \otimes \Sigma^{-1} \mathbf{A} \\ \Sigma_{\mathbf{X}} \mathbf{B}^T \otimes \mathbf{A} \Sigma^{-1} & \Sigma_{\mathbf{X}} \otimes \mathbf{A}^T \Sigma^{-1} \mathbf{A} \end{pmatrix}. \quad (\text{C.9})$$

If we known $\mathbf{A}$, then we could cross the first row and the first column, and hence

$$\text{avar}[\sqrt{n} \text{vec}(\hat{\mathbf{B}}_A)] = \mathbf{J}_B^{-1} = \Sigma_{\mathbf{X}}^{-1} \otimes (\mathbf{A}^T \Sigma^{-1} \mathbf{A})^{-1}. \quad (\text{C.10})$$

Similarly,

$$\text{avar}[\sqrt{n} \text{vec}(\hat{\mathbf{A}}_B)] = \mathbf{J}_A^{-1} = (\mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T)^{-1} \otimes \Sigma. \quad (\text{C.11})$$

Then by using the fact that $\text{vec}(\hat{\beta}_A) = (\mathbf{I}_p \otimes \mathbf{A}) \text{vec}(\hat{\mathbf{B}}_A)$ and that $\text{vec}(\hat{\beta}_B) = (\mathbf{B}^T \otimes \mathbf{I}_r) \text{vec}(\hat{\mathbf{A}}_B)$,

we have

$$\begin{aligned} \text{avar}[\sqrt{n} \text{vec}(\hat{\beta}_A)] &= \Sigma_{\mathbf{X}}^{-1} \otimes \mathbf{M}_A \\ \text{avar}[\sqrt{n} \text{vec}(\hat{\beta}_B)] &= \mathbf{M}_B \otimes \Sigma. \end{aligned}$$

By noticing $\mathbf{P}_{\mathbf{A}(\Sigma^{-1})} = \mathbf{M}_A \Sigma^{-1}$ and $\mathbf{P}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})} = \mathbf{M}_B \Sigma_{\mathbf{X}}$, we have

$$\begin{aligned} \text{avar}[\sqrt{n} \text{vec}(\hat{\beta}_A \mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})}^T)] &= [\mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})} \otimes \mathbf{I}_r] (\Sigma_{\mathbf{X}}^{-1} \otimes \mathbf{M}_A) [\mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})}^T \otimes \mathbf{I}_r] \\ &= [(\mathbf{I}_p - \mathbf{M}_B \Sigma_{\mathbf{X}}) \Sigma_{\mathbf{X}}^{-1} (\mathbf{I}_p - \Sigma_{\mathbf{X}}^T \mathbf{M}_B)] \otimes \mathbf{M}_A \\ &= (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes \mathbf{M}_A \\ \text{avar}[\sqrt{n} \text{vec}(\mathbf{Q}_{\mathbf{A}(\Sigma^{-1})} \hat{\beta}_B)] &= [\mathbf{I}_p \otimes \mathbf{Q}_{\mathbf{A}(\Sigma^{-1})}] (\mathbf{M}_B \otimes \Sigma) [\mathbf{I}_p \otimes \mathbf{Q}_{\mathbf{A}(\Sigma^{-1})}^T] \\ &= \mathbf{M}_B \otimes (\mathbf{I}_r - \mathbf{M}_A \Sigma^{-1}) \Sigma (\mathbf{I}_r - \Sigma^{-1} \mathbf{M}_A) \\ &= \mathbf{M}_B \otimes (\Sigma - \mathbf{M}_A). \end{aligned}$$

The proof of Proposition 6 is then completed by compare the above quantities with (C.8).

D Proof for Proposition 7

The role of η is analogous to $\mathbf{A}$ given Γ , thus we define $\mathbf{M}_\eta := \eta(\eta^T \Omega^{-1} \eta)^{-1} \eta^T \leq \Omega$. Note

that the projection matrices $\mathbf{P}_{\eta(\Omega^{-1})} = \mathbf{M}_\eta \Omega^{-1}$.

D.1 Explicit expression for $\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}})]$

By noticing $\mathbf{h}_1 = \text{vec}(\beta) = \text{vec}(\Gamma\eta\mathbf{B}) = (\eta^T\mathbf{B}^T \otimes \mathbf{I}_r)\text{vec}(\Gamma) = (\mathbf{B}^T \otimes \Gamma)\text{vec}(\eta) = (\mathbf{I}_p \otimes \Gamma\eta)\text{vec}(\mathbf{B})$, we have

$$\mathbf{R} = \begin{pmatrix} \mathbf{B}^T\eta^T \otimes \mathbf{I}_r & \mathbf{B}^T \otimes \Gamma & \mathbf{I}_p \otimes \Gamma\eta & 0 & 0 \\ 2\mathbf{C}_r(\Gamma\Omega \otimes \mathbf{I}_r - \Gamma \otimes \Gamma_0\Omega_0\Gamma_0^T) & 0 & 0 & \mathbf{C}_r(\Gamma \otimes \Gamma)\mathbf{E}_u & \mathbf{C}_r(\Gamma_0 \otimes \Gamma_0)\mathbf{E}_{r-u} \end{pmatrix}. \quad (\text{D.1})$$

The asymptotic covariance $\text{avar}(\sqrt{n}\mathbf{h}(\hat{\phi})) = \mathbf{R}(\mathbf{R}^T\mathbf{J}_h\mathbf{R})^\dagger\mathbf{R}^T = \tilde{\mathbf{R}}(\tilde{\mathbf{R}}^T\mathbf{J}_h\tilde{\mathbf{R}})^\dagger\tilde{\mathbf{R}}^T$ for any $\tilde{\mathbf{R}}$ such that $\mathbf{R} = \tilde{\mathbf{R}}\mathbf{T}$ for a full row rank matrix $\mathbf{T}$. We choose $\tilde{\mathbf{R}}$ to make $\tilde{\mathbf{R}}^T\mathbf{J}_h\tilde{\mathbf{R}}$ block-diagonal as follows.

$$\tilde{\mathbf{R}} = \begin{pmatrix} \mathbf{B}^T\eta^T \otimes \Gamma_0 & \mathbf{B}^T \otimes \Gamma & (\mathbf{I}_p - \mathbf{M}_B\boldsymbol{\Sigma}_X) \otimes \Gamma\eta & 0 & 0 \\ 2\mathbf{C}_r(\Gamma\Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0\Omega_0) & 0 & 0 & \mathbf{C}_r(\Gamma \otimes \Gamma)\mathbf{E}_u & \mathbf{C}_r(\Gamma_0 \otimes \Gamma_0)\mathbf{E}_{r-u} \end{pmatrix}, \quad (\text{D.2})$$

$$\mathbf{T} = \begin{pmatrix} \mathbf{I}_u \otimes \Gamma_0^T & 0 & 0 & 0 & 0 \\ \eta^T \otimes \Gamma^T & \mathbf{I}_{ud} & (\mathbf{B}\boldsymbol{\Sigma}_X\mathbf{B}^T)^{-1}\mathbf{B}\boldsymbol{\Sigma}_X \otimes \eta & 0 & 0 \\ 0 & 0 & \mathbf{I}_{pd} & 0 & 0 \\ 2\mathbf{C}_u(\Omega \otimes \Gamma^T) & 0 & 0 & \mathbf{I}_{\frac{1}{2}r(r+1)} & 0 \\ 0 & 0 & 0 & 0 & \mathbf{I}_{\frac{1}{2}(r-u)(r-u+1)} \end{pmatrix}. \quad (\text{D.3})$$

Next, we calculate $\tilde{\mathbf{R}}^T\mathbf{J}_h\tilde{\mathbf{R}}$ and verify that it is block-diagonal. We decompose $\tilde{\mathbf{G}}$ by it 2×5 blocks as $\tilde{\mathbf{R}} := (\tilde{\mathbf{G}}_1, \tilde{\mathbf{G}}_2, \tilde{\mathbf{G}}_3, \tilde{\mathbf{G}}_4, \tilde{\mathbf{G}}_5)$. We first calculate $\mathbf{J}_h\tilde{\mathbf{R}}$ and write down the 2×5 blocks by column:

$$\mathbf{J}_h\tilde{\mathbf{R}}_1 = \begin{pmatrix} \boldsymbol{\Sigma}_X\mathbf{B}^T\eta^T \otimes \Gamma_0\Omega_0^{-1} \\ \mathbf{E}_r^T(\Gamma \otimes \Gamma_0\Omega_0^{-1} - \Gamma\Omega^{-1} \otimes \Gamma_0) \end{pmatrix}, \quad (\text{D.4})$$

$$\mathbf{J}_h[\tilde{\mathbf{R}}_2, \tilde{\mathbf{R}}_3] = \begin{pmatrix} \boldsymbol{\Sigma}_X\mathbf{B}^T \otimes \Gamma\Omega^{-1} & (\boldsymbol{\Sigma}_X - \boldsymbol{\Sigma}_X\mathbf{M}_B\boldsymbol{\Sigma}_X) \otimes \Gamma\Omega^{-1}\eta \\ 0 & 0 \end{pmatrix}, \quad (\text{D.5})$$

$$\mathbf{J}_h[\tilde{\mathbf{G}}_4, \tilde{\mathbf{G}}_5] = \begin{pmatrix} 0 & 0 \\ \frac{1}{2}\mathbf{E}_r^T(\Gamma\Omega^{-1} \otimes \Gamma\Omega^{-1})\mathbf{E}_u & \frac{1}{2}\mathbf{E}_r^T(\Gamma_0\Omega_0^{-1} \otimes \Gamma_0\Omega_0^{-1})\mathbf{E}_{r-u} \end{pmatrix}. \quad (\text{D.6})$$

Then $\tilde{\mathbf{R}}^T\mathbf{J}_h\tilde{\mathbf{R}}$ equals to a block-diagonal matrix with five blocks: $\tilde{\mathbf{R}}_i^T\mathbf{J}_h\tilde{\mathbf{R}}_i$, $i = 1, \dots, 5$.

677 The explicit expressions are given as follows.

$$\begin{aligned}
\tilde{\mathbf{R}}_1^T \mathbf{J}_h \tilde{\mathbf{R}}_1 &= \boldsymbol{\eta} \mathbf{B} \boldsymbol{\Sigma}_X \mathbf{B}^T \boldsymbol{\eta}^T \otimes \boldsymbol{\Omega}_0^{-1} \\
&\quad + 2(\boldsymbol{\Omega} \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}_0^T - \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T) \mathbf{C}_r^T \mathbf{E}_r^T (\boldsymbol{\Gamma} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} - \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma}_0) \\
&= \boldsymbol{\eta} \mathbf{B} \boldsymbol{\Sigma}_X \mathbf{B}^T \boldsymbol{\eta}^T \otimes \boldsymbol{\Omega}_0^{-1} + \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1} - 2\mathbf{I}_{u(r-u)} + \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}_0, \\
\tilde{\mathbf{R}}_2^T \mathbf{J}_h \tilde{\mathbf{R}}_2 &= \mathbf{B} \boldsymbol{\Sigma}_X \mathbf{B}^T \otimes \boldsymbol{\Omega}^{-1}, \\
\tilde{\mathbf{R}}_3^T \mathbf{J}_h \tilde{\mathbf{R}}_3 &= \boldsymbol{\Sigma}_X \otimes \boldsymbol{\eta}^T \boldsymbol{\Omega}^{-1} \boldsymbol{\eta}, \\
\tilde{\mathbf{R}}_4^T \mathbf{J}_h \tilde{\mathbf{R}}_4 &= \mathbf{E}_u^T (\boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}^T) \mathbf{C}_r^T \cdot \frac{1}{2} \mathbf{E}_r^T (\boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1}) \mathbf{E}_u \\
&= \frac{1}{2} \mathbf{E}_u^T (\boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}^T) (\boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1}) \mathbf{E}_u \\
&= \frac{1}{2} \mathbf{E}_u^T (\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}) \mathbf{E}_u.
\end{aligned}$$

$$\begin{aligned}
\tilde{\mathbf{R}}_5^T \mathbf{J}_h \tilde{\mathbf{R}}_5 &= \mathbf{E}_{r-u}^T (\boldsymbol{\Gamma}_0^T \otimes \boldsymbol{\Gamma}_0^T) \mathbf{C}_r^T \cdot \frac{1}{2} \mathbf{E}_r^T (\boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1}) \mathbf{E}_{r-u} \\
&= \frac{1}{2} \mathbf{E}_{r-u}^T (\boldsymbol{\Omega}_0^{-1} \otimes \boldsymbol{\Omega}_0^{-1}) \mathbf{E}_{r-u}
\end{aligned}$$

678 Then the asymptotic covariance is

$$\text{avar}[\sqrt{n} \mathbf{h}(\hat{\boldsymbol{\phi}})] = \sum_{i=1}^5 \tilde{\mathbf{R}}_i (\tilde{\mathbf{R}}_i^T \mathbf{J}_h \tilde{\mathbf{R}}_i)^\dagger \tilde{\mathbf{R}}_i^T$$

679 We are only interested in the asymptotic covariance of $\text{avar}[\sqrt{n} \text{vec}(\hat{\boldsymbol{\beta}}_{\text{RE}})]$, which is the
680 upper left block of $\text{avar}[\sqrt{n} \mathbf{h}(\hat{\boldsymbol{\phi}})]$. And $\tilde{\mathbf{R}}_4 (\tilde{\mathbf{R}}_4^T \mathbf{J}_h \tilde{\mathbf{R}}_4)^\dagger \tilde{\mathbf{R}}_4^T$ and $\tilde{\mathbf{R}}_5 (\tilde{\mathbf{R}}_5^T \mathbf{J}_h \tilde{\mathbf{R}}_5)^\dagger \tilde{\mathbf{R}}_5^T$ have no
681 contribution to that. So we will focus our attention on the upper left block of $\tilde{\mathbf{R}}_i (\tilde{\mathbf{R}}_i^T \mathbf{J}_h \tilde{\mathbf{R}}_i)^\dagger \tilde{\mathbf{R}}_i^T$,
682 $i = 1, 2, 3$. The upper left block of $\tilde{\mathbf{R}}_1 (\tilde{\mathbf{R}}_1^T \mathbf{J}_h \tilde{\mathbf{R}}_1)^\dagger \tilde{\mathbf{R}}_1^T$ is

$$\begin{aligned}
&(\mathbf{B}^T \boldsymbol{\eta}^T \otimes \boldsymbol{\Gamma}_0) (\tilde{\mathbf{R}}_1^T \mathbf{J}_h \tilde{\mathbf{R}}_1)^\dagger (\boldsymbol{\eta} \mathbf{B} \otimes \boldsymbol{\Gamma}_0^T) \\
&= (\mathbf{B}^T \boldsymbol{\eta}^T \otimes \boldsymbol{\Gamma}_0) (\boldsymbol{\eta} \mathbf{B} \boldsymbol{\Sigma}_X \mathbf{B}^T \boldsymbol{\eta}^T \otimes \boldsymbol{\Omega}_0^{-1} + \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1} - 2\mathbf{I}_{u(r-u)} + \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}_0)^\dagger (\boldsymbol{\eta} \mathbf{B} \otimes \boldsymbol{\Gamma}_0^T).
\end{aligned}$$

683 The upper left block of $\tilde{\mathbf{R}}_2(\tilde{\mathbf{R}}_2^T \mathbf{J}_h \tilde{\mathbf{R}}_2)^\dagger \tilde{\mathbf{R}}_2^T$ is

$$\begin{aligned} & (\mathbf{B}^T \otimes \Gamma)(\mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T \otimes \Omega^{-1})^\dagger (\mathbf{B} \otimes \Gamma^T) \\ &= (\mathbf{B}^T \otimes \Gamma)[(\mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T)^{-1} \otimes \Omega](\mathbf{B} \otimes \Gamma^T) \\ &= \mathbf{M}_B \otimes \Gamma \Omega \Gamma^T. \end{aligned}$$

684 The upper left block of $\tilde{\mathbf{R}}_3(\tilde{\mathbf{R}}_3^T \mathbf{J}_h \tilde{\mathbf{R}}_3)^\dagger \tilde{\mathbf{R}}_3^T$ is

$$\begin{aligned} & [(\mathbf{I}_p - \mathbf{M}_B \Sigma_{\mathbf{X}}) \otimes \Gamma \eta](\Sigma_{\mathbf{X}} \otimes \eta^T \Omega^{-1} \eta)^\dagger [(\mathbf{I}_p - \Sigma_{\mathbf{X}} \mathbf{M}_B) \otimes \eta^T \Gamma^T] \\ &= (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes \Gamma \mathbf{M}_\eta \Gamma^T, \end{aligned}$$

685 where $\mathbf{M}_\eta = \eta(\eta^T \Omega^{-1} \eta)^{-1} \eta^T$.

686 Hence, the asymptotic covariance $\text{avar}[\sqrt{n} \text{vec}(\hat{\beta}_{\text{RE}})]$ equals to

$$\begin{aligned} & (\mathbf{B}^T \eta^T \otimes \Gamma_0)(\eta \mathbf{B} \Sigma_{\mathbf{X}} \mathbf{B}^T \eta^T \otimes \Omega_0^{-1} + \Omega \otimes \Omega_0^{-1} - 2\mathbf{I}_{u(r-u)} + \Omega^{-1} \otimes \Omega_0)^\dagger (\eta \mathbf{B} \otimes \Gamma_0^T) \\ &+ \mathbf{M}_B \otimes \Gamma \Omega \Gamma^T + (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes \Gamma \mathbf{M}_\eta \Gamma^T. \end{aligned} \tag{D.7}$$

687 D.2 Interpretation

688 The Fisher information matrix for $\hat{\phi}$ is simply $\mathbf{R}^T \mathbf{J}_h \mathbf{R}$:

$$\mathbf{R}^T \mathbf{J} \mathbf{R} := \begin{pmatrix} \mathbf{J}_\Gamma & \mathbf{J}_{\Gamma\eta} & \mathbf{J}_{\Gamma\mathbf{B}} & \mathbf{J}_{\Gamma\Omega} & 0 \\ \mathbf{J}_{\eta\Gamma} & \mathbf{J}_\eta & \mathbf{J}_{\eta\mathbf{B}} & 0 & 0 \\ \mathbf{J}_{\mathbf{B}\Gamma} & \mathbf{J}_{\mathbf{B}\eta} & \mathbf{J}_{\mathbf{B}} & 0 & 0 \\ \mathbf{J}_{\Omega\Gamma} & 0 & 0 & \mathbf{J}_\Omega & 0 \\ 0 & 0 & 0 & 0 & \mathbf{J}_{\Omega_0} \end{pmatrix}. \tag{D.8}$$

689

Each nonzero block is

$$\begin{aligned}
\mathbf{J}_\Gamma &= \boldsymbol{\eta} \mathbf{B} \boldsymbol{\Sigma}_\mathbf{X} \mathbf{B}^T \boldsymbol{\eta}^T \otimes \boldsymbol{\Sigma}^{-1} + \boldsymbol{\Omega} \otimes \boldsymbol{\Sigma}^{-1} + (\boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}) \mathbf{K}_{ru} \\
&\quad + \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T - 2\mathbf{I}_u \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Gamma}_0^T \\
\mathbf{J}_\eta &= \mathbf{B} \boldsymbol{\Sigma}_\mathbf{X} \mathbf{B}^T \otimes \boldsymbol{\Omega}^{-1} \\
\mathbf{J}_\mathbf{B} &= \boldsymbol{\Sigma}_\mathbf{X} \otimes \boldsymbol{\eta}^T \boldsymbol{\Omega}^{-1} \boldsymbol{\eta} \\
\mathbf{J}_\Omega &= \frac{1}{2} \mathbf{E}_u^T (\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}) \mathbf{E}_u \\
\mathbf{J}_{\Omega_0} &= \frac{1}{2} \mathbf{E}_{r-u}^T (\boldsymbol{\Omega}_0^{-1} \otimes \boldsymbol{\Omega}_0^{-1}) \mathbf{E}_{r-u} \\
\mathbf{J}_{\Gamma\eta} &= \boldsymbol{\eta} \mathbf{B} \boldsymbol{\Sigma}_\mathbf{X} \mathbf{B}^T \otimes \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \\
\mathbf{J}_{\Gamma\mathbf{B}} &= \boldsymbol{\eta} \mathbf{B} \boldsymbol{\Sigma}_\mathbf{X} \otimes \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\eta} \\
\mathbf{J}_{\Gamma\Omega} &= (\mathbf{I}_u \otimes \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1}) \mathbf{E}_u \\
\mathbf{J}_{\eta\mathbf{B}} &= \mathbf{B} \boldsymbol{\Sigma}_\mathbf{X} \otimes \boldsymbol{\Omega}^{-1} \boldsymbol{\eta}
\end{aligned}$$

690 D.2.1 Asymptotic covariance when $\boldsymbol{\eta}$ and $\mathbf{B}$ are known

691 The asymptotic covariance for $\text{vec}(\hat{\boldsymbol{\Gamma}}_{\boldsymbol{\eta}, \mathbf{B}})$ is

$$\text{avar}[\sqrt{n} \text{vec}(\hat{\boldsymbol{\Gamma}}_{\boldsymbol{\eta}, \mathbf{B}})] = (\mathbf{J}_\Gamma - \mathbf{J}_{\Gamma\Omega} \mathbf{J}_\Omega^{-1} \mathbf{J}_{\Omega\Gamma})^{-1}.$$

692 Follow Cook et al. (2010), we have

$$\begin{aligned}
\text{avar}[\sqrt{n} \text{vec}(\hat{\boldsymbol{\Gamma}}_{\boldsymbol{\eta}, \mathbf{B}})] &= [\boldsymbol{\eta} \mathbf{B} \boldsymbol{\Sigma}_\mathbf{X} \mathbf{B}^T \boldsymbol{\eta}^T \otimes \boldsymbol{\Sigma}^{-1} + \boldsymbol{\Omega} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T - 2\mathbf{I}_u \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Gamma}_0^T \\
&\quad + \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T]^\dagger,
\end{aligned}$$

693 and by replacing $\boldsymbol{\eta} \mathbf{B} \rightarrow \boldsymbol{\eta}$ in Cook et al. (2010), it is easy to obtain the following results

$$\text{avar}[\sqrt{n} \text{vec}(\mathbf{Q}_\Gamma \hat{\boldsymbol{\beta}}_{\boldsymbol{\eta}, \mathbf{B}})] = \left[\tilde{\mathbf{R}}_1 (\tilde{\mathbf{R}}_1^T \mathbf{J}_h \tilde{\mathbf{R}}_1)^\dagger \tilde{\mathbf{R}}_1^T \right]_{11},$$

694 where $\square_{11}$ means the upper left block of a block-wise matrix. The above equality explains the
695 contribution from the first column block of $\tilde{\mathbf{R}}$, which is the first term in (D.7).

Therefore, (D.7) can be written as

$$\begin{aligned} \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\text{RE}})] &= \text{avar}[\sqrt{n}\text{vec}(\mathbf{Q}_{\Gamma}\hat{\beta}_{\eta,\mathbf{B}})] \\ &\quad + \mathbf{M}_B \otimes \Gamma\Omega\Gamma^T + (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes \Gamma\mathbf{M}_{\eta}\Gamma^T \end{aligned} \quad (\text{D.9})$$

$$= \text{avar}[\sqrt{n}\text{vec}(\mathbf{Q}_{\Gamma}\hat{\beta}_{\eta,\mathbf{B}})] + \text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma})], \quad (\text{D.10})$$

where the last equality follows from the asymptotic covariance of $\text{vec}(\hat{\beta}_{\text{RR}})$ in (C.8) and from Lemma 1 that $\hat{\beta}_{\Gamma}$ is Γ times the reduced-rank regression estimator of regression $\Gamma^T \mathbf{Y}$ on $\mathbf{X}$.

D.2.2 Asymptotic covariance when Γ and $\mathbf{B}$ are known

The asymptotic covariance for $\text{vec}(\hat{\eta}_{\Gamma,\mathbf{B}})$ is

$$\text{avar}[\sqrt{n}\text{vec}(\hat{\eta}_{\Gamma,\mathbf{B}})] = \mathbf{J}_{\eta}^{-1} = (\mathbf{B}\Sigma_{\mathbf{X}}\mathbf{B}^T)^{-1} \otimes \Omega. \quad (\text{D.11})$$

Notice that $\text{vec}(\hat{\beta}_{\Gamma,\mathbf{B}}) = \text{vec}(\Gamma\hat{\eta}_{\Gamma,\mathbf{B}}\mathbf{B}) = (\mathbf{B}^T \otimes \Gamma)\text{vec}(\hat{\eta}_{\Gamma,\mathbf{B}})$, we have

$$\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma,\mathbf{B}})] = \mathbf{M}_B \otimes \Gamma\Omega\Gamma^T. \quad (\text{D.12})$$

D.2.3 Asymptotic covariance when Γ and η are known

The asymptotic covariance for $\text{vec}(\hat{\mathbf{B}}_{\Gamma,\eta})$ is

$$\text{avar}[\sqrt{n}\text{vec}(\hat{\mathbf{B}}_{\Gamma,\eta})] = \mathbf{J}_{\mathbf{B}}^{-1} = \Sigma_{\mathbf{X}}^{-1} \otimes (\eta^T \Omega^{-1} \eta)^{-1}. \quad (\text{D.13})$$

Notice that $\text{vec}(\hat{\beta}_{\Gamma,\eta}) = \text{vec}(\Gamma\eta\hat{\mathbf{B}}_{\Gamma,\eta}) = (\mathbf{I}_p \otimes \Gamma\eta)\text{vec}(\hat{\mathbf{B}}_{\Gamma,\eta})$, we have

$$\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma,\eta})] = \Sigma_{\mathbf{X}}^{-1} \otimes \Gamma\mathbf{M}_{\eta}\Gamma^T. \quad (\text{D.14})$$

$$\text{avar}[\sqrt{n}\text{vec}(\hat{\beta}_{\Gamma,\eta}\mathbf{Q}_{\mathbf{B}^T(\Sigma_{\mathbf{X}})}^T)] = (\Sigma_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes \Gamma\mathbf{M}_{\eta}\Gamma^T. \quad (\text{D.15})$$

D.2.4 Decomposition

Finally, plugging (D.12) and (D.15) into (D.9), we have proven this Proposition.

E Proof for Corollary 1

By noticing $\mathbf{A} = \Gamma\boldsymbol{\eta}$, we can write

$$\begin{aligned}\mathbf{P}_{\mathbf{A}(\boldsymbol{\Sigma}^{-1})} &= \Gamma\boldsymbol{\eta}(\boldsymbol{\eta}^T\Gamma^T\boldsymbol{\Sigma}^{-1}\Gamma\boldsymbol{\eta})^{-1}\boldsymbol{\eta}^T\Gamma^T\boldsymbol{\Sigma}^{-1} = \Gamma\mathbf{P}_{\boldsymbol{\eta}(\boldsymbol{\Omega}^{-1})}\Gamma^T. \\ \Gamma\mathbf{M}_{\boldsymbol{\eta}}\Gamma^T &= \Gamma\mathbf{P}_{\boldsymbol{\eta}(\boldsymbol{\Omega}^{-1})}\boldsymbol{\Omega}\Gamma^T = \Gamma\mathbf{P}_{\boldsymbol{\eta}(\boldsymbol{\Omega}^{-1})}\Gamma^T \cdot \Gamma\boldsymbol{\Omega}\Gamma = \mathbf{P}_{\mathbf{A}(\boldsymbol{\Sigma}^{-1})}\mathbf{P}_{\Gamma}\boldsymbol{\Sigma}.\end{aligned}$$

Then, from (D.9), we have

$$\begin{aligned}\text{avar}[\sqrt{n}\text{vec}(\hat{\boldsymbol{\beta}}_{\Gamma})] &= \mathbf{M}_B \otimes \Gamma\boldsymbol{\Omega}\Gamma^T + (\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} - \mathbf{M}_B) \otimes \Gamma\mathbf{M}_{\boldsymbol{\eta}}\Gamma^T \\ &= \mathbf{P}_{\mathbf{B}^T(\boldsymbol{\Sigma}_{\mathbf{X}})}\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \mathbf{P}_{\Gamma}\boldsymbol{\Sigma} + \mathbf{Q}_{\mathbf{B}^T(\boldsymbol{\Sigma}_{\mathbf{X}})}\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \mathbf{P}_{\mathbf{A}(\boldsymbol{\Sigma}^{-1})}\mathbf{P}_{\Gamma}\boldsymbol{\Sigma} \\ &= \{\mathbf{P}_{\mathbf{B}^T(\boldsymbol{\Sigma}_{\mathbf{X}})} \otimes \mathbf{I}_r + \mathbf{Q}_{\mathbf{B}^T(\boldsymbol{\Sigma}_{\mathbf{X}})} \otimes \mathbf{P}_{\mathbf{A}(\boldsymbol{\Sigma}^{-1})}\} \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \mathbf{P}_{\Gamma}\boldsymbol{\Sigma} \\ &= \{\mathbf{I}_{pr} - \mathbf{Q}_{\mathbf{B}^T(\boldsymbol{\Sigma}_{\mathbf{X}})} \otimes \mathbf{Q}_{\mathbf{A}(\boldsymbol{\Sigma}^{-1})}\} \boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \mathbf{P}_{\Gamma}\boldsymbol{\Sigma} \\ &= \{\mathbf{I}_{pr} - \mathbf{Q}_{\mathbf{B}^T(\boldsymbol{\Sigma}_{\mathbf{X}})} \otimes \mathbf{Q}_{\mathbf{A}(\boldsymbol{\Sigma}^{-1})}\} (\mathbf{I}_p \otimes \mathbf{P}_{\Gamma})\boldsymbol{\Sigma}_{\mathbf{X}}^{-1} \otimes \boldsymbol{\Sigma}.\end{aligned}$$

F Proof for Proposition 8

From Proposition 2 and Proposition 4, we see that the minimizer $\hat{\mathbf{h}}_{\text{RE}} = \mathbf{h}(\hat{\boldsymbol{\phi}})$ of $\mathcal{F}(\mathbf{h}(\boldsymbol{\phi}), \hat{\mathbf{h}}_{\text{OLS}})$ is Fisher consistent. 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 first can check that $\mathcal{F}(\mathbf{h}, \hat{\mathbf{h}}_{\text{OLS}})$ satisfies: (1) $\mathcal{F}(\mathbf{h}, \hat{\mathbf{h}}_{\text{OLS}}) \geq 0$ for all $\hat{\mathbf{h}}_{\text{OLS}}$ and $\mathbf{h}$; (2) $\mathcal{F}(\mathbf{h}, \hat{\mathbf{h}}_{\text{OLS}}) = 0$ if and only if $\hat{\mathbf{h}}_{\text{OLS}} = \mathbf{h}$; and (3) $\mathcal{F}(\mathbf{h}, \hat{\mathbf{h}}_{\text{OLS}})$ is twice continuously differentiable in $\mathbf{h}$ and $\hat{\mathbf{h}}_{\text{OLS}}$. Recall from Section 4.2 that we use the subscript 0 to emphasize the true parameter: $\mathbf{h}_0$ and $\boldsymbol{\phi}_0$ correspond to the true distribution of $\boldsymbol{\epsilon}$. Then $\hat{\mathbf{h}}_{\text{OLS}}$ is $\sqrt{n}$ -consistent for $\mathbf{h}_0$. Notice that $\hat{\mathbf{h}}_{\text{OLS}}$ is a smooth function of the sample covariance matrices which converges in distribution to the population covariance matrices, then by the delta method we know $\sqrt{n}(\hat{\mathbf{h}}_{\text{OLS}} - \mathbf{h}_0) \rightarrow N(0, \mathbf{K})$, for some positive definite covariance $\mathbf{K}$. Using Shapiro's (1986) Proposition 3.1 and Proposition 4.1, we will have $\sqrt{n}$ -consistency results for $\hat{\mathbf{h}}_{\text{RE}} = \mathbf{h}(\hat{\boldsymbol{\phi}})$ as shown in Proposition 8.