

NONLINEAR INTERACTION DETECTION THROUGH MODEL-BASED SUFFICIENT DIMENSION REDUCTION

Guoliang Fan^{1,2}, Liping Zhu² and Shujie Ma³

¹*Shanghai Maritime University*, ²*Renmin University of China* and

³*University of California, Riverside*

Abstract: In this paper we propose an efficient model-based sufficient dimension reduction method to detect interactions. We introduce a new class of multivariate adaptive varying index models (MAVIM) to investigate nonlinear interaction effects of the grouped covariates on multivariate response variables. Grouping the covariates through linear combinations in the MAVIM accommodates weak individual interaction effects as long as their joint interaction effects are strong enough to be detectable. We estimate the joint interaction effects by a weighted-profile least squares method that is numerically stable and computationally fast. The resultant profile least squares estimate is root- n consistent and asymptotically normal. We discuss how to choose an optimal weight to improve the estimation efficiency. We determine the structural dimension with a BIC-type criterion, and establish its consistency. The effectiveness of our proposal is illustrated through simulation studies and an analysis of Framingham heart study.

Key words and phrases: Central mean subspace, dimension determination, high-dimensionality, interaction detection, sufficient dimension reduction.

1. Introduction

With the advance of information technology, high-dimensional data are effectively collected at a low cost in many scientific fields. Regression analysis is perhaps one of the most popular tools that help us gain insight into the relationship between two sets of high-dimensional variables. Suppose $\mathbf{x} = (X_1, \dots, X_p)^T \in \mathbb{R}^p$ is the covariate vector and $\mathbf{y} = (Y_1, \dots, Y_r)^T \in \mathbb{R}^r$ is the response vector. In general, the goal of regression analysis is to study how the conditional mean function of $\mathbf{y}$, denoted by $E(\mathbf{y}|\mathbf{x})$, varies with $\mathbf{x}$. Nonparametric regression is a flexible and effective approach to estimating the conditional mean. However, it suffers from the “curse of dimensionality” when the dimension of $\mathbf{x}$ increases. To reduce the covariate dimension, sufficient dimension reduction (Cook (1998)) is an effective tool that combines the concept of sufficiency with the idea of dimension reduction. It achieves the goal of dimension reduction through replacing the

high-dimensional $\mathbf{x}$ with d linear combinations, denoted as $(\boldsymbol{\alpha}^\top \mathbf{x})$ for $\boldsymbol{\alpha} \in \mathbb{R}^{p \times d}$, without losing information of $(\mathbf{y}|\mathbf{x})$. In other words, replacing $\mathbf{x}$ with $(\boldsymbol{\alpha}^\top \mathbf{x})$ is “sufficient” in the sense that $E(\mathbf{y}|\mathbf{x}) = E(\mathbf{y}|\boldsymbol{\alpha}^\top \mathbf{x})$. By the very purpose of dimension reduction, d is assumed to be small and hence estimating $E(\mathbf{y}|\boldsymbol{\alpha}^\top \mathbf{x})$ via nonparametric smoothing is straightforward. In practice, d is unknown and needs to be estimated from the data. This differentiates the dimension reduction model $E(\mathbf{y}|\mathbf{x}) = E(\mathbf{y}|\boldsymbol{\alpha}^\top \mathbf{x})$ from the single- or multiple-index models in the literature Ichimura (1993); Carroll et al. (1997); Ma and Zhu (2014). Seeking for an appropriate $\boldsymbol{\alpha}$ such that $E(\mathbf{y}|\mathbf{x}) = E(\mathbf{y}|\boldsymbol{\alpha}^\top \mathbf{x})$, is the central goal of sufficient dimension reduction when estimating $E(\mathbf{y}|\mathbf{x})$ is concerned. Because $\boldsymbol{\alpha}$ is not identifiable, the space spanned by $\boldsymbol{\alpha}$, denoted as $\text{span}(\boldsymbol{\alpha})$, with the minimal column dimension, is the parameter of primary interest and referred to as the central mean space (Cook and Li (2002)) in sufficient dimension reduction.

In the present article we consider the problem of sufficient dimension reduction in the presence of high-dimensional controlling variables $\mathbf{z} = (Z_1, \dots, Z_q)^\top$. This falls into the framework of partial mean dimension reduction (Li, Cook and Chiaromonte (2003)), which seeks a $p \times d_0$ matrix $\boldsymbol{\beta}$ such that $E(\mathbf{y}|\mathbf{x}, \mathbf{z}) = E(\mathbf{y}|\boldsymbol{\beta}^\top \mathbf{x}, \mathbf{z})$. Similar to $\boldsymbol{\alpha}$, $\boldsymbol{\beta}$ is not identifiable either. Therefore, our primary goal is to seek for the minimal column space of $\boldsymbol{\beta}$, denoted by $\text{span}(\boldsymbol{\beta})$, such that $E(\mathbf{y}|\mathbf{x}, \mathbf{z}) = E(\mathbf{y}|\boldsymbol{\beta}^\top \mathbf{x}, \mathbf{z})$. Following the convention of sufficient dimension reduction (Li, Cook and Chiaromonte (2003)), we refer to the column space of $\boldsymbol{\beta}$, denoted by $\text{span}(\boldsymbol{\beta})$, such that $E(\mathbf{y}|\mathbf{x}, \mathbf{z}) = E(\mathbf{y}|\boldsymbol{\beta}^\top \mathbf{x}, \mathbf{z})$ as the partial central mean dimension reduction subspace. Its column dimension, denoted by d_0 , is referred to as the structural dimension of $\text{span}(\boldsymbol{\beta})$. Existing methods require $r = q = 1$. Specifically, if $\mathbf{z}$ and $\mathbf{y}$ are univariate and $\mathbf{z}$ is categorical taking a small number of values, Li, Cook and Chiaromonte (2003) suggested applying an existing sufficient dimension reduction method to $E(\mathbf{y}|\mathbf{x}, \mathbf{z} = \mathbf{z}_0)$ within each category of $\mathbf{z}$, say, $\mathbf{z} = \mathbf{z}_0$, and sum up all estimates to form an estimate of $\text{span}(\boldsymbol{\beta})$; if $\mathbf{y}$ is univariate and $\mathbf{z}$ is continuous and low dimensional, Feng et al. (2013) suggested discretizing $\mathbf{z}$ into a series of binary variables and then applying Li, Cook and Chiaromonte (2003)’s method to form an estimate of $\text{span}(\boldsymbol{\beta})$. Hilafu and Wu (2017) suggest recovering $\text{span}(\boldsymbol{\beta})$ through regressing $\tilde{\mathbf{y}} \stackrel{\text{def}}{=} (\mathbf{y}^\top, \mathbf{z}^\top)^\top$ onto $\mathbf{x}$. We consider a more general situation in which $\mathbf{z}$ and $\mathbf{y}$ are allowed to be multivariate and components in $\mathbf{z}$ can be either categorical or continuous. Such considerations are motivated by the Framingham Heart Study, where 304 subjects were collected to evaluate the effects of physical exercises on the blood pressures. The Framingham Heart Study Data were downloaded from NCBI db-

GaP with an IRB number HS-11-159. The systolic (Y_1) and the diastolic blood pressures (Y_2), and several measurements of such physical exercises, as hours for heavy, moderate and light activities per day, denoted by Z_1, Z_2 and Z_3 , respectively, were recorded for each subject. There is a public health concern regarding the developmental effects resulting from the lack of physical exercises. It is believed that a moderate amount of physical exercise helps to relieve stress, and hence are beneficial to controlling for the blood pressures. Therefore, our goal is to investigate whether or not the physical exercises, Z_1, Z_2 and Z_3 , affect the blood pressures. The blood pressures are also relevant to the degree of obesity; they are measured by weight (X_1), height (X_2), bi-deltoid girth (X_3), right arm girth-upper third (X_4), waist girth (X_5), hip girth (X_6) and thigh girth (X_7).

In general, to understand how the conditional mean functions of $\mathbf{y} = (Y_1, \dots, Y_r)^T$ vary with $\mathbf{x} = (X_1, \dots, X_p)^T$ in the presence of $\mathbf{z} = (Z_1, \dots, Z_q)^T$, and to simultaneously accommodate the interaction effects between $\mathbf{x}$ and $\mathbf{z}$, we consider a multivariate adaptive varying index model (MAVIM for short),

$$E(\mathbf{y}|\mathbf{x}, \mathbf{z}) = \sum_{k=1}^q \mathbf{m}_k(\beta^T \mathbf{x}) Z_k, \text{ or equivalently,} \quad (1.1)$$

$$E(Y_l|\mathbf{x}, \mathbf{z}) = \sum_{k=1}^q m_{kl}(\beta^T \mathbf{x}) Z_k, \text{ for } l = 1, \dots, r, \quad (1.2)$$

where $\mathbf{m}_k(\beta^T \mathbf{x}) = (m_{k1}(\beta^T \mathbf{x}), \dots, m_{kr}(\beta^T \mathbf{x}))^T$, $k = 1, \dots, q$, β is a $p \times d_0$ matrix with an unknown d_0 . All $\mathbf{m}_k$'s, β and d_0 have to be estimated from data. All $\mathbf{m}_k$'s share an identical β to ensure that $\text{span}(\beta)$ is identifiable (Zhu and Zhong (2015)). We group $\mathbf{x}$ through $(\beta^T \mathbf{x})$ to augment the interaction effects between $\mathbf{x}$ and Z_k , a useful strategy if the joint interaction effects between $\mathbf{x}$ and Z_k are strong enough but the individual interaction effects between X_i and Z_k are too weak to be detectable.

The functions $\mathbf{m}_k$ accommodate nonlinear interaction effects between $\mathbf{x}$ and Z_k (Ma et al. (2011)). To be precise, if $\mathbf{m}_k$ is constant, $\mathbf{x}$ does not interact with Z_k . In addition, if the i -th row of β is zero, X_i does not interact with $\mathbf{z}$. In other words, the i -th row of β describes the joint interaction effects between X_i and $\mathbf{z}$. Model (1.1), or equivalently, (1.2), characterizes the interaction effects between $\mathbf{x}$ and $\mathbf{z}$. We achieve the goal of dimension reduction through replacing $\mathbf{x}$ with $(\beta^T \mathbf{x})$ in the presence of $\mathbf{z}$ and such a reduction is sufficient in the sense that (1.1), or equivalently, (1.2), holds almost surely. Because we allow for a general d_0 , that all $\mathbf{m}_k$'s share a common β in model (1.1) is a necessary assumption for the purposes of identifiability.

Our goal is to estimate and make inference on $\text{span}(\beta)$. Towards this goal, we recast the problem of estimating $\text{span}(\beta)$ to the problem of estimating an identifiable basis matrix β . We propose a weighted-profile least squares estimation procedure for β in which each $\mathbf{m}_k$ is approximated by the local linear regression (Fan and Gijbels (1996)). An important methodological merit of our approach is the ease of simultaneously approximating multiple nonparametric functions to create a single objective function for β , so that the profile least squares estimation can be established in a straightforward manner. The resultant estimate of β is root- n consistent and asymptotically normal. We devise a Wald chi-square testing procedure for β based on the asymptotic distribution of the profile least squares estimate. In the Framingham Heart Study, the systolic (Y_1) and the diastolic blood pressures (Y_2) are highly correlated and hence are considered jointly to improve the efficiency of the profile least squares estimate. We also discuss how to choose an optimal weight to improve efficiency of estimating β , a particular basis matrix of $\text{span}(\beta)$. Because the model structure is assumed in (1.1), we refer to our proposal as a model-based sufficient dimension reduction method.

This article is organized as follows. Section 2 introduces the weighted-profile least squares estimation and presents asymptotic properties of the proposed estimate. We also discuss how to choose an optimal weight matrix. In Section 3 we evaluate finite sample properties of the proposed estimation and inference procedures via comprehensive simulation studies. We also illustrate the usefulness of our proposals through an analysis of the Framingham Heart Study. Some concluding remarks are given in Section 4. All technical details and additional discussions are given in an online supplement.

2. Methodology Development

We seek β with the minimal column dimension d_0 such that (1.1) holds, then replacing $\mathbf{x}$ with $(\beta^T \mathbf{x})$ is sufficient to describe how $E(\mathbf{y}|\mathbf{x}, \mathbf{z})$ varies with $\mathbf{x}$ and $\mathbf{z}$. We further assume that the upper $d_0 \times d_0$ submatrix of β is an identity matrix. In other words, $\beta = (\mathbf{I}_{d_0 \times d_0}, \beta_{-d_0}^T)^T$, where $\mathbf{I}_{d_0 \times d_0}$ is a $d_0 \times d_0$ identity matrix and β_{-d_0} is a $(p - d_0) \times d_0$ matrix composed of the lower $(p - d_0)$ rows of β . In single index models (Ichimura (1993)) where $d_0 = 1$, there are two options to ensure that β is identifiable. The first is to restrict that β is of unit-length and the first entry of β is strictly positive. The second is to simply set the first entry of β to be 1 and thus all other entries are free parameters. These options are, in

spirit, equivalent. Requiring the upper $d_0 \times d_0$ submatrix of β to be an identity matrix is an extension of the second option. Such a parameterization is also used by Ma and Zhu (2013) and implies that the first d_0 covariates of $\mathbf{x}$ contribute to model (1.1). If this is not the case, one can always rotate the order of the entries in $\mathbf{x}$ to guarantee that the first d_0 components of $\mathbf{x}$ are useful. Through parameterizing $\text{span}(\beta)$ with a particular basis matrix $\beta = (\mathbf{I}_{d_0 \times d_0}, \beta_{-d_0}^T)^T$, we convert the problem of estimating β into a problem of estimating the $(p-d_0) \times d_0$ matrix β_{-d_0} , the free parameters in β .

Because the structural dimension d_0 of $\text{span}(\beta)$ is unknown a priori, we illustrate our proposed estimation procedure with a working dimension d . In this section we propose a profile least squares method to estimate β , or equivalently, β_{-d} . Let $\mathbf{x}_d = (X_1, \dots, X_d)^T$ and $\mathbf{x}_{-d} = (X_{d+1}, \dots, X_p)^T$. Hence, $\beta^T \mathbf{x} = \mathbf{x}_d + \beta_{-d}^T \mathbf{x}_{-d}$ and $\mathbf{m}_k(\beta^T \mathbf{x}) = \mathbf{m}_k(\mathbf{x}_d + \beta_{-d}^T \mathbf{x}_{-d})$. Suppose that $\{(\mathbf{x}_i, \mathbf{z}_i, \mathbf{y}_i), i = 1, \dots, n\}$ is a random sample of $(\mathbf{x}, \mathbf{z}, \mathbf{y})$ that follows model (1.1). For a given β , we estimate $\mathbf{m}_k$, $k = 1, \dots, q$, using the local linear approximation (Fan and Gijbels (1996)). Specifically, for $\mathbf{U} = (\beta^T \mathbf{x})$ in a small neighborhood of $\mathbf{u}$, one can approximate $\mathbf{m}_k(\mathbf{U}) \approx \mathbf{m}_k(\mathbf{u}) + \mathbf{m}_k^{(1)}(\mathbf{u})(\mathbf{U} - \mathbf{u}) \stackrel{\text{def}}{=} \mathbf{a}_k + \mathbf{B}_k(\mathbf{U} - \mathbf{u})$, for $k = 1, \dots, q$, where $\mathbf{m}_k^{(1)}(\mathbf{u})$, for $k = 1, \dots, q$, denotes the first derivative of $\mathbf{m}_k(\mathbf{u})$ with respect to $\mathbf{u}$ and hence all of them are $r \times d$ matrices. The local linear estimators for $\mathbf{m}_k(\mathbf{u})$ and $\mathbf{m}_k^{(1)}(\mathbf{u})$ are defined as $\hat{\mathbf{m}}_k(\mathbf{u}, \beta) = \hat{\mathbf{a}}_k$ and $\hat{\mathbf{m}}_k^{(1)}(\mathbf{u}, \beta) = \hat{\mathbf{B}}_k$ at the fixed point β , where $\{(\hat{\mathbf{a}}_k, \hat{\mathbf{B}}_k), k = 1, \dots, q\}$ minimize the sum of the weighted least squares

$$\sum_{i=1}^n \left[\mathbf{y}_i - \sum_{k=1}^q \{ \mathbf{a}_k + \mathbf{B}_k(\beta^T \mathbf{x}_i - \mathbf{u}) \} Z_{ik} \right]^2 K_h(\beta^T \mathbf{x}_i - \mathbf{u}),$$

where $K_h(\cdot) = K(\cdot/h)/h^d$ is a product of d univariate kernel functions and h is a bandwidth. By some straightforward algebraic calculations, we derive that

$$\begin{aligned} & \left\{ \hat{\mathbf{m}}_1(\mathbf{u}, \beta), \dots, \hat{\mathbf{m}}_q(\mathbf{u}, \beta), h\hat{\mathbf{m}}_1^{(1)}(\mathbf{u}, \beta), \dots, h\hat{\mathbf{m}}_q^{(1)}(\mathbf{u}, \beta) \right\}^T \\ &= \mathbf{S}_n^{-1}(\mathbf{u}, \beta) \boldsymbol{\xi}_n(\mathbf{u}, \beta), \end{aligned} \quad (2.1)$$

where $\mathbf{S}_n(\mathbf{u}, \beta) \stackrel{\text{def}}{=} \begin{pmatrix} \mathbf{S}_{n0}(\mathbf{u}, \beta) & \mathbf{S}_{n1}^T(\mathbf{u}, \beta) \\ \mathbf{S}_{n1}(\mathbf{u}, \beta) & \mathbf{S}_{n2}(\mathbf{u}, \beta) \end{pmatrix}$ and $\boldsymbol{\xi}_n(\mathbf{u}, \beta) \stackrel{\text{def}}{=} \begin{pmatrix} \boldsymbol{\xi}_{n0}(\mathbf{u}, \beta) \\ \boldsymbol{\xi}_{n1}(\mathbf{u}, \beta) \end{pmatrix}$, with

$$\mathbf{S}_{nj}(\mathbf{u}, \beta) \stackrel{\text{def}}{=} n^{-1} \sum_{i=1}^n \mathbf{z}_i \mathbf{z}_i^T \otimes \left(\frac{\beta^T \mathbf{x}_i - \mathbf{u}}{h} \right)^j K_h(\beta^T \mathbf{x}_i - \mathbf{u}) \quad \text{and}$$

$$\xi_{nj}(\mathbf{u}, \boldsymbol{\beta}) \stackrel{\text{def}}{=} n^{-1} \sum_{i=1}^n \mathbf{z}_i \otimes \left\{ \left(\frac{\boldsymbol{\beta}^T \mathbf{x}_i - \mathbf{u}}{h} \right)^j \mathbf{y}_i^T \right\} K_h(\boldsymbol{\beta}^T \mathbf{x}_i - \mathbf{u}).$$

Here $\mathbf{A} \otimes \mathbf{B} = (a_{ij} \mathbf{B})$ for $\mathbf{A} = (a_{ij})$, and $\mathbf{A}^0 = 1$, $\mathbf{A}^1 = \mathbf{A}$ and $\mathbf{A}^2 = \mathbf{A} \mathbf{A}^T$. For a fixed $\boldsymbol{\beta}$, $\mathbf{m}_k$ is now profiled out. Subsequently, we estimate $\boldsymbol{\beta}_{-d}$ through minimizing

$$\sum_{i=1}^n \left\{ \mathbf{y}_i - \sum_{k=1}^q \hat{\mathbf{m}}_k(\mathbf{x}_{d,i} + \boldsymbol{\beta}_{-d}^T \mathbf{x}_{-d,i}, \boldsymbol{\beta}) Z_{ik} \right\}^T \mathbf{W} \left\{ \mathbf{y}_i - \sum_{k=1}^q \hat{\mathbf{m}}_k(\mathbf{x}_{d,i} + \boldsymbol{\beta}_{-d}^T \mathbf{x}_{-d,i}, \boldsymbol{\beta}) Z_{ik} \right\}, \quad (2.2)$$

where $\mathbf{W}$ is a user-specified $r \times r$ positive-definite weight matrix. Let $\hat{\boldsymbol{\beta}}_{-d, \mathbf{W}}$ be the profile least squares estimate of $\hat{\boldsymbol{\beta}}_{-d, \mathbf{W}}$ if the working weight matrix $\mathbf{W}$ is used.

We need regularity conditions to establish the asymptotic normality property for $\hat{\boldsymbol{\beta}}_{-d, \mathbf{W}}$. For notational clarity, let $f(\boldsymbol{\beta}^T \mathbf{x})$ be the density function of $(\boldsymbol{\beta}^T \mathbf{x})$, $\mathbf{m}(\boldsymbol{\beta}^T \mathbf{x}) = (\mathbf{m}_1(\boldsymbol{\beta}^T \mathbf{x}), \dots, \mathbf{m}_r(\boldsymbol{\beta}^T \mathbf{x}))^T$, $\mathbf{m}_k^{(1)}(\boldsymbol{\beta}^T \mathbf{x}) = (\mathbf{m}_{k1}^{(1)}(\boldsymbol{\beta}^T \mathbf{x}), \dots, \mathbf{m}_{kr}^{(1)}(\boldsymbol{\beta}^T \mathbf{x}))^T$ be the first derivative of $\mathbf{m}_k(\boldsymbol{\beta}^T \mathbf{x})$ with respect to $(\boldsymbol{\beta}^T \mathbf{x})$ for $k = 1, \dots, q$.

(C1) (*The Lipschitz Continuity*) The density function $f(\boldsymbol{\beta}^T \mathbf{x})$ of $(\boldsymbol{\beta}^T \mathbf{x})$ is locally Lipschitz continuous, and bounded away from zero and infinity. In addition, $\mathbf{m}(\boldsymbol{\beta}^T \mathbf{x})$, $E(\mathbf{x}|\boldsymbol{\beta}^T \mathbf{x})$ and $\boldsymbol{\Omega}(\boldsymbol{\beta}^T \mathbf{x}) = E(\mathbf{z}\mathbf{z}^T|\boldsymbol{\beta}^T \mathbf{x})$ are locally Lipschitz continuous.

(C2) (*The Kernel Function*) The univariate kernel function $K(\cdot)$ is symmetric, has a compact support and is Lipschitz continuous. In addition, $\int K(u)du = 1$, $\int u^k K(u)du = 0$, for $k = 1, \dots, s-1$, and $0 \neq \int u^s K(u)du < \infty$. The d -dimensional kernel is a product of d univariate kernels. We abuse the notation of K here when it is sufficiently clear from the context.

(C3) (*The Bandwidth*) The bandwidth $h = O(n^{-\delta})$ for $(4s)^{-1} < \delta < (2d)^{-1}$.

(C4) (*The Moment Condition*) All the involved moments, $E[\{\mathbf{m}_k(\boldsymbol{\beta}^T \mathbf{x})\}^T \{\mathbf{m}_k(\boldsymbol{\beta}^T \mathbf{x})\}]$, $E(\mathbf{x}^T \mathbf{x})$, $E\{\mathbf{y}^T \mathbf{y}\}^{\kappa_1}$ and $E\left[\left\{\mathbf{m}_k^{(1)}(\boldsymbol{\beta}^T \mathbf{x})\right\}^T \left\{\mathbf{m}_k^{(1)}(\boldsymbol{\beta}^T \mathbf{x})\right\}\right]$, exist for so-me $\kappa_1 \geq 3/2$ and $k = 1, \dots, q$.

These conditions are generally regarded as mild. In particular, condition (C1) imposes smoothness conditions on the mean and density functions that allow us to implement such local smoothers as kernel and local polynomial regressions (Fan and Gijbels (1996)). Condition (C2) states that an s -th order kernel function is used. Condition (C3) specifies the order of the bandwidth,

whose range is fairly wide and, more importantly, contains an optimal order. We assume moment conditions in condition (C4) to establish the asymptotic normality. Similar conditions appear in Ma and Zhu (2012, 2013).

Define

$$\mathbf{A}_{\mathbf{w}} \stackrel{\text{def}}{=} E \left[\left\{ \sum_{k=1}^q \mathbf{m}_k^{(1),\top} (\boldsymbol{\beta}^\top \mathbf{x}) Z_k \otimes \tilde{\mathbf{x}}_{-d} \right\} \mathbf{W} \left\{ \sum_{k=1}^q \mathbf{m}_k^{(1)} (\boldsymbol{\beta}^\top \mathbf{x}) Z_k \otimes \tilde{\mathbf{x}}_{-d}^\top \right\} \right], \text{ and}$$

$$\mathbf{B}_{\mathbf{w}} \stackrel{\text{def}}{=} E \left[\left\{ \sum_{k=1}^q \mathbf{m}_k^{(1),\top} (\boldsymbol{\beta}^\top \mathbf{x}) Z_k \otimes \tilde{\mathbf{x}}_{-d} \right\} \mathbf{W} \boldsymbol{\Sigma} \mathbf{W} \left\{ \sum_{k=1}^q \mathbf{m}_k^{(1)} (\boldsymbol{\beta}^\top \mathbf{x}) Z_k \otimes \tilde{\mathbf{x}}_{-d}^\top \right\} \right].$$

Theorem 1. *If conditions (C1)-(C4) in the Appendix hold, then*

$$n^{1/2} \left\{ \text{vec}(\hat{\boldsymbol{\beta}}_{-d,\mathbf{w}}) - \text{vec}(\boldsymbol{\beta}_{-d}) \right\} \xrightarrow{d} \mathcal{N}(\mathbf{0}, \mathbf{A}_{\mathbf{w}}^{-1} \mathbf{B}_{\mathbf{w}} \mathbf{A}_{\mathbf{w}}^{-1}),$$

where $\tilde{\mathbf{x}}_{-d} = \mathbf{x}_{-d} - E(\mathbf{x}_{-d} | \boldsymbol{\beta}^\top \mathbf{x})$ and “ $\xrightarrow{d}$ ” stands for “convergence in distribution”.

How to specify the working weight matrix $\mathbf{W}$ is an interesting issue. As long as $\mathbf{W}$ is positive definite, $\hat{\boldsymbol{\beta}}_{-d,\mathbf{w}}$ is root- n consistent and asymptotically normal. However, choosing an appropriate $\mathbf{W}$ may improve the efficiency of estimating $\boldsymbol{\beta}_{-d,\mathbf{w}}$. We compare two options: $\mathbf{W} = \mathbf{I}_{r \times r}$ and $\mathbf{W} = \hat{\boldsymbol{\Sigma}}^{-1}$, where

$$\hat{\boldsymbol{\Sigma}} \stackrel{\text{def}}{=} n^{-1} \sum_{i=1}^n \hat{\boldsymbol{\varepsilon}}_i \hat{\boldsymbol{\varepsilon}}_i^\top \text{ and } \hat{\boldsymbol{\varepsilon}}_i \stackrel{\text{def}}{=} \mathbf{y}_i - \sum_{k=1}^q \hat{\mathbf{m}}_k \left(\hat{\boldsymbol{\beta}}_1^\top \mathbf{x}_i, \hat{\boldsymbol{\beta}}_1 \right) Z_{ik}.$$

Theorem 2 indicates that using $\mathbf{W} = \hat{\boldsymbol{\Sigma}}^{-1}$ yields a more efficient estimate of $\boldsymbol{\beta}_{-d,\mathbf{w}}$ than using $\mathbf{W} = \mathbf{I}_{r \times r}$.

Theorem 2. $\mathbf{A}_{\mathbf{I}}^{-1} \mathbf{B}_{\mathbf{I}} \mathbf{A}_{\mathbf{I}}^{-1} \geq \mathbf{A}_{\boldsymbol{\Sigma}^{-1}}^{-1} \mathbf{B}_{\boldsymbol{\Sigma}^{-1}} \mathbf{A}_{\boldsymbol{\Sigma}^{-1}}^{-1} = \mathbf{A}_{\boldsymbol{\Sigma}^{-1}}^{-1}$.

For the asymptotic normality of $\hat{\boldsymbol{\beta}}_{-d,\mathbf{w}}$ to be useful, we provide a consistent estimate for the asymptotic covariance matrix. Let $\hat{\boldsymbol{\beta}}_{\mathbf{w}} \stackrel{\text{def}}{=} (\mathbf{I}_{d \times d}, \hat{\boldsymbol{\beta}}_{-d,\mathbf{w}}^\top)^\top$, where $\mathbf{W}$ can be $\hat{\boldsymbol{\Sigma}}^{-1}$ or $\mathbf{I}$. Take

$$\hat{\tilde{\mathbf{x}}}_{-d,i} \stackrel{\text{def}}{=} \mathbf{x}_{-d,i} - \frac{\sum_{j=1, j \neq i}^n K_h(\hat{\boldsymbol{\beta}}_{\mathbf{w}}^\top \mathbf{x}_j - \hat{\boldsymbol{\beta}}_{\mathbf{w}}^\top \mathbf{x}_i) \mathbf{x}_{-d,i}}{\sum_{j=1, j \neq i}^n K_h(\hat{\boldsymbol{\beta}}_{\mathbf{w}}^\top \mathbf{x}_j - \hat{\boldsymbol{\beta}}_{\mathbf{w}}^\top \mathbf{x}_i)},$$

$$\hat{\mathbf{A}}_{\mathbf{w}} \stackrel{\text{def}}{=} n^{-1} \sum_{i=1}^n \left\{ \sum_{k=1}^q \hat{\mathbf{m}}_k^{(1),\top} \left(\hat{\boldsymbol{\beta}}_{\mathbf{w}}^\top \mathbf{x}_i, \hat{\boldsymbol{\beta}}_{\mathbf{w}} \right) Z_{ik} \otimes \hat{\tilde{\mathbf{x}}}_{-d,i} \right\}$$

$$\mathbf{W} \left\{ \sum_{k=1}^q \hat{\mathbf{m}}_k^{(1)} \left(\hat{\boldsymbol{\beta}}_{\mathbf{w}}^\top \mathbf{x}_i, \hat{\boldsymbol{\beta}}_{\mathbf{w}} \right) Z_{ik} \otimes \hat{\tilde{\mathbf{x}}}_{-d,i}^\top \right\},$$

$$\begin{aligned}\widehat{\mathbf{B}}_{\mathbf{w}} &\stackrel{\text{def}}{=} n^{-1} \sum_{i=1}^n \left\{ \sum_{k=1}^q \widehat{\mathbf{m}}_k^{(1),\top} \left(\widehat{\boldsymbol{\beta}}_{\mathbf{w}}^{\top} \mathbf{x}_i, \widehat{\boldsymbol{\beta}}_{\mathbf{w}} \right) Z_{ik} \otimes \widehat{\mathbf{x}}_{-d,i} \right\} \\ &\quad \mathbf{W} \widehat{\boldsymbol{\Sigma}} \mathbf{W} \left\{ \sum_{k=1}^q \widehat{\mathbf{m}}_k^{(1)} \left(\widehat{\boldsymbol{\beta}}_{\mathbf{w}}^{\top} \mathbf{x}_i, \widehat{\boldsymbol{\beta}}_{\mathbf{w}} \right) Z_{ik} \otimes \widehat{\mathbf{x}}_{-d,i}^{\top} \right\}.\end{aligned}$$

Theorem 3. *If conditions (C1)-(C4) hold, then $\widehat{\mathbf{A}}_{\mathbf{w}} \xrightarrow{p} \mathbf{A}_{\mathbf{w}}$, $\widehat{\mathbf{B}}_{\mathbf{w}} \xrightarrow{p} \mathbf{B}_{\mathbf{w}}$, and hence $\widehat{\mathbf{A}}_{\mathbf{w}}^{-1} \widehat{\mathbf{B}}_{\mathbf{w}} \widehat{\mathbf{A}}_{\mathbf{w}}^{-1} \xrightarrow{p} \mathbf{A}_{\mathbf{w}}^{-1} \mathbf{B}_{\mathbf{w}} \mathbf{A}_{\mathbf{w}}^{-1}$.*

Testing whether there exist interaction effects between X_i and $\mathbf{z}$ amounts to testing whether all components of the i -th row of $\boldsymbol{\beta}$ in model (1.1) are simultaneously zero. In a general context, we consider the hypothesis testing problem

$$H_0 : \mathbf{Q}\boldsymbol{\beta}_{-d} = \mathbf{q}_0 \quad \text{versus} \quad H_1 : \mathbf{Q}\boldsymbol{\beta}_{-d} \neq \mathbf{q}_0,$$

where $\mathbf{Q}$ is a user-specified $q_0 \times (p-d)$ matrix and $\mathbf{q}_0$ is another user-specified $q_0 \times d$ matrix. This problem is general enough to include a variety of hypothesis of interest. For example, we are free to choose $\mathbf{Q} = (1, 0, \dots, 0)_{1 \times (p-d)}$ and $\mathbf{q}_0 = \mathbf{0}_{1 \times d}$, aiming to test whether there exist interaction effects between X_{d_0+1} and $\mathbf{z}$. In general, we devise the Wald chi-square test

$$\begin{aligned}T_{\mathbf{w}} &= n \left\{ (\mathbf{I}_{d \times d} \otimes \mathbf{Q}) \text{vec}(\widehat{\boldsymbol{\beta}}_{-d,\mathbf{w}}) - \text{vec}(\mathbf{q}_0) \right\}^{\top} \\ &\quad \left\{ (\mathbf{I}_{d \times d} \otimes \mathbf{Q}) \widehat{\mathbf{A}}_{\mathbf{w}}^{-1} \widehat{\mathbf{B}}_{\mathbf{w}} \widehat{\mathbf{A}}_{\mathbf{w}}^{-1} (\mathbf{I}_{d \times d} \otimes \mathbf{Q}^{\top}) \right\}^{-1} \left\{ (\mathbf{I}_{d \times d} \otimes \mathbf{Q}) \text{vec}(\widehat{\boldsymbol{\beta}}_{-d,\mathbf{w}}) - \text{vec}(\mathbf{q}_0) \right\}.\end{aligned}$$

A direct application of Theorem 1 yields the following corollary. Its proof is omitted.

Corollary 1. *If conditions (C1)-(C4) hold, then under H_0 , $T_{\mathbf{w}} \xrightarrow{d} \chi^2(q_0 d)$, where $\chi^2(q_0 d)$ is the central chi-square distribution with $(q_0 d)$ degrees of freedom.*

It remains to estimate the structural dimension of $\text{span}(\boldsymbol{\beta})$, the minimum column dimension of $\boldsymbol{\beta}$, such that (1.1) holds. Following Zhu, Miao and Peng (2006) and Xu et al. (2016), we suggest a BIC-type criterion. Specifically, for a working dimension d , we take

$$\begin{aligned}\mathcal{L}(d) &\stackrel{\text{def}}{=} \frac{\sum_{i=1}^n \{ \mathbf{y}_i - \sum_{k=1}^q \widehat{\mathbf{m}}_k(\widehat{\boldsymbol{\beta}}_{d,\mathbf{w}}^{\top} \mathbf{x}_i, \widehat{\boldsymbol{\beta}}_{d,\mathbf{w}}) Z_{ik} \}^{\top} \{ \mathbf{y}_i - \sum_{k=1}^q \widehat{\mathbf{m}}_k(\widehat{\boldsymbol{\beta}}_{d,\mathbf{w}}^{\top} \mathbf{x}_i, \widehat{\boldsymbol{\beta}}_{d,\mathbf{w}}) Z_{ik} \}}{\{ \sum_{i=1}^n (\mathbf{y}_i - \bar{\mathbf{y}})^{\top} (\mathbf{y}_i - \bar{\mathbf{y}}) \}^{1/2}} \\ \text{and } \mathcal{L}^*(d) &\stackrel{\text{def}}{=} \mathcal{L}(d) + (pd)\lambda_n,\end{aligned}$$

where $\widehat{\boldsymbol{\beta}}_{d,\mathbf{w}} = (\mathbf{I}_{d \times d}, \widehat{\boldsymbol{\beta}}_{-d,\mathbf{w}})^{\top}$. The estimated structural dimension is then given by

$$\hat{d} \stackrel{\text{def}}{=} \underset{1 \leq d \leq p}{\operatorname{argmin}} \mathcal{L}^*(d). \quad (2.3)$$

Theorem 4. *Under the conditions of Theorem 1, if $\lambda_n / \log n \rightarrow \infty$ and $\lambda_n n^{-1/2} \rightarrow 0$, then $\operatorname{pr}(\hat{d} = d_0) \rightarrow 1$.*

Thus the BIC-type criterion enables us to select the true structural dimensional of $\operatorname{span}(\beta)$ consistently. The penalty term λ_n is allowed to vary over a wide enough range for $\hat{d}$ to be consistent. How to choose an optimal λ_n is challenging. Our limited simulations show that $\lambda_n = n^{2/5}$ works well. We use this choice of λ_n throughout our numerical studies.

An algorithm for estimating β is as follows, starting with a working dimension d and a user-specified initial value of β .

1. Estimate $\mathbf{m}_k$ and $\mathbf{m}_k^{(1)}$ with (2.1) for a given β .
2. Set $\mathbf{W} = \mathbf{I}_{r \times r}$. Estimate β with (2.2) for given $\mathbf{m}_k$ and $\mathbf{m}_k^{(1)}$.
3. Repeat these two steps until convergence. The resultant estimate, denoted by $\hat{\beta}_{\mathbf{I}} = (\mathbf{I}_{d \times d}, \hat{\beta}_{-d, \mathbf{I}}^T)^T$, is referred to as the unweighted-profile least squares estimate.
4. We vary the working dimension d from 1 through p and repeat the above three steps. The estimated dimension $\hat{d}$ is given in (2.3).
5. Set $\mathbf{W} = \hat{\Sigma}^{-1}$ and $d = \hat{d}$ in the second step. Repeat the first two steps until convergence. The final estimate, denoted by $\hat{\beta}_{\hat{\Sigma}^{-1}} = (\mathbf{I}_{\hat{d} \times \hat{d}}, \hat{\beta}_{-\hat{d}, \hat{\Sigma}^{-1}}^T)$, is referred to as the weighted-profile least squares estimate.

3. Numerical Studies

In this section we demonstrate the performance of our proposals through comprehensive simulations and an application to the Framingham Heart Study. Because existing methods cannot be used directly if $\mathbf{y}$ is multivariate, we only report the simulation results of our proposal in Section 3.1 when $\mathbf{y}$ is multivariate. In Section 3.2, we compare our proposal with methods proposed by Li, Cook and Chiaromonte (2003), Ma and Song (2015) and Liu, Cui and Li (2016) when both $\mathbf{y}$ and $\mathbf{z}$ are univariate.

3.1. Simulation experiments for multivariate response data

We conducted simulation studies to evaluate the performance of our proposed methodology when the response is multivariate. Throughout our simulations we drew $\mathbf{x}$ and $\mathbf{z}$ independently from multivariate normal distribution with zero mean and covariance matrix $(0.5^{|k-l|})$. We fixed $r = 3$, and generated $\boldsymbol{\varepsilon} = (\varepsilon_1, \varepsilon_2, \varepsilon_3)^\top$ from $\mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma})$, where

$$\boldsymbol{\Sigma} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 4 \end{pmatrix} \begin{pmatrix} 1 & \rho & \rho^2 \\ \rho & 1 & \rho \\ \rho^2 & \rho & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 4 \end{pmatrix}.$$

We considered four simulated models.

Model I: A single-index model structure with a linear link function:

$$\begin{cases} Y_1 = 2(\boldsymbol{\beta}^\top \mathbf{x})Z_1 + (\boldsymbol{\beta}^\top \mathbf{x})Z_2 + \varepsilon_1, \\ Y_2 = (\boldsymbol{\beta}^\top \mathbf{x})Z_1 + 3(\boldsymbol{\beta}^\top \mathbf{x})Z_2 + \varepsilon_2, \\ Y_3 = \varepsilon_3, \end{cases}$$

Model II: A single-index model structure with a nonlinear link function:

$$\begin{cases} Y_1 = \sin(4\boldsymbol{\beta}^\top \mathbf{x})Z_1 + 2(\boldsymbol{\beta}^\top \mathbf{x})Z_2 + \varepsilon_1, \\ Y_2 = \cos(2\boldsymbol{\beta}^\top \mathbf{x})Z_2 + \varepsilon_2, \\ Y_3 = 2(\boldsymbol{\beta}^\top \mathbf{x})Z_1 + \sin(2\boldsymbol{\beta}^\top \mathbf{x})Z_2 + \varepsilon_3. \end{cases}$$

Model III: A multiple-index model structure with a linear link function:

$$\begin{cases} Y_1 = \{(\boldsymbol{\beta}_1^\top \mathbf{x}) + (\boldsymbol{\beta}_2^\top \mathbf{x})\}Z_1 + (\boldsymbol{\beta}_1^\top \mathbf{x})Z_2 + \varepsilon_1, \\ Y_2 = (\boldsymbol{\beta}_2^\top \mathbf{x})Z_1 + \{(2\boldsymbol{\beta}_1^\top \mathbf{x}) - 3(\boldsymbol{\beta}_2^\top \mathbf{x})\}Z_2 + \varepsilon_2, \\ Y_3 = 2(\boldsymbol{\beta}_1^\top \mathbf{x})Z_1 + 4(\boldsymbol{\beta}_2^\top \mathbf{x})Z_2 + \varepsilon_3. \end{cases}$$

Model IV: A multiple-index model structure with a nonlinear link function:

$$\begin{cases} Y_1 = \frac{(\boldsymbol{\beta}_1^\top \mathbf{x})Z_1}{\{0.5 + (\boldsymbol{\beta}_2^\top \mathbf{x} + 1.5)^2\}} + (\boldsymbol{\beta}_2^\top \mathbf{x})Z_2 + \varepsilon_1, \\ Y_2 = \sin^2(\boldsymbol{\beta}_1^\top \mathbf{x})Z_1 + \cos^2(\boldsymbol{\beta}_2^\top \mathbf{x})Z_2 + \varepsilon_2, \\ Y_3 = \{(2\boldsymbol{\beta}_1^\top \mathbf{x}) - (\boldsymbol{\beta}_2^\top \mathbf{x})\}^2 Z_1 + \varepsilon_3. \end{cases}$$

We set $p = 10$, $q = 2$ and $\boldsymbol{\beta} = (1, 0.8, 0.6, 0.4, 0.2, -0.2, -0.4, -0.6, -0.8, 0)^\top$ in Models (I)-(II), and set $p = 7$, $q = 2$, $\boldsymbol{\beta}_1 = (1, 0, 0.8, -0.6, 0.4, -0.2, 0)^\top$ and $\boldsymbol{\beta}_2 = (0, 1, -0.8, 0.6, -0.4, 0.2, 0)^\top$ in Models (III)-(IV). We chose the sample size $n = 200$ and 500 and repeated each simulation 1,000 times. We used the Gaussian kernel and chose the bandwidth $h = (4/3n)^{1/(d+4)}s$, where s is the median of the robust estimators of the standard deviation of $(\boldsymbol{\beta}^\top \mathbf{x})$.

The average of estimation bias (“bias”), the Monte Carlo standard deviation (“std”), the average of estimated standard deviation (“ $\widehat{\text{std}}$ ”), and the empirical coverage probability (“cvp”) at the nominal 95% confidence level for all free parameter are summarized in Tables 1–4 for models (I)–(IV), respectively. These estimates have very small biases, and the biases decrease as the sample size increases. This phenomenon provides strong evidence that both the weighted and the unweighted estimates are consistent, the theoretical result of Theorem 1.

In terms of the Monte Carlo standard deviation and the average of estimated standard deviation, the weighted estimate performs competitively in comparison with the unweighted one. The empirical coverage probabilities for the weighted and unweighted estimators are close to the nominal level, which implies that our inferential results are fairly reliable. The Monte Carlo standard deviations are close to the average of the estimated standard deviations especially for large n . This finding means that the standard deviations have been estimated precisely, which verifies Theorem 3.

To demonstrate the performance of our proposed Wald test statistic $T_{\mathbf{w}}$, we tested whether X_7 interacts with $\mathbf{z}$ in Model (IV). Towards this end we simply chose $\mathbf{Q} = (0, \dots, 0, 1)_{1 \times 5}$, $\mathbf{q}_0 = \mathbf{0}_{1 \times 2}$ in our testing problem. To investigate the size and the power performance of our proposed Wald test, we changed the last row of $\boldsymbol{\beta} = (\boldsymbol{\beta}_1, \boldsymbol{\beta}_2)$ to (a, a) and reestimated all parameters, with $a = -0.10 : 0.02 : 0.10$. Apparently, $a = 0.00$ corresponds to the case that X_7 does not interact with $\mathbf{z}$. The power curves at the significance level 0.05 are reported in Figure 1 based on 1,000 replications. All the power curves increase quickly as $|a|$ increases, indicating that our proposed test approach can detect the interaction effects effectively.

We evaluated the performance of the BIC-type criterion at (2.3) in estimating the structural dimension of $\text{span}(\boldsymbol{\beta})$. The structure dimension was $d_0 = 1$ in models (I)–(II) and $d_0 = 2$ in models (III) and (IV). For Models (I)–(IV), the percentages for each estimated dimension are charted in Table 5. Our proposed BIC-type criterion works pretty well: with high probability the estimated and the true structural dimension are equal in all models. The performance of our BIC-type criterion also improves gradually as the sample size increases.

3.2. Comparison with existing methods for univariate response data

In this section we compare our proposal (NEW for short) with existing methods of Li, Cook and Chiaromonte (2003) (LCC for short), Ma and Song (2015) (MS for short) and Liu, Cui and Li (2016) (LCL for short) when both $\mathbf{y}$ and $\mathbf{z}$

Table 2. Simulation results for Model II: the average bias of the estimators (“bias”), the Monte Carlo standard deviation (“std”), the average of the estimated standard deviation (“ $\widehat{\text{std}}$ ”) based on the theoretical calculation, and the empirical coverage probability (“cvp”) at the nominal 95% confidence level. All simulation results reported below are multiplied by 100.

True value		$\widehat{\beta}_2$ 0.8	$\widehat{\beta}_3$ 0.6	$\widehat{\beta}_4$ 0.4	$\widehat{\beta}_5$ 0.2	$\widehat{\beta}_6$ -0.2	$\widehat{\beta}_7$ -0.4	$\widehat{\beta}_8$ -0.6	$\widehat{\beta}_9$ -0.8	$\widehat{\beta}_{10}$ 0
W $\rho = 0.5, n = 200$										
I	bias	2.18	0.87	0.49	0.32	-0.31	-0.84	-1.14	-0.81	-0.12
	std	11.61	9.69	8.81	8.66	8.85	9.10	9.90	10.30	7.84
	$\widehat{\text{std}}$	12.62	9.89	9.34	8.96	8.94	9.34	9.94	10.66	7.90
	cvp	95.90	95.20	95.90	96.10	94.60	95.00	94.60	95.50	95.50
$\widehat{\Sigma}^{-1}$	bias	1.59	0.50	0.06	0.19	-0.15	-0.44	-0.83	-0.65	0.01
	std	6.79	5.67	5.26	4.98	5.21	5.33	5.65	5.97	4.52
	$\widehat{\text{std}}$	7.10	5.56	5.24	5.04	5.04	5.26	5.59	6.00	4.44
	cvp	95.30	95.30	94.70	95.90	94.30	93.80	95.50	95.30	95.60
W $\rho = 0.5, n = 500$										
I	bias	1.58	0.47	0.50	0.11	-0.08	-0.69	-0.24	-0.89	-0.21
	std	7.54	6.23	5.71	5.57	5.52	5.84	6.18	6.52	4.96
	$\widehat{\text{std}}$	8.00	6.24	5.87	5.62	5.63	5.87	6.24	6.74	4.96
	cvp	96.40	95.50	95.90	95.30	94.70	94.30	95.20	95.90	94.70
$\widehat{\Sigma}^{-1}$	bias	1.40	0.57	0.34	0.30	-0.20	-0.48	-0.58	-0.70	-0.05
	std	3.77	3.26	3.08	3.02	3.03	3.09	3.27	3.43	2.68
	$\widehat{\text{std}}$	4.29	3.34	3.14	3.01	3.02	3.14	3.35	3.61	2.66
	cvp	96.50	95.20	94.80	93.70	94.40	94.40	94.50	95.90	94.90
W $\rho = 0.8, n = 200$										
I	bias	2.19	0.80	0.44	0.15	-0.40	-0.86	-0.91	-0.76	-0.25
	std	12.47	10.18	9.42	9.17	9.31	9.57	10.46	10.48	8.08
	$\widehat{\text{std}}$	13.25	10.38	9.80	9.39	9.37	9.80	10.43	11.19	8.29
	cvp	96.10	95.00	96.20	95.80	95.30	95.90	95.10	96.20	95.40
$\widehat{\Sigma}^{-1}$	bias	1.53	0.45	0.21	0.22	-0.21	-0.37	-0.84	-0.61	-0.07
	std	5.69	4.80	4.34	4.20	4.22	4.44	4.67	5.01	3.77
	$\widehat{\text{std}}$	5.88	4.61	4.35	4.18	4.18	4.35	4.64	4.97	3.68
	cvp	95.00	95.90	95.00	95.20	94.90	93.60	94.20	95.60	94.10
W $\rho = 0.8, n = 500$										
I	bias	1.62	0.49	0.43	0.04	-0.11	-0.68	-0.24	-0.74	-0.27
	std	7.82	6.49	6.00	6.04	5.91	6.13	6.41	6.63	5.06
	$\widehat{\text{std}}$	8.38	6.53	6.13	5.88	5.89	6.14	6.53	7.05	5.18
	cvp	96.00	94.90	96.00	94.40	94.80	95.10	95.10	96.30	95.40
$\widehat{\Sigma}^{-1}$	bias	1.36	0.53	0.35	0.25	-0.17	-0.43	-0.57	-0.68	-0.03
	std	3.01	2.56	2.40	2.38	2.39	2.42	2.56	2.70	2.18
	$\widehat{\text{std}}$	3.48	2.71	2.54	2.44	2.44	2.55	2.71	2.93	2.15
	cvp	95.70	94.90	95.70	94.40	94.50	95.10	93.90	95.80	93.80

Table 3. Simulation results for Model III: the average bias of the estimators (“bias”), the Monte Carlo standard deviation (“std”), the average of the estimated standard deviation (“ $\widehat{\text{std}}$ ”) based on the theoretical calculation, and the empirical coverage probability (“cvp”) at the nominal 95% confidence level. All simulation results reported below are multiplied by 100.

True value		$\widehat{\beta}_{13}$	$\widehat{\beta}_{14}$	$\widehat{\beta}_{15}$	$\widehat{\beta}_{16}$	$\widehat{\beta}_{17}$	$\widehat{\beta}_{23}$	$\widehat{\beta}_{24}$	$\widehat{\beta}_{25}$	$\widehat{\beta}_{26}$	$\widehat{\beta}_{27}$
		0.8	-0.6	0.4	-0.2	0	-0.8	0.6	-0.4	0.2	0
W		$\rho = 0.5, n = 200$									
I	bias	2.11	-1.33	0.80	-0.33	-0.12	-0.85	0.93	-0.69	0.46	-0.19
	std	5.96	5.86	5.30	4.95	4.14	5.07	4.92	4.41	3.93	3.34
	$\widehat{\text{std}}$	7.38	7.25	6.28	5.58	4.79	5.99	5.92	5.11	4.54	3.91
	cvp	98.40	98.20	98.20	97.40	97.60	97.50	97.50	97.00	97.20	97.10
$\widehat{\Sigma}^{-1}$	bias	1.21	-0.92	0.56	-0.22	-0.09	-0.86	0.81	-0.52	0.30	-0.12
	std	3.91	3.77	3.28	3.03	2.50	3.22	3.23	2.81	2.46	1.99
	$\widehat{\text{std}}$	4.31	4.24	3.68	3.26	2.80	3.52	3.48	3.01	2.67	2.30
	cvp	97.30	97.10	96.70	96.40	96.60	95.70	95.90	96.00	96.50	97.90
W		$\rho = 0.5, n = 500$									
I	bias	1.62	-1.23	0.90	-0.37	-0.04	-0.73	0.89	-0.43	0.25	0.00
	std	4.27	4.27	3.68	3.22	2.93	3.59	3.47	3.09	2.76	2.40
	$\widehat{\text{std}}$	5.41	5.31	4.56	4.04	3.46	4.45	4.38	3.76	3.33	2.86
	cvp	98.80	98.00	97.90	98.60	97.60	98.40	98.50	97.80	98.30	97.80
$\widehat{\Sigma}^{-1}$	bias	0.88	-0.69	0.45	-0.24	0.02	-0.65	0.71	-0.47	0.23	0.02
	std	2.72	2.77	2.43	2.05	1.91	2.32	2.17	1.88	1.71	1.39
	$\widehat{\text{std}}$	3.11	3.05	2.62	2.33	1.99	2.57	2.52	2.17	1.92	1.65
	cvp	96.90	96.90	96.40	96.90	96.10	96.80	97.00	97.50	97.90	97.90
W		$\rho = 0.8, n = 200$									
I	bias	2.13	-1.42	0.82	-0.38	-0.01	-0.73	0.93	-0.79	0.44	-0.23
	std	6.51	6.56	5.80	5.23	4.52	5.78	5.61	4.93	4.27	3.78
	$\widehat{\text{std}}$	8.07	7.93	6.89	6.11	5.24	6.56	6.46	5.57	4.96	4.27
	cvp	98.30	98.00	98.60	98.30	98.20	96.80	97.30	96.60	96.50	96.30
$\widehat{\Sigma}^{-1}$	bias	0.89	-0.64	0.39	-0.13	-0.06	-0.52	0.46	-0.30	0.17	-0.08
	std	3.27	3.21	2.74	2.57	2.11	2.41	2.42	2.11	1.85	1.51
	$\widehat{\text{std}}$	3.60	3.55	3.07	2.73	2.34	2.61	2.59	2.23	1.98	1.71
	cvp	97.60	96.50	97.10	95.80	96.80	95.90	95.90	96.00	96.90	97.60
W		$\rho = 0.8, n = 500$									
I	bias	1.66	-1.25	0.88	-0.32	0.04	-0.79	0.91	-0.42	0.31	-0.07
	std	4.56	4.65	4.07	3.62	3.22	3.89	3.88	3.40	3.01	2.62
	$\widehat{\text{std}}$	5.95	5.83	5.02	4.44	3.80	4.87	4.79	4.12	3.66	3.13
	cvp	98.70	97.90	98.30	97.60	97.20	99.20	98.00	98.10	98.60	97.90
$\widehat{\Sigma}^{-1}$	bias	0.57	-0.39	0.26	-0.17	0.02	-0.35	0.39	-0.24	0.12	0.00
	std	2.33	2.31	2.01	1.73	1.57	1.69	1.57	1.39	1.27	1.03
	$\widehat{\text{std}}$	2.55	2.51	2.15	1.91	1.63	1.85	1.82	1.56	1.39	1.19
	cvp	97.10	96.60	96.30	96.10	95.70	97.10	98.30	96.90	97.10	97.60

Table 4. Simulation results for Model IV: the average bias of the estimators (“bias”), the Monte Carlo standard deviation (“std”), the average of the estimated standard deviation (“ $\widehat{\text{std}}$ ”) based on the theoretical calculation, and the empirical coverage probability (“cvp”) at the nominal 95% confidence level. All simulation results reported below are multiplied by 100.

True value		$\widehat{\beta}_{13}$	$\widehat{\beta}_{14}$	$\widehat{\beta}_{15}$	$\widehat{\beta}_{16}$	$\widehat{\beta}_{17}$	$\widehat{\beta}_{23}$	$\widehat{\beta}_{24}$	$\widehat{\beta}_{25}$	$\widehat{\beta}_{26}$	$\widehat{\beta}_{27}$
		0.8	-0.6	0.4	-0.2	0	-0.8	0.6	-0.4	0.2	0
W $\rho = 0.5, n = 200$											
I	bias	-0.33	0.47	-0.29	0.22	0.00	-0.21	0.54	-0.30	0.15	0.06
	std	4.32	4.12	3.64	3.29	2.89	5.73	5.63	4.86	4.34	3.68
	$\widehat{\text{std}}$	5.10	5.04	4.34	3.90	3.34	6.54	6.46	5.58	4.98	4.27
	cvp	97.80	98.00	97.30	97.70	98.00	97.20	97.70	96.90	96.40	96.30
$\widehat{\Sigma}^{-1}$	bias	0.02	0.11	-0.02	0.11	-0.06	-0.24	0.43	-0.22	0.08	0.07
	std	3.06	2.98	2.64	2.28	2.04	3.33	3.31	2.86	2.63	2.18
	$\widehat{\text{std}}$	3.66	3.61	3.12	2.79	2.39	3.89	3.84	3.31	2.94	2.53
	cvp	97.90	97.40	97.30	97.90	97.90	97.30	97.20	97.00	97.50	98.00
W $\rho = 0.5, n = 500$											
I	bias	-0.52	0.62	-0.28	0.14	-0.08	-0.31	0.57	-0.16	0.12	-0.02
	std	2.91	2.90	2.55	2.16	1.96	4.25	4.15	3.53	3.03	2.67
	$\widehat{\text{std}}$	3.46	3.41	2.95	2.63	2.24	4.88	4.82	4.15	3.70	3.16
	cvp	97.50	97.20	97.60	98.30	96.50	97.00	97.60	98.10	98.10	97.60
$\widehat{\Sigma}^{-1}$	bias	-0.28	0.37	-0.18	0.09	-0.05	-0.31	0.49	-0.31	0.11	0.05
	std	2.06	1.99	1.69	1.51	1.36	2.44	2.30	1.98	1.83	1.60
	$\widehat{\text{std}}$	2.44	2.40	2.07	1.84	1.57	2.79	2.74	2.36	2.10	1.79
	cvp	97.80	97.80	98.60	98.30	97.00	97.40	97.30	98.80	97.50	97.60
W $\rho = 0.8, n = 200$											
I	bias	-0.33	0.45	-0.23	0.19	-0.05	-0.20	0.62	-0.22	0.20	-0.03
	std	4.54	4.43	3.92	3.58	3.02	6.51	6.05	5.48	4.85	4.27
	$\widehat{\text{std}}$	5.28	5.22	4.54	4.09	3.50	7.11	7.02	6.09	5.47	4.71
	cvp	97.40	97.60	97.60	97.70	97.40	96.50	97.50	96.00	96.90	96.70
$\widehat{\Sigma}^{-1}$	bias	-0.31	0.31	-0.18	0.15	-0.04	0.05	0.10	-0.02	-0.02	0.08
	std	2.40	2.37	2.11	1.90	1.59	2.91	2.89	2.55	2.21	1.81
	$\widehat{\text{std}}$	2.84	2.81	2.45	2.20	1.88	3.22	3.18	2.77	2.48	2.13
	cvp	97.80	97.60	96.80	98.30	98.00	97.20	96.60	97.00	97.80	97.90
W $\rho = 0.8, n = 500$											
I	bias	-0.40	0.50	-0.30	0.11	0.02	-0.40	0.50	-0.30	0.11	0.02
	std	2.98	2.99	2.70	2.34	2.03	4.44	4.50	3.88	3.41	2.79
	$\widehat{\text{std}}$	3.64	3.58	3.10	2.76	2.36	5.36	5.28	4.55	4.07	3.48
	cvp	98.20	98.00	97.60	98.20	97.20	97.80	98.10	98.20	98.40	98.10
$\widehat{\Sigma}^{-1}$	bias	-0.35	0.34	-0.19	0.13	-0.03	-0.35	0.34	-0.19	0.13	-0.03
	std	1.60	1.53	1.32	1.17	1.03	2.00	1.87	1.61	1.53	1.31
	$\widehat{\text{std}}$	1.86	1.83	1.58	1.41	1.20	2.27	2.23	1.92	1.71	1.46
	cvp	97.10	97.20	97.80	98.00	97.00	97.30	97.40	98.20	96.90	97.70

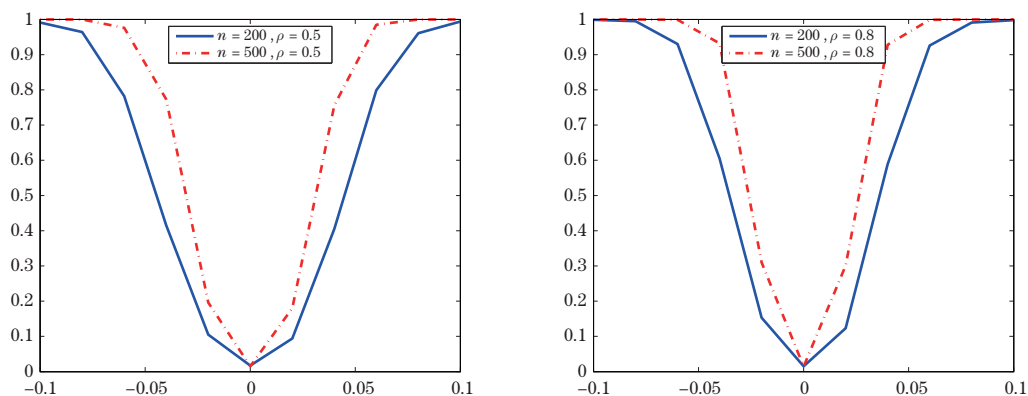


Figure 1. The power curves of $T_{\mathbf{w}}$ for $\rho = 0.5$ (left panel) and $\rho = 0.8$ (right panel) with $n = 200$ (solid line) and 500 (dot dash line).

Table 5. The frequency (%) of the estimated structural dimension $\hat{d}$.

Model	$\hat{d} = 1$	$\hat{d} = 2$	$\hat{d} \geq 3$	$\hat{d} = 1$	$\hat{d} = 2$	$\hat{d} \geq 3$
	$\rho = 0.5, n = 200$			$\rho = 0.8, n = 200$		
I	95.90	4.10	0.00	95.30	4.70	0.00
II	87.20	12.80	0.00	84.90	15.10	0.00
III	0.90	99.10	0.00	0.50	99.50	0.00
IV	0.00	100.00	0.00	0.00	100.00	0.00
	$\rho = 0.5, n = 500$			$\rho = 0.8, n = 500$		
I	98.80	1.20	0.00	98.80	1.20	0.00
II	94.90	5.10	0.00	95.30	4.70	0.00
III	0.20	99.80	0.00	0.40	99.60	0.00
IV	0.00	100.00	0.00	0.00	100.00	0.00

are univariate. When the response is univariate, the weighted and unweighted estimates of our proposal are identical, thus we only report the unweighted estimate. The MS method yields two estimates of β . We report the MS estimate with smaller bias and standard deviation. We compared their performance through a single index model.

Model V: A single-index model structure with a nonlinear link function

$$\mathbf{y} = 2(\beta^T \mathbf{x}) + \sin(4\beta^T \mathbf{x})\mathbf{z} + \varepsilon,$$

where $\mathbf{x}$ is drawn independently from multivariate normal distribution with zero mean and covariance matrix $(0.5^{|k-l|})$ and ε follows standard normal distribution. We set $p = 10$ and $\beta = (1, 0.8, 0.6, 0.4, 0.2, -0.2, -0.4, -0.6, -0.8, 0)^T$. We considered (i) $\mathbf{z} \sim \text{Bernoulli}(0.5)$, and (ii) $\mathbf{z} \sim \mathcal{N}(0, 0.5)$. To implement the LCC method when $\mathbf{z}$ is continuous, we discretize $\mathbf{z}$ into a series of binary variables

$I(\mathbf{z} \leq \tilde{\mathbf{z}})$, where $\tilde{\mathbf{z}}$ is an independent copy of $\mathbf{z}$ and $I(A)$ is an indicator function. For each given $\tilde{\mathbf{z}}$, we have an estimate of $\text{span}(\boldsymbol{\beta})$. We then pool all estimates to yield an integrated estimate of $\text{span}(\boldsymbol{\beta})$. For fair comparison, we rescaled the resulting estimate, $\hat{\boldsymbol{\beta}}$, obtained through existing methods so that the first entry of $\hat{\boldsymbol{\beta}}$, $\hat{\beta}_1$, is one. We repeated these scenario 1,000 times and report the biases and the Monte Carlo standard deviations of $(\hat{\boldsymbol{\beta}}/\hat{\beta}_1)$ in Table 6. The performance of all methods improves when the sample size n is increased from 200 to 500. In both cases, all methods perform comparatively, although our proposed NEW estimate has slightly smaller biases and standard deviations.

3.3. Application to Framingham Heart Study

In this section we revisit the Framingham Heart Study described in Section 1. Let $\mathbf{y} = (Y_1, Y_2)$, $\mathbf{z} = (1, Z_1, Z_2, Z_3)^\top$, $\mathbf{x} = (X_1, \dots, X_7)^\top$ in model (1.1). We added a column of ones in $\mathbf{z}$ to include an intercept in model (1.1). The BIC-type criterion finds $\hat{d} = 2$. The unweighted and the weighted-profile least squares estimates, along with their standard deviations and p-values, are given in Table 7. The weighted-profile least squares estimates have smaller standard deviations than those of the unweighted ones. In effect, $\text{corr}(Y_1, Y_2) = 0.4159$ and the p-value is less than 10^{-4} in the test for significance of Pearson's correlation coefficient. Thus, the systolic and diastolic blood pressures are highly correlated. It is then not surprising to see that the weighted-profile least squares estimates are significantly more efficient than the unweighted ones. For $k = 3, \dots, 7$, at least one p-value of X_k is significant at the significance level 0.05, indicating that the interactions between $\mathbf{x}$ and $\mathbf{z}$ are all significant. Therefore, we can conclude that healthy daily life styles, including a moderate amount of physical exercise, helps to control for the blood pressures. To show the interactions between $\mathbf{x}$ and $\mathbf{z}$ graphically, the estimated surfaces $\hat{m}_{ij}(\hat{\boldsymbol{\beta}}^\top \mathbf{x})$ of $m_{ij}(\boldsymbol{\beta}^\top \mathbf{x})$ with $\hat{\boldsymbol{\beta}} = (\hat{\beta}_1, \hat{\beta}_2)^\top$, for $i = 1, 2$ and $j = 2, 3, 4$ are shown in Figure 2, which clearly reveals the nonlinear modulating effect of the degree of obesity on physical exercises. Such dynamic effects are helpful in designing a moderate amount of physical exercise to control for the blood pressures.

4. Concluding Remarks

There are ancillary covariates $\mathbf{z} = (Z_1, Z_2, Z_3)^\top$ in our motivating example. These ancillary covariates are weakly correlated in that $\text{corr}(Z_1, Z_2) = -0.017$, $\text{corr}(Z_1, Z_3) = -0.060$ and $\text{corr}(Z_2, Z_3) = 0.101$, with p-values of 0.770, 0.297 and 0.079. Their correlations are not significant, thus we do not consider the

Table 6. Simulated results for Model V when $\mathbf{z} \sim \text{Bernoulli}(0.5)$ and $\mathbf{z} \sim \mathcal{N}(0, 0.5)$, respectively: the average bias (“bias”) and the Monte Carlo standard deviation (“std”) of $(\hat{\beta}/\hat{\beta}_1)$. All simulation results reported below are multiplied by 100.

method	n	$\mathbf{z}$	$\hat{\beta}_2$	$\hat{\beta}_3$	$\hat{\beta}_4$	$\hat{\beta}_5$	$\hat{\beta}_6$	$\hat{\beta}_7$	$\hat{\beta}_8$	$\hat{\beta}_9$	$\hat{\beta}_{10}$		
			0.8	0.6	0.4	0.2	-0.2	-0.4	-0.6	-0.8	0		
NEW	200	Bernoulli	bias	0.05	0.04	0.16	0.09	-0.20	0.08	-0.21	-0.12	-0.05	
			std	6.97	5.75	5.08	4.94	4.93	5.19	5.61	5.88	4.32	
		Normal	bias	0.01	0.30	-0.15	0.09	-0.06	-0.01	-0.23	-0.02	-0.07	
			std	6.69	5.10	5.10	4.95	4.95	5.09	5.37	5.59	4.17	
	500	Bernoulli	bias	0.44	0.10	0.25	0.07	-0.08	-0.21	-0.08	-0.39	-0.02	
			std	4.37	3.41	3.27	3.15	3.17	3.28	3.48	3.72	2.80	
		Normal	bias	0.30	0.32	-0.07	0.18	0.05	-0.20	-0.12	-0.21	0	
			std	4.11	3.25	3.19	2.95	3.04	3.18	3.33	3.57	2.71	
	LCC	200	Bernoulli	bias	0.36	0.17	-0.36	0.02	0.17	-0.36	-0.16	-0.39	-0.03
				std	8.18	6.15	5.85	5.72	5.60	5.94	6.39	6.76	4.79
			Normal	bias	0.57	0.33	0.06	0.01	0.08	-0.16	-0.05	-0.48	0.15
				std	7.96	6.04	5.62	5.24	5.37	5.49	5.99	6.33	4.53
500		Bernoulli	bias	0.15	-0.11	0.24	0.05	-0.17	0.00	-0.03	0.03	0.07	
			std	4.95	3.66	3.45	3.24	3.49	3.44	3.77	3.89	2.94	
		Normal	bias	-0.03	0.05	-0.05	-0.06	0.05	-0.00	0.11	-0.13	0.12	
			std	4.64	3.51	3.20	3.11	3.21	3.37	3.50	3.90	2.78	
MS		200	Bernoulli	bias	0.26	0.32	0.11	0.15	0.06	-0.02	-0.11	-0.24	-0.12
				std	9.35	6.69	6.53	6.01	5.98	6.22	6.75	7.54	5.34
			Normal	bias	0.38	0.24	0.18	0.06	0.07	-0.17	-0.33	0.04	-0.13
				std	8.27	6.11	5.92	5.57	5.61	5.74	6.35	6.83	5.06
	500	Bernoulli	bias	0.21	0.09	0.05	0.14	-0.10	-0.21	0.15	-0.15	-0.09	
			std	5.06	3.85	3.60	3.37	3.48	3.70	4.02	4.27	3.14	
		Normal	bias	0.02	-0.04	0.03	-0.03	0.04	-0.08	-0.05	0.05	0.04	
			std	4.61	3.63	3.45	3.35	3.32	3.53	3.55	4.03	2.88	
	LCL	200	Bernoulli	bias	0.65	0.34	-0.02	0.17	0.09	-0.28	-0.51	-0.45	6.43
				std	8.37	6.10	5.91	5.58	5.40	5.89	6.37	7.05	12.22
			Normal	bias	0.65	0.63	0.11	0.06	-0.08	-0.29	-0.13	-0.57	3.69
				std	7.89	6.16	5.84	5.61	5.31	5.65	6.13	6.73	9.41
500		Bernoulli	bias	0.36	0.05	0.08	0.11	0.06	-0.23	-0.26	0.19	0.39	
			std	4.69	3.57	3.17	3.05	3.22	3.26	3.69	3.75	2.42	
		Normal	bias	0.22	0.18	-0.14	-0.09	0.12	-0.12	-0.13	-0.07	0.32	
			std	4.67	3.50	3.21	3.20	3.26	3.38	3.65	3.94	2.21	

interactions among these ancillary covariates here. In model (1.1) we assume implicitly that the effects of the ancillary covariates are additive. Our proposed methodology and associated theoretical results are applicable when the ancillary covariates $\mathbf{z}$ are moderately correlated. If the ancillary covariates $\mathbf{z}$ are highly correlated, it is recommended considering the interactions among the $\mathbf{z}$ as well.

Table 7. Both the unweighted and the weighted profile least squares estimates, along with the standard errors and the p-values.

W		X_3		X_4		X_5		X_6		X_7	
		β_1	β_2	β_1	β_2	β_1	β_2	β_1	β_2	β_1	β_2
I	coef	0.4285	0.5149	1.0917	0.4939	0.9340	0.5008	0.5007	0.5024	0.2376	0.4919
	std	0.3805	0.0542	0.6388	0.0925	0.2754	0.0374	0.3115	0.0470	0.3343	0.0485
	p-value	0.2602	0.0000	0.0875	0.0000	0.0007	0.0000	0.1079	0.0000	0.4773	0.0000
$\hat{\Sigma}^{-1}$	coef	0.3811	0.5215	0.8881	0.5140	1.0031	0.5090	0.4575	0.4732	0.4355	0.4988
	std	0.3119	0.0440	0.5384	0.0741	0.2431	0.0306	0.2615	0.0382	0.2819	0.0402
	p-value	0.2217	0.0000	0.0990	0.0000	0.0000	0.0000	0.0802	0.0000	0.1224	0.0000

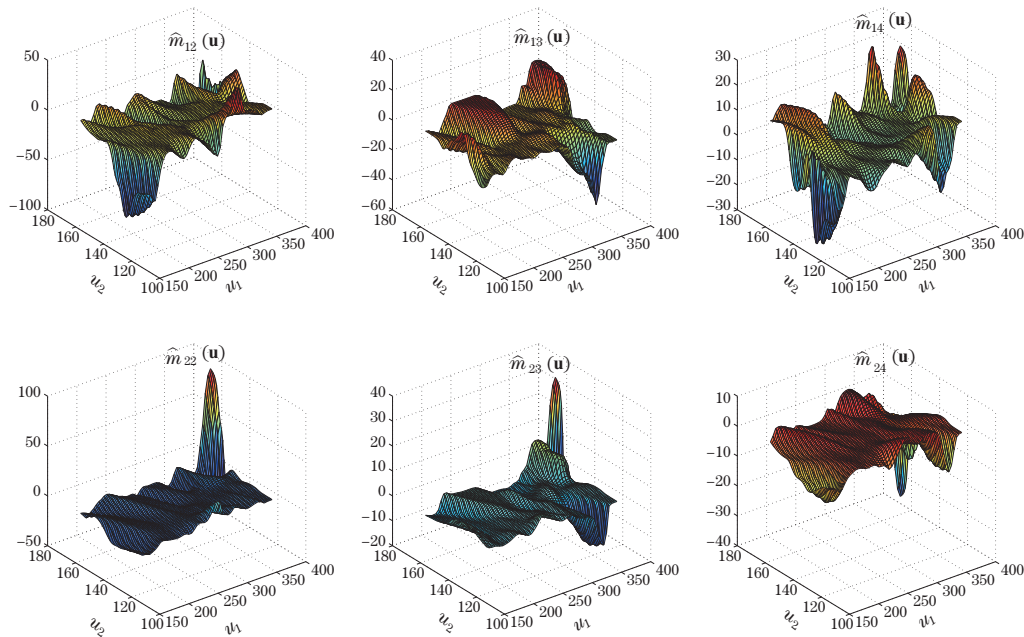


Figure 2. The estimated surfaces $\hat{m}_{ij}(\hat{\beta}^T \mathbf{x})$ of $m_{ij}(\hat{\beta}^T \mathbf{x})$ with $\hat{\beta} = (\hat{\beta}_1, \hat{\beta}_2)$, for $i = 1, 2$ and $j = 2, 3, 4$.

This leads to quite different model structures. Accordingly, new algorithms and estimation procedures are needed. Future research along this line is warranted.

Supplementary Materials

Some related models and comments on relevant methods, additional simulations and proofs of theorems can be found in the Supplementary Materials.

Acknowledgment

The authors thank the Editor, an Associated Editor and two reviewers for their constructive comments, which have led to a dramatic improvement of the earlier version of this article. Fan's work is supported by National Natural Science Foundation of China (11401006), Chinese Postdoctoral Science Foundation (2017M611083) and a Visiting Program for Young Scholar in Universities (gxfx2017042), Anhui Provincial Education Department, China. Zhu's work is supported by National Natural Science Foundation of China (11731011), Ministry of Education Project of Key Research Institute of Humanities and Social Sciences at Universities (16JJD910002) and National Youth Top-notch Talent Support Program, China.

References

- Carroll, R. J., Fan, J., Gijbels, I. and Wand, M. P. (1997). Generalized partially linear single-index models. *Journal of the American Statistical Association* **92**, 477–489.
- Cook, R. D. (1998). *Regression Graphics*. Wiley, New York.
- Cook, R. D. and Li, B. (2002). Dimension reduction for conditional mean in regression. *The Annals of Statistics* **30**, 455–474.
- Fan, J. and Gijbels, I. (1996). *Local Polynomial Modelling and Its Applications*. Chapman and Hall, London.
- Feng, Z., Wen, X. M., Yu, Z. and Zhu, L. (2013). On partial sufficient dimension reduction with applications to partially linear multi-index models. *Journal of the American Statistical Association* **108**, 237–246.
- Hilafu, H. and Wu, W. (2017). Partial projective resampling method for dimension reduction: with applications to partially linear models. *Computational Statistics & Data Analysis* **109**, 1–14.
- Ichimura, H. (1993). Semiparametric least squares (SLS) and weighted SLS estimation of single-index models. *Journal of Econometrics* **58**, 71–120.
- Li, B., Cook, R. D. and Chiaromonte, F. (2003). Dimension reduction for the conditional mean in regressions with categorical predictors. *The Annals of Statistics* **31**, 1636–1668.
- Liu, X., Cui, Y. and Li, R. (2016). Partial linear varying multi-index coefficient model for integrative gene-environment interactions. *Statistica Sinica* **26**, 1037–1060.
- Ma, S. and Song, P. X. K. (2015). Varying index coefficient models. *Journal of the American Statistical Association* **110**, 341–356.
- Ma, S., Yang, L., Romero, R. and Cui, Y. (2011). Varying coefficient model for gene-environment interaction: a non-linear look. *Bioinformatics* **27**, 2119–2126.
- Ma, Y. and Zhu, L. (2012). A semiparametric approach to dimension reduction. *Journal of the American Statistical Association* **107**, 168–179.
- Ma, Y. and Zhu, L. (2014). On estimation efficiency of the central mean subspace. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* **76**, 885–901.

- Ma, Y. and Zhu, L. P. (2013). Efficiency loss caused by linearity condition in dimension reduction. *Biometrika* **100**, 371–383..
- Xu, K., Guo, W., Xiong, M., Zhu, L. and Jin, L. (2016). An estimating equation approach to dimension reduction for longitudinal data. *Biometrika* **103**, 189–203.
- Zhu, L., Miao, B. and Peng, H. (2006). On sliced inverse regression with high-dimensional covariates. *Journal of the American Statistical Association* **101**, 630–643.
- Zhu, L. and Zhong, W. (2015). Estimation and inference on central mean subspace for multivariate response data. *Computational Statistics & Data Analysis* **92**, 68–83.

School of Economics and Management, Shanghai Maritime University, Shanghai 201306, P. R. China and Institute of Statistics and Big Data, Renmin University of China, 59 Zhongguancun Avenue, Beijing 100872, P. R. China.

E-mail: guoliangfan@yahoo.com

Center for Applied Statistics and Institute of Statistics and Big Data, Renmin University of China, 59 Zhongguancun Avenue, Haidian District, Beijing 100872, P. R. China.

E-mail: zhu.liping@ruc.edu.cn

Department of Statistics, University of California Riverside, CA 92521, USA.

E-mail: shujie.ma@ucr.edu

(Received May 2017; accepted October 2017)