

Fused estimators of the central subspace in sufficient dimension reduction

R. Dennis Cook^{*} and Xin Zhang[†]

Abstract

When studying the regression of a univariate variable Y on a vector $\mathbf{x}$ of predictors, most existing sufficient dimension reduction (SDR) methods require the construction of *slices* of Y in order to estimate moments of the conditional distribution of $\mathbf{X}$ given Y . But there is no widely accepted method for choosing the number of slices, while a poorly chosen slicing scheme may produce miserable results. We propose a novel and easily implemented fusing method that can mitigate the problem of choosing a slicing scheme and improve estimation efficiency at the same time. We develop two fused estimators – called FIRE and DIRE – based on an optimal inverse regression estimator. The asymptotic variance of FIRE is no larger than that of the original methods regardless of the choice of slicing scheme, while DIRE is less computationally intense and more robust. Simulation studies show that the fused estimators perform effectively the same as or substantially better than the parent methods. Fused estimators based on other methods can be developed in parallel: fused SIR, fused CSS-SIR and fused LAD are introduced briefly. Simulation studies of the fused CSS-SIR and fused LAD estimators show substantial gain over their parent methods. A real data example is also presented for illustration and comparison.

Key Words: Central solution space; central subspace; cumulative mean estimation; inverse regression estimator; sliced average variance estimation; sliced inverse regression; sufficient dimension reduction.

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

[†]Xin Zhang is Assistant Professor, Department of Statistics, Florida State University, Tallahassee, FL, 32306 (Email: henry@stat.fsu.edu).

1 Introduction

It is often useful in regression to concentrate the information about a univariate response Y in a low-dimensional function of the random predictor vector $\mathbf{X} \in \mathbb{R}^p$ without pre-specifying a parsimonious parametric model. This may be useful for mitigating the effects of collinearity, facilitating model specification by allowing visualization of the regression in low dimensions (Cook, 1998) and providing a relatively small set of “composite” predictors on which to base prediction or interpretation. When visualization is the goal, reducing the dimensionality of $\mathbf{X}$ may be useful when p exceeds 2 or 3 since these bounds represent the limits of our ability to view a data set in full.

In particular, the goal of sufficient dimension reduction (SDR) is to replace $\mathbf{X}$ by a lower dimensional function $\mathbf{R}(\mathbf{X}) \in \mathbb{R}^q$, $q < p$, without loss of information on the regression. Usually, $\mathbf{R}(\mathbf{X})$ is constrained to be a linear function $\boldsymbol{\eta}^T \mathbf{X}$, where $\boldsymbol{\eta} \in \mathbb{R}^{p \times q}$, $q < p$. The space spanned by the columns of $\boldsymbol{\eta}$ is then called a dimension reduction subspace for the regression of Y on $\mathbf{X}$. In most SDR contexts, we are interested in finding a population meta-parameter, the central subspaces (CS), represented by $\mathcal{S}_{Y|\mathbf{X}}$ and define as the intersection of all dimension reduction subspaces when this intersection is also a dimension reduction subspace. It is well-known that the CS exists under mild conditions (Cook, 1998). Let $d = \dim(\mathcal{S}_{Y|\mathbf{X}})$. Any basis $\boldsymbol{\beta} \in \mathbb{R}^{p \times d}$, $d < p$, of the CS satisfies three equivalent properties, (a) $Y \perp\!\!\!\perp \mathbf{X}|\boldsymbol{\beta}^T \mathbf{X}$, (b) $\mathbf{X}|(Y, \boldsymbol{\beta}^T \mathbf{X}) \sim \mathbf{X}|\boldsymbol{\beta}^T \mathbf{X}$, and (c) $Y|\mathbf{X} \sim Y|\boldsymbol{\beta}^T \mathbf{X}$. This implies that we can replace $\mathbf{X}$ by a lower dimensional projection $\boldsymbol{\beta}^T \mathbf{X}$ without loss of information on the regression, and that the CS is the smallest of the dimension reduction subspaces.

Sliced inverse regression (SIR; Li, 1991) and sliced average variance estimation (SAVE; Cook and Weisberg, 1991) were the first methods proposed for SDR. For regressions with both quantitative and categorical predictors, Chiamonte, Cook and Li (2002) extended SIR to partial SIR for estimating a partial central subspace that reduces only the quantitative predictors. Ye and Weiss (2003) proposed to combine SIR and SAVE by using

a bootstrap method to select the coefficients in a convex combination method. Li and Wang (2007) proposed directional regression (DR) as a dimension reduction method derived from empirical directions. DR synthesizes SIR and SAVE but requires substantially less computation than the bootstrap method proposed by Ye and Weiss (2003). Cook and Ni (2005) found the optimal method in the inverse regression family in terms of asymptotic efficiency. Cook and Forzani (2009) developed a likelihood-based method (LAD), which inherits properties and methods from general likelihood theory. LAD is based on normality assumptions, but is $\sqrt{n}$ consistent and able to perform better than classical methods with deviations from normality.

Most of these and other dimension reduction methods, including cluster-based methods (Setodji and Cook, 2004) and methods that aim to estimate the central solution space (CSS; Li and Dong, 2009; Dong and Li, 2010), hinge on slicing a quantitative response; that is, on replacing Y with a discrete version $\tilde{Y}$ in order to approximate the conditional moments of $\mathbf{X}|Y$ from $\mathbf{X}|\tilde{Y}$. This in effect approximates the moments of $\mathbf{X}|Y$ with step functions. The effectiveness of a slicing scheme depends on the number of slices and on their arrangement, which in turn can depend on the sample size n , the dimension d of the CS, the joint distribution the data and other characteristics of the regression. How to choose a good slicing scheme is a theoretically difficult but crucial issue that has resisted practically useful solutions over the past 20 years and that is often a severe nag in applications. In this article we propose a practically useful solution to the slicing dilemma.

The problem of choosing the number of slices was first noticed when Li (1991) established SIR. Nevertheless, it was not paid much attention, perhaps since Li claimed in the same article that even if the slice number is $n/2$, so that each slice contains only two observations, the resulting estimate will still be $\sqrt{n}$ consistent. Hsing and Carroll (1992) derived asymptotic properties for the special case where each slice has only two observations. This was generalized by Zhu and Ng (1995) to the case that each slice has an arbitrary but fixed number of observations, and to the case that the number of observations

within each slice goes to infinity with increasing sample size. They derived the asymptotic variance of SIR's kernel matrix with a fixed number c of observation per slice and observed that if $c = O(n^\alpha)$, $\alpha \in (0, \frac{1}{2})$, then the asymptotic variance could be reduced to a minimum. The choice of α depends on the distribution of $\mathbf{X}$ and on the smoothness of $\mathbf{m}(y) = E(\mathbf{X}|Y = y)$. This can be interpreted as requiring that the number of observations per slice be large enough to yield efficient estimates, but still be relatively small when comparing with n . This result from Zhu and Ng (1995) suggests that slicing schemes with too many slices having too few observations per slice should be avoided. Nevertheless, it is still difficult to turn such qualitative conclusions into practically useful rules.

In practice, slicing Y according to its quantiles, so we have approximately the same number of observations in each slice, may produce satisfactory results. We will refer to this as *quantile slicing*. However, the performance of quantile slicing still depends on the number of slices especially for higher-order methods (Li and Zhu, 2007). Such issues make choosing a good slicing scheme crucial, but there is no widely accepted rule for choosing the number of slices.

On the other hand, fusing information from different slicing schemes could effectively circumvent the pursuit of a best slicing scheme. There have been many studies on forecast combinations since Bates and Granger (1969); see Timmermann (2006) for an overview. However, forecast combination methods can not be applied directly in SDR without a clear predictive goal. Recently, Zhu, Zhu and Feng (2010) developed cumulative mean estimation (CUME), which uses a weighted average of SIR kernel matrices from all possible slicing schemes with two slices. CUME will be used for comparisons to our proposed method in simulation studies.

In this article we propose a method of fusing quantile slicing schemes that largely avoids the long-standing problem of selecting the number of slices and can result in substantially improved estimators of the CS. Our formal development is based on fusing

through minimum discrepancy functions similar to the asymptotically optimal inverse regression estimator (IRE; Cook and Ni, 2005), although the methodology can be applied straightforwardly to most SDR methods that rely on slicing a quantitative response, including SIR, SAVE, DR, and CSS-based methods. Our fusing method not only recovers intraslice information and therefore mitigates the complexity of choosing a good slicing scheme, but can also substantially improve estimation efficiency over methods with single slicing schemes.

The rest of this paper is organized as follows. In Section 2, the inverse regression methodology is reviewed, and two fusing methods corresponding to different inner product matrices are introduced. Theoretic properties of our fusing methods are given in Section 3, where we compare the asymptotic variances of our fused estimators to estimators based on a single quantile slicing scheme. In Section 4, we discuss computing and the use of robust inner product matrices. Section 5 includes simulation studies and Section 6 gives a brief discussion. Proofs and additional simulations are contained in the Supplement to this article.

Throughout this article, we will use bold symbols for matrices and vectors, and β for a semi-orthogonal basis of the CS, so $\text{span}(\beta) = \mathcal{S}_{Y|\mathbf{X}}$. Also we will use Σ for the covariance matrix of the predictor vector $\mathbf{X}$, and Γ for covariance matrices in fusing methods. We use quantile slicing for all methods, unless otherwise specified.

2 Methodology and Estimation

Most of the existing dimension reduction methods rely on an assumption about the marginal distribution of $\mathbf{X}$, called the linearity condition. The linearity condition requires that $E(\mathbf{X}|\beta^T \mathbf{X})$ is a linear function of $\beta^T \mathbf{X}$. This condition holds for all elliptically distributed $\mathbf{X}$ (Eaton, 1986) and it is also asymptotically true when the number of predictors $p \rightarrow \infty$ with $d = \dim(\mathcal{S}_{Y|\mathbf{X}})$ fixed (Hall and Li, 1993). When it holds, $\Sigma^{-1}\{E(\mathbf{X}|Y = y) - E(\mathbf{X})\} \in \mathcal{S}_{Y|\mathbf{X}}$ for all y . Thus by varying y in the domain of Y and estimating

$\Sigma^{-1}\{E(\mathbf{X}|Y = y) - E(\mathbf{X})\}$, we can obtain directions in the CS. When Y is categorical, it is straightforward to construct a sample version of $E(\mathbf{X}|Y = y)$ by averaging $\mathbf{X}$ within each class of Y . When Y is continuous, most SDR methods rely on nonparametric estimation of $E(\mathbf{X}|Y = y)$. Such nonparametric estimation requires partitioning the range of Y into h slices and then constructing a discrete version $\tilde{Y}$ of Y in order to obtain directions in the CS. It is always true that $\mathcal{S}_{\tilde{Y}|\mathbf{X}} \subseteq \mathcal{S}_{Y|\mathbf{X}}$. The accuracy of estimating the CS via inverse regression thus depends on slicing the response Y , as mentioned in the Introduction. At this point, given a slicing scheme, the inverse regression estimator (IRE) is an optimal member of the inverse regression (IR) family, and it has at least three desirable properties: its estimated basis is asymptotically efficient given the slicing scheme; it has an asymptotic chi-squared statistic for testing hypotheses about d ; and it provides for testing conditional independence hypotheses that the response is independent of a selected subset of predictors given the remaining predictors. Motivated by this, we introduce in this section a fusing method based on IRE objective functions following a review of the IRE method.

2.1 Review of IRE

Consider a univariate response variable Y , a $p \times 1$ vector of predictors $\mathbf{X}$ and the standardized predictor $\mathbf{Z} = \Sigma^{-1/2}(\mathbf{X} - E(\mathbf{X}))$, where $\Sigma = \text{cov}(\mathbf{X}) > 0$. Then $\mathcal{S}_{Y|\mathbf{X}} = \Sigma^{-1/2}\mathcal{S}_{Y|\mathbf{Z}}$ (Cook 1998; Proposition 6.1). Further assume that Y has a finite support $\{1, 2, \dots, h\}$ after slicing, and then define

$$\boldsymbol{\xi}_y = \Sigma^{-1}(E(\mathbf{X}|Y = y) - E(\mathbf{X})), \text{ and } \mathcal{S}_{\boldsymbol{\xi}} = \sum_{y=1}^h \text{span}(\boldsymbol{\xi}_y).$$

The linearity condition implies $\mathcal{S}_{\boldsymbol{\xi}} \subseteq \mathcal{S}_{Y|\mathbf{X}}$. We further require the coverage condition to enforce the equality of the two subspaces $\mathcal{S}_{\boldsymbol{\xi}} = \mathcal{S}_{Y|\mathbf{X}}$. The coverage condition requires that the target subspace $\mathcal{S}_{\boldsymbol{\xi}}$ has the same dimension as $\mathcal{S}_{Y|\mathbf{X}}$. See Cook and Ni (2005) and Li and Wang (2007) for discussion on the coverage condition. To focus our discussion on

fusing, we assume that the dimension of the CS is known: $d = \dim(\mathcal{S}_{Y|X}) = \dim(\mathcal{S}_\xi)$. By definition, for each y , there exists a vector γ_y such that $\xi_y = \beta\gamma_y$. Define

$$\xi = (\xi_1, \dots, \xi_h) = \beta\gamma = \beta(\gamma_1, \dots, \gamma_h).$$

and let $\mathbf{f} = (f_1, \dots, f_h)^T$ with $f_y = \Pr(Y = y)$. It is easy to see that $\xi\mathbf{f} = \beta\gamma\mathbf{f} = 0$, called the intrinsic location constraint by Cook and Ni (2005).

After gathering a simple random sample $(\mathbf{X}_i, Y_i)$, $i = 1, 2, \dots, n$, from $(\mathbf{X}, Y)$, we let $\mathbf{X}_{yj}$ denote the j -th observation on $\mathbf{X}$ in slice y , $y = 1, \dots, h$, $j = 1, \dots, n_y$, and $\sum_{y=1}^h n_y = n$. Further, let $\bar{\mathbf{X}}$ be the overall average of $\mathbf{X}$, and let $\bar{\mathbf{X}}_y$ be the average of the n_y points with $Y = y$. Let $\hat{\mathbf{f}} = (n_1/n, \dots, n_h/n)^T$, and let $\hat{\Sigma} > 0$ be the usual sample covariance matrix for $\mathbf{X}$. The sample version of ξ is

$$\hat{\xi} = (\hat{\xi}_1, \dots, \hat{\xi}_h) \in \mathbb{R}^{p \times h}, \quad \hat{\xi}_y = \hat{\Sigma}^{-1}(\bar{\mathbf{X}}_y - \bar{\mathbf{X}}), \quad y = 1, \dots, h.$$

Let $\mathbf{D}_{\hat{\mathbf{f}}}$ denote a diagonal matrix with the elements of the vector $\hat{\mathbf{f}}$ on the diagonal. Because of the intrinsic location constraint, $\hat{\xi}\mathbf{D}_{\hat{\mathbf{f}}}\mathbf{1}_h = \hat{\xi}\hat{\mathbf{f}} = 0$. Construct $\mathbf{A} \in \mathbb{R}^{h \times (h-1)}$, such that $\mathbf{A}^T \mathbf{A} = \mathbf{I}_{h-1}$ and $\mathbf{A}^T \mathbf{1}_h = 0$, then use the reduced matrix $\zeta \equiv \beta\gamma\mathbf{D}_{\hat{\mathbf{f}}}\mathbf{A} = \beta\nu$ and its sample version $\hat{\zeta} = \hat{\xi}\mathbf{D}_{\hat{\mathbf{f}}}\mathbf{A}$ in the construction of discrepancy functions without loss of generality.

With these characters established, the quadratic discrepancy function of IRE is defined as

$$F_d(\mathbf{B}, \mathbf{C}) = (\text{vec}(\hat{\zeta}) - \text{vec}(\mathbf{B}\mathbf{C}))^T \hat{\Gamma}_{\hat{\zeta}}^{-1} (\text{vec}(\hat{\zeta}) - \text{vec}(\mathbf{B}\mathbf{C})), \quad (2.1)$$

where the columns of $\mathbf{B} \in \mathbb{R}^{p \times d}$ represent an orthogonal basis for $\text{span}(\zeta)$, which equals $\mathcal{S}_{Y|X}$ under the linearity and coverage conditions, but may be a proper subset of $\mathcal{S}_{Y|X}$ with linearity but not coverage; $\mathbf{C} \in \mathbb{R}^{d \times (h-1)}$ is used only in fitting to represent the coordinates of ζ relative to $\mathbf{B}$, and $\hat{\Gamma}_{\hat{\zeta}} \in \mathbb{R}^{p(h-1) \times p(h-1)}$ is a consistent estimator of the asymptotic variance of $\text{vec}(\hat{\zeta})$. Replacing $\hat{\Gamma}_{\hat{\zeta}}^{-1}$ in (2.1) with an arbitrary nonsingular matrix, we still obtain a $\sqrt{n}$ consistent estimator of $\mathcal{S}_{Y|X} = \text{span}(\zeta)$, but using $\hat{\Gamma}_{\hat{\zeta}}^{-1}$ guarantees an asymptotically efficient estimator. For instance, Cook and Ni (2005) showed

that SIR is a member of a suboptimal class using an inner product matrix different from $\hat{\mathbf{\Gamma}}_{\hat{\boldsymbol{\zeta}}}^{-1}$. Minimizing $F_d(\mathbf{B}, \mathbf{C})$ produces the IRE.

A distinct advantage of IRE is its ability to properly weight the slice means when the intra-slice variance $\text{var}(\mathbf{X}|Y = j)$ changes appreciably with the slice. In contrast, SIR and other methods like CUME that neglect the intra-slice variances can produce poor results when there are substantial differences in variability across the slices.

However, the performance of IRE is still based on the tuning parameter h , the number of slices, and on the method for constructing the slices themselves. Instead of pursuing the “best” h , we propose a fusing approach in the following section where a set of possible choices of slice numbers $H = \{h_1, \dots, h_K\}$ is used rather than a particular value h .

2.2 Fused Inverse Regression Estimators

With a set of $K \geq 2$ quantile slicing schemes with slice numbers $H = \{h_1, \dots, h_K\}$, we have K different data matrices, denoted by

$$\hat{\boldsymbol{\xi}}^{(j)} \in \mathbb{R}^{p \times h_j}, \hat{\boldsymbol{\zeta}}^{(j)} \equiv \hat{\boldsymbol{\xi}}^{(j)} \mathbf{D}_j \mathbf{A}_j \in \mathbb{R}^{p \times (h_j - 1)}, j = 1, \dots, K,$$

where $\hat{\mathbf{D}}_j \equiv \text{diag}(n_1^{(j)}/n, \dots, n_{h_j}^{(j)}/n)$ is a diagonal matrix consisting of the sample proportions of slices in the j -th slicing scheme. Constant matrices $\mathbf{A}_j \in \mathbb{R}^{h_j \times (h_j - 1)}$, $j = 1, \dots, K$, are associated with the intrinsic location constraint in each slicing scheme, so that $\mathbf{A}_j^T \mathbf{A}_j = \mathbf{I}_{h_j - 1}$ and $\mathbf{A}_j^T \mathbf{1}_{h_j} = 0$. Here the data matrices $\hat{\boldsymbol{\zeta}}^{(j)}$ converge in probability at rate $\sqrt{n}$ to

$$\boldsymbol{\zeta}^{(j)} \equiv \boldsymbol{\beta} \boldsymbol{\gamma}_j \mathbf{D}_j \mathbf{A}_j = \boldsymbol{\beta} \boldsymbol{\nu}_j \in \mathbb{R}^{p \times (h_j - 1)}, j = 1, \dots, K.$$

In this equation, $\hat{\mathbf{D}}_j$ converges to $\mathbf{D}_j$, and $\boldsymbol{\beta} \in \mathbb{R}^{p \times d}$ is a basis for $\mathcal{S}_{Y|\mathbf{X}}$ such that the matrices $\boldsymbol{\zeta}^{(j)}$, $j = 1, 2, \dots, K$, can be written as the products of $\boldsymbol{\beta}$ and the coordinate matrices $\boldsymbol{\nu}_j \in \mathbb{R}^{d \times (h_j - 1)}$. The coordinate matrices $\boldsymbol{\nu}_j$ will share the same row dimension $d = \dim(\mathcal{S}_{Y|\mathbf{X}})$, but different column dimension due to varied slice numbers.

In order to fuse information from various slicing schemes, the most fundamental step is to fuse the K data matrices by row, since they have the same number of rows and each row corresponds to one dimension in the CS. Therefore the fused data matrix is $\hat{\zeta}_F = (\hat{\zeta}^{(1)}, \dots, \hat{\zeta}^{(K)}) \in \mathbb{R}^{d \times (\sum_{j=1}^K h_j - K)}$ which converges in probability to $\zeta_F = (\zeta^{(1)}, \dots, \zeta^{(K)}) = \beta(\nu_1, \dots, \nu_K)$. The fused inverse regression estimator follows from minimizing a quadratic discrepancy function with the fused data matrix:

$$\mathcal{F}_d(\mathbf{B}, \mathbf{C}; \Psi) = (\text{vec}(\hat{\zeta}_F) - \text{vec}(\mathbf{B}\mathbf{C}))^T \hat{\Psi} (\text{vec}(\hat{\zeta}_F) - \text{vec}(\mathbf{B}\mathbf{C})), \quad (2.2)$$

where the columns of $\mathbf{B} \in \mathbb{R}^{p \times d}$ still represent an orthogonal basis of our estimate of the CS, $\mathbf{C} \in \mathbb{R}^{d \times \sum_j (h_j - 1)} = \mathbb{R}^{d \times (\sum_j h_j - K)}$ is used only in fitting and it represents the coordinates of ζ_F relative to $\mathbf{B}$. Moreover, $\hat{\Psi}$ is a consistent estimator of a positive semi-definite matrix $\Psi \in \mathbb{R}^{p(\sum_j h_j - K) \times p(\sum_j h_j - K)}$ and is essentially the inner product matrix of this weighted least squares route. The estimate of the CS is then constructed by minimizing the quadratic discrepancy function in (2.2), over $\mathbf{B}$ and $\mathbf{C}$. Since $\mathcal{F}_d(\mathbf{B}, \mathbf{C}; \Psi)$ is an over-parameterized discrepancy function, the minimizer $(\hat{\mathbf{B}}, \hat{\mathbf{C}})$ is not unique: for any orthogonal matrix Λ that $\Lambda\Lambda^T = \mathbf{I}_d$, $(\hat{\mathbf{B}}\Lambda, \Lambda^T\hat{\mathbf{C}})$ is also a minimizer. However, the product $\hat{\mathbf{B}}\hat{\mathbf{C}}$ is unique, as well as the CS estimator $\text{span}(\hat{\mathbf{B}})$ and its projection $\hat{\mathbf{B}}\hat{\mathbf{B}}^T$.

3 Asymptotic properties

We study the asymptotic properties of the fused inverse regression estimators in this section. We re-write the quadratic discrepancy function of IRE with slicing scheme j as

$$F_d^{(j)}(\mathbf{B}, \mathbf{C}) = (\text{vec}(\hat{\zeta}^{(j)}) - \text{vec}(\mathbf{B}\mathbf{C}))^T \hat{\Gamma}_j^{-1} (\text{vec}(\hat{\zeta}^{(j)}) - \text{vec}(\mathbf{B}\mathbf{C})), \quad j = 1, 2, \quad (3.1)$$

which is just (2.1) with specific slicing schemes, and where $\hat{\Gamma}_j \equiv \hat{\Gamma}_{\hat{\zeta}_j}$ is a consistent estimator of the non-singular positive definite matrix $\Gamma_j \equiv \Gamma_{\zeta_j}$. Without loss of generality, we will discuss fusing only two schemes, i.e. $K = 2$ and $H = \{h_1, h_2\}$, along with comparisons to IRE with $h = h_1$. The results hold for $K > 2$ automatically.

3.1 Asymptotic normality

Prior to deriving the asymptotic properties of the fused estimators, we need to find the asymptotic distribution of the fused data matrix $\hat{\boldsymbol{\zeta}}_F$. As stated in Theorem 1 from Cook and Ni (2005), $\sqrt{n}(\text{vec}(\hat{\boldsymbol{\zeta}}^{(j)}) - \text{vec}(\boldsymbol{\zeta}^{(j)}))$, $j = 1, 2$, converges in distribution to a normal with mean 0 and nonsingular covariance matrix $\boldsymbol{\Gamma}_j$. For each slicing scheme, define h_j random variables $J_y^{(j)}$ such that $J_y^{(j)}$ equals 1 if an observation is in the y -th slice of j -th slicing scheme and 0 otherwise, $j = 1, 2$, $y = 1, \dots, h_j$. Then $E(J_y^{(j)}) = (\mathbf{D}_j)_{yy} = f_y^{(j)}$. Also define the random vector $\boldsymbol{\epsilon}^{(j)} = (\epsilon_1^{(j)}, \dots, \epsilon_{h_j}^{(j)})^T$ with elements

$$\epsilon_y^{(j)} = J_y^{(j)} - E(J_y^{(j)}) - \mathbf{Z}^T E(\mathbf{Z} J_y^{(j)})$$

that are the population residuals from the ordinary least squares fit of $J_y^{(j)}$ on $\mathbf{Z}$. Denote the fused population residual vector by

$$\mathbf{E} \equiv (\boldsymbol{\epsilon}^{(1)T}, \boldsymbol{\epsilon}^{(2)T})^T = (\epsilon_1^{(1)}, \dots, \epsilon_{h_1}^{(1)}, \epsilon_1^{(2)}, \dots, \epsilon_{h_2}^{(2)})^T \in \mathbb{R}^{(h_1+h_2) \times 1}. \quad (3.2)$$

The following theorem gives the asymptotic distribution of the data matrix $\hat{\boldsymbol{\zeta}}_F$ that we have used in the fused estimators.

Theorem 1. *Assume that the data $(\mathbf{X}_i, Y_i)$, $i = 1, \dots, n$, is a simple random sample of $(\mathbf{X}, Y)$ with finite fourth moments. Then*

$$\sqrt{n}(\text{vec}(\hat{\boldsymbol{\xi}}^{(1)} \hat{\mathbf{D}}_1, \hat{\boldsymbol{\xi}}^{(2)} \hat{\mathbf{D}}_2) - \text{vec}(\boldsymbol{\xi}^{(1)} \mathbf{D}_1, \boldsymbol{\xi}^{(2)} \mathbf{D}_2)) \longrightarrow N(0, \boldsymbol{\Gamma}_H),$$

where $\boldsymbol{\Gamma}_H = \text{cov}(\text{vec}(\boldsymbol{\Sigma}^{-1/2} \mathbf{Z} \mathbf{E}^T))$ and $\mathbf{E}$ is given in (3.2). Consequently,

$$\sqrt{n}(\text{vec}(\hat{\boldsymbol{\zeta}}_F) - \text{vec}(\boldsymbol{\zeta}_F)) \longrightarrow N(0, \boldsymbol{\Gamma}_F),$$

where $\boldsymbol{\Gamma}_F = (\mathbf{A}_F^T \otimes \mathbf{I}_p) \boldsymbol{\Gamma}_H (\mathbf{A}_F \otimes \mathbf{I}_p)$, and $\mathbf{A}_F = \text{diag}(\mathbf{A}_1, \mathbf{A}_2)$ is a non-stochastic block diagonal matrix.

The inverse of $\boldsymbol{\Gamma}_F$ will be used later, but $\boldsymbol{\Gamma}_F$ is possibly singular even though all $\boldsymbol{\Gamma}_j$, $j = 1, \dots, K$, are assumed to be non-singular. The singularity of $\boldsymbol{\Gamma}_F$ occurs when one of the

K data matrices $\hat{\zeta}^{(j)}$ can be expressed as a linear combination of the other data matrices. For example, suppose we have a 5-slice scheme and a 10-slice scheme obtained by splitting each of the 5 slices. Then every element of the data matrix in the 5-slice scheme will be a linear combination of the data matrix in the 10-slice scheme. Fusing these two will lead to a singular Γ_F because of the 5 redundant slices. Singularity of Γ_F can be avoided by using quantile slicing schemes that do not share the same slice boundaries. We can fuse a 5-slice scheme and a 11-slice scheme without encounter singular Γ_F . Unless indicated otherwise, we will assume that the numbers of slices $H = \{h_1, \dots, h_K\}$ are properly chosen such that Γ_F is non-singular. See Section 5 for more discussion on choosing the set of numbers of slices $H = \{h_1, \dots, h_K\}$.

3.2 Consistency of the fused estimators

With the asymptotic distribution of the data matrix $\hat{\zeta}_F$, we can establish the $\sqrt{n}$ consistency of the fused estimator from (2.2). Let $(\hat{\beta}_j, \hat{\nu}_j)$ be minimizer of the IRE discrepancy function with $h = h_j$, i.e. (3.1); let $(\hat{\beta}_F, \hat{\omega}_F) \equiv (\hat{\beta}_F, \hat{\omega}_1, \hat{\omega}_2)$ be a minimizer of the fused discrepancy function (2.2). As discussed in Section 2, the minimizers $(\hat{\beta}_j, \hat{\nu}_j)$, $j = 1, 2$, and $(\hat{\beta}_F, \hat{\omega}_F)$ are not unique. But we have uniqueness in the products: $\hat{\beta}_j \hat{\nu}_j$ and $\hat{\beta}_F \hat{\omega}_F$. Moreover, the two products: $\hat{\beta}_j \hat{\nu}_j$ and $\hat{\beta}_F \hat{\omega}_j$ both converge to the matrix $\zeta^{(j)}$ for $j = 1, 2$. Therefore, we can judge our fused estimators by comparing asymptotic behaviors of $\hat{\beta}_F \hat{\omega}_j$ to $\hat{\beta}_1 \hat{\nu}_1$. The result is summarized in the following theorem.

Theorem 2. *Assume that the data $(\mathbf{X}_i, Y_i)$, $i = 1, \dots, n$, are a simple random sample of $(\mathbf{X}, Y)$ with finite fourth moments and assume that $\hat{\Psi}$ converges in probability to $\Psi > 0$ as n goes to infinity. Then $\hat{\beta}_F \hat{\omega}_F = (\hat{\beta}_F \hat{\omega}_1, \hat{\beta}_F \hat{\omega}_2)$ and $(\hat{\beta}_1 \hat{\nu}_1, \hat{\beta}_2 \hat{\nu}_2)$ are both $\sqrt{n}$ -consistent estimators for $\zeta_F = (\zeta^{(1)}, \zeta^{(2)})$ such that for $j = 1, 2$,*

$$\begin{aligned} \sqrt{n}(\text{vec}(\hat{\beta}_j \hat{\nu}_j) - \text{vec}(\zeta^{(j)})) &\longrightarrow N(0, \mathbf{V}_j), \\ \sqrt{n}(\text{vec}(\hat{\beta}_F \hat{\omega}_j) - \text{vec}(\zeta^{(j)})) &\longrightarrow N(0, \mathbf{W}_j), \end{aligned}$$

where the explicit expressions for $\mathbf{V}_j$ and $\mathbf{W}_j$ are given in the Appendix. In particular, if $\Psi = \Gamma_F^{-1}$, then the fused estimator is asymptotically efficient with respect to the fused data $\hat{\zeta}_F$ and the asymptotic variances satisfy

$$\mathbf{W}_j \leq \mathbf{V}_j, \quad j = 1, 2. \quad (3.3)$$

All the results of Theorem 2 extend straightforwardly to slicing schemes with $K \geq 2$. The estimators of the CS constructed from minimizing $\mathcal{F}(\mathbf{B}, \mathbf{C}; \Gamma_F^{-1})$ are called FIRE.

A direct consequence of (3.3) is that the asymptotic variances of $\text{vec}(\hat{\beta}_F \hat{\beta}_F^T)$ for FIRE will be no larger than the asymptotic variances of $\text{vec}(\hat{\beta}_1 \hat{\beta}_1^T)$ or $\text{vec}(\hat{\beta}_2 \hat{\beta}_2^T)$. This shows that FIRE is asymptotically at least as efficient at estimating the CS as IRE with the corresponding slicing schemes. Another straightforward generalization of this theorem is obtained by considering $H = \{h_1, \dots, h_K\}$ and $H_1 \subset H$, a proper subset of H . The asymptotic variances of FIRE based on H are smaller than or equal to the asymptotic variances of FIRE based on H_1 . Notice that IRE is asymptotically efficient for only a single data matrix $\hat{\zeta}^{(j)}$, but not necessarily asymptotically efficient with respect to the fused data matrix $\hat{\zeta}_F$. On the other hand, Theorem 2 states FIRE is asymptotically efficient with respect to $\hat{\zeta}_F$, because this method is essentially fusing a set of quantile slicing schemes H to a single fused slicing scheme.

Although the boundaries of a fused slicing scheme are at quantiles, these quantiles are not arranged so each slice has the same number of observations and thus fused slicing is not the same as a single quantile slicing scheme. When the number of observations in each slice goes to infinity, this fused slicing scheme has asymptotic variance no larger than that for IRE.

An example may help fix ideas. Assume we have 120 iid samples with distinct Y_i 's and $H = \{4, 5\}$. Let y_p denote the p -th quantile of Y , i.e. $y_p = \inf\{y : \sum_{i=1}^n I(Y_i \leq y)/n \leq p\}$. The quantile slicing scheme for the IRE estimator with $h = 4$ (similarly for $h = 5$) is obtained from $\{y_{0.25}, y_{0.5}, y_{0.75}\}$ with four slices consist of 30 observations each. Therefore, the fused data matrix will consist of four slices with 30 observations

each and 5 slices with 24 observations each. FIRE uses the full inner product matrix. A full rank linear transformation of the fused data matrix $\hat{\zeta}_F$ will result in a corresponding linear transformation of the inner product matrix. Since there is a linear transformation between the fused data matrix and a single data matrix obtained from eight slices with slice boundaries $\{y_{0.2}, y_{0.25}, y_{0.4}, y_{0.5}, y_{0.6}, y_{0.75}, y_{0.8}\}$, FIRE is exactly the same as IRE with this (non-quantile) fused slicing scheme having eight slices with 24, 6, 18, 12, 12, 18, 6, and 24 observations per slice. However, IRE with this fused slicing scheme is different from IRE with a single quantile slicing scheme having $h = 8$ with 15 observations per slice. In our simulation studies (Section 5) we see clear advantages of fused slicing schemes over quantile slicing schemes with the same numbers of slices.

The improvement from fusing over simply increasing the number of quantile slices is remarkable. Slicing on Y and then using the discrete approximation of $E(\mathbf{X}|Y)$ in the objective function is essentially an approximation to some integral with respect to the probability density function of Y . Results from Luceno (1999) imply that a properly chosen slicing scheme with h slices could guarantee that the discrete version $\tilde{Y}$ has the same r -th moments as Y for $r = 0, \dots, 2h - 1$. Since high order moments are not important, composite rules with small numbers of grid points in numerical integration are often more accurate than directly applying the integration formula with a large number of grid points. Fusing estimators are analogous to composite rules in numerical integration. Fusing several crude slicing schemes with small values of h often leads to more accurate basis estimation than IRE with a large h . We will demonstrate this result in Section 5.

3.3 Inner product matrix Ψ

Our estimation based on minimizing the quadratic discrepancy function defined in (2.2) is largely related to the choice of Ψ . FIRE uses the full inverse of the asymptotic covariance matrix of the vectorized fused data matrix $\text{vec}(\hat{\zeta}_F) = \text{vec}(\hat{\zeta}^{(1)}, \hat{\zeta}^{(2)})$, which makes itself asymptotically efficient. However, estimation of such inner product matrix is difficult

because of its dimensionality. Because a poor estimate of the inner product matrix could lead to poor performance, we propose simpler alternatives to the full inner product matrix $\Psi = \Gamma_F^{-1}$ in this section.

The first natural simplification is by taking $\Psi = \text{diag}\{\Gamma_1^{-1}, \Gamma_2^{-1}\} \equiv \Gamma_D^{-1}$. The fused estimators constructed from minimizing $\mathcal{F}(\mathbf{B}, \mathbf{C}; \Gamma_D^{-1})$ are called DIRE. The block diagonal inner product matrix is obtained by pretending that the slicing schemes to be fused are independent, which reduces the number of parameters in the inner product matrix and may improve estimation accuracy. Additionally, using a block diagonal inner product matrix corresponds to simply adding the objective functions for the individual slicing schemes to obtain the fused objective function, which is a procedure that can be used straightforwardly to fuse quantile slicing schemes for any method based on an objective function. Moreover, we demonstrate in Section 6 that this block-diagonal fusing approach is actually a composite likelihood approach (Lindsay 1988) in likelihood-based dimension reduction (Cook and Forzani 2009).

For small sample problems, Ni and Cook (2007) introduced a robust version of the IRE inner product matrix, which requires only second order moments rather than fourth order moments. They showed that the robust IRE produces a $\sqrt{n}$ consistent basis estimator of the CS and was quite robust in small sample simulations. Therefore, we will use this robust version of the inner product matrices in our fused methods for small sample problems. Replacing the previous residual variables $\epsilon_y^{(j)} = J_y^{(j)} - \mathbb{E}(J_y^{(j)}) - \mathbf{Z}^T \mathbb{E}(\mathbf{Z} J_y^{(j)})$ by $\varepsilon_y^{(j)} = J_y^{(j)} - \mathbb{E}(J_y^{(j)})$, we will obtain the robust version of inner product matrices:

$$\mathbf{G}_j = (\mathbf{A}_j^T \otimes \Sigma^{-1/2}) \text{cov}\{\mathbf{Z}(\varepsilon^{(j)})^T\}(\mathbf{A}_j \otimes \Sigma^{-1/2}), \quad (3.4)$$

where $\varepsilon^{(j)} = (\varepsilon_1^{(j)}, \dots, \varepsilon_{h_j}^{(j)})^T \in \mathbb{R}^{h_j}$ is the vector of modified residuals. Then we similarly define the fused inner product matrix $\mathbf{G}_F^{-1}$ and $\mathbf{G}_D^{-1}$ in responding to Γ_F^{-1} and Γ_D^{-1} . The estimators of the CS obtained from minimizing $\mathcal{F}_d(\mathbf{B}, \mathbf{C}; \mathbf{G}_F^{-1})$ and $\mathcal{F}_d(\mathbf{B}, \mathbf{C}; \mathbf{G}_D^{-1})$ will be called the robust FIRE and robust DIRE. And their $\sqrt{n}$ -consistency follows directly from Theorem 2. We will show simulation results in Section 5 regarding both the

original fused estimators and the robust fused estimators.

FIRE preserves asymptotic optimality while DIRE and the robust estimators are all asymptotically suboptimal. However, in a finite sample problem, DIRE often performs better than both FIRE and IRE because estimating Γ_D^{-1} well requires no more observations than estimating Γ_j with the largest $h_j \in H$. Simulation studies and real data application in Section 5 also indicate that DIRE is better than or close to FIRE in most situations. FIRE is preferred over DIRE only when n is very large comparing to p , see Section 5.3 for a real data example where $n = 1030$ and $p = 8$. Moreover, it is straightforward to generalize DIRE to other methods by pretending independent slicing schemes. At the same time, the $\sqrt{n}$ consistency and asymptotic variance properties of DIRE could be maintained. We illustrate this by adapting DIRE for fusing SIR estimators in the next section.

3.4 Degenerate slicing schemes and fused SIR

From Section 3.2, we know that FIRE is asymptotically at least as efficient as IRE with any single slicing scheme in H . As a referee pointed out, it would be helpful to guide practice if we know when a single slicing scheme IRE could be asymptotically efficient with respect to the fused data matrix. Then IRE based on this slicing scheme would have the same asymptotic performance as FIRE.

Corollary 1. *Under the conditions of Theorem 2, the following statements are equivalent.*

1. *IRE with slicing scheme h_k is asymptotic efficient with respect to $\hat{\zeta}_F$.*
2. $\text{avar}(\sqrt{n}\hat{\beta}_k\hat{\beta}_k^T) = \text{avar}(\sqrt{n}\hat{\beta}_F\hat{\beta}_F^T).$
3. $\Gamma_j - \Gamma_{kj}^T\Gamma_k^{-1}\Gamma_{kj} = 0$ for all $j \neq k$.
4. $\Gamma_F = \mathbf{O} \begin{pmatrix} \Gamma_k & 0 \\ 0 & 0 \end{pmatrix} \mathbf{O}^T$ for some full rank transformation $\mathbf{O}$.

The situation in Corollary 1 happens when H is a set of poorly chosen degenerated slicing schemes, as discussed in Section 3.1. Another possible cause of such situation is that $\mathbf{X}$ depends on Y only through some latent slices of Y and one slicing scheme h_k happens to coincide with that latent structure. In our first simulation example in Section 5, we simulated data in this way and observed that FIRE had approximately the same performance to IRE with the latent slicing scheme, while DIRE, being asymptotically suboptimal, had even improved finite performance.

To elaborate on the above discussion, we use the inverse regression model of Cook et al. (2012). In addition to the linearity and the coverage conditions, they assumed the covariance condition that $\text{var}(\mathbf{X}|Y)$ is non-stochastic, which led to the model

$$\mathbf{X}_i = \boldsymbol{\mu} + \boldsymbol{\eta} \mathbf{f}(Y_i) + \boldsymbol{\varepsilon}_i, \quad i = 1, \dots, n, \quad (3.5)$$

where $\boldsymbol{\mu} = \mathbb{E}(\mathbf{X})$ and $\boldsymbol{\varepsilon}_i \sim N(0, \boldsymbol{\Phi})$ and is independent of Y and the vector-valued functions $\mathbf{f}(Y_i) \in \mathbb{R}^d$. Since we are focusing on slicing instead of a general functional estimation, we assume $\mathbf{f}(Y_i)$ is a linear function of slice indicators $J_y^{(k)}, y = 1, \dots, h_k$, of some particular single slicing scheme h_k . Neither $\boldsymbol{\eta}$ nor $\mathbf{f}(Y_i)$ is identified in model (3.5), because for any non-singular $d \times d$ matrix $\mathbf{A}$, $\boldsymbol{\eta} \mathbf{A} \mathbf{A}^{-1} \mathbf{f}(Y_i) = \boldsymbol{\eta} \mathbf{f}(Y_i)$ leads to a different parameterization. We could add additional constraint $\boldsymbol{\eta}^T \boldsymbol{\eta} = \mathbf{I}_d$ to get uniqueness. Nevertheless, the central subspace is affected only by $\text{span}(\boldsymbol{\eta} \mathbf{f}(Y_i))$.

Lemma 1. *Under model (3.5), IRE with slicing scheme h_k is asymptotically efficient with respect to $\hat{\boldsymbol{\zeta}}_F$, where H is any set of slicing schemes such that $h_k \in H$. Moreover, if the errors in (3.5) are isotropic, $\boldsymbol{\Phi} = \sigma^2 \mathbf{I}_p$, then IRE is asymptotically equivalent to SIR which is the MLE under these assumptions.*

Lemma 1 implies that we could judge the necessity of fusing based on model (3.5): if this model holds for some known h_k , then fusing would not be necessary. However, in practice, it is difficult to check such conditions and almost impossible to find such h_k for a continuous Y , and then fusing will alleviate the burden of choosing slicing schemes.

Lemma 1 also implies that SIR could be used instead of IRE for isotropic errors. SIR is based on the spectral decomposition of a kernel matrix $\widehat{\mathbf{M}}_{\text{SIR}} = \sum_{y=1}^h \hat{f}_y \bar{\mathbf{Z}}_y \bar{\mathbf{Z}}_y^T$, where $\bar{\mathbf{Z}}_y = \hat{\Sigma}^{-1/2}(\bar{\mathbf{X}}_y - \bar{\mathbf{X}})$ is the sample average of the n_y points in slice y computed in the $\mathbf{Z}$ -scale. Cook and Ni (2005) shows that SIR is covered by the inverse regression (IR) family with the following objective function,

$$F_d^{\text{SIR}}(\mathbf{B}, \mathbf{C}; h) = (\text{vec}(\hat{\xi} \mathbf{D}_{\hat{f}}) - \text{vec}(\mathbf{B}\mathbf{C}))^T \text{diag}\{\hat{f}_y^{-1} \hat{\Sigma}\} (\text{vec}(\hat{\xi} \mathbf{D}_{\hat{f}}) - \text{vec}(\mathbf{B}\mathbf{C})), \quad (3.6)$$

where the columns of $\mathbf{B} \in \mathbb{R}^{p \times d}$ represent a basis of SIR's estimator of the CS and $\mathbf{C} \in \mathbb{R}^{d \times h}$ is again a coordinate matrix used only in fitting. Let $(\hat{\mathbf{b}}, \hat{\mathbf{c}})$ be the minimizer of (3.6). Then $\text{span}(\hat{\mathbf{b}})$ is the same as the span of the d eigenvectors of $\widehat{\mathbf{M}}_{\text{SIR}}$ corresponding to its largest d eigenvalues.

Let $H = \{h_1, \dots, h_K\}$ denote a set of K quantile slicing schemes. Analogous to DIRE, we define the fused SIR estimator via its objective function

$$F_d^{\text{FSIR}}(\mathbf{B}, \mathbf{C}) = \sum_{j=1}^K F_d^{\text{SIR}}(\mathbf{B}, \mathbf{C}_j; h_j),$$

where $\mathbf{C} = (\mathbf{C}_1, \dots, \mathbf{C}_K) \in \mathbb{R}^{d \times \sum h_j}$. Let $(\hat{\mathbf{b}}_1, \hat{\mathbf{c}}_1)$ be minimizer of $F_d^{\text{SIR}}(\mathbf{B}, \mathbf{C}_1; h_1)$; let $(\hat{\mathbf{b}}_F, \hat{\mathbf{c}}_F) \equiv (\hat{\mathbf{b}}_F, \hat{\mathbf{v}}_1, \dots, \hat{\mathbf{v}}_K)$ be a minimizer of the fused SIR objective function $F_d^{\text{FSIR}}(\mathbf{B}, \mathbf{C})$. In practice, it is unnecessary to numerically minimize the objective function $F_d^{\text{FSIR}}(\mathbf{B}, \mathbf{C})$ since we can obtain an explicit form of the minimizer. The next lemma provides a spectral decomposition approach for fused SIR.

Lemma 2. *Assume that the data $(\mathbf{X}_i, Y_i)$, $i = 1, \dots, n$, are a simple random sample of $(\mathbf{X}, Y)$ with finite second moments. Let $\hat{\mathbf{u}}_1, \dots, \hat{\mathbf{u}}_p$ be the eigenvectors of $\widehat{\mathbf{M}}_{\text{FSIR}} = \sum_{j=1}^K \sum_{y=1}^{h_j} \hat{f}_y^{(j)} \bar{\mathbf{Z}}_y^{(j)} \bar{\mathbf{Z}}_y^{(j)T}$ corresponding to eigenvalues $\hat{\lambda}_1 > \dots > \hat{\lambda}_d > \hat{\lambda}_{d+1} \geq \dots \geq \hat{\lambda}_p$, where $\bar{\mathbf{Z}}_y^{(j)}$ is the sample average in the y -th slice of slicing scheme h_j . Then the minimizer $\text{span}(\hat{\mathbf{b}}_F) = \text{span}(\hat{\mathbf{u}}_1, \dots, \hat{\mathbf{u}}_d)$.*

The kernel matrix $\widehat{\mathbf{M}}_{\text{FSIR}}$ is the sum of K kernel matrices for SIR, which is very intuitive by treating the K slicing schemes independently. In the same way, fused estimators

of SAVE (Cook and Weisberg 1991) and DR (Li and Wang 2007) can be constructed by spectral decomposing the sum of kernel matrices with different slicing schemes.

4 Estimation procedure and computing algorithm

We will use the data matrix $\hat{\zeta}_F$ and a consistent estimator of Γ_F^{-1} or a consistent estimator of Γ_D^{-1} for $\mathcal{F}_d(\mathbf{B}, \mathbf{C}; \Psi)$ in (2.2). Also, we adopt the same computing algorithm, the alternating least squares algorithm, as in Cook and Ni (2005), to minimize the objective function. Since Γ_D^{-1} is block diagonal with Γ_j^{-1} on the diagonal, we can simply replace it by the consistent estimator used in the IRE, $\hat{\Gamma}_j^{-1}$. So that $\text{diag}(\hat{\Gamma}_1^{-1}, \dots, \hat{\Gamma}_K^{-1})$ is a consistent estimator of Γ_D^{-1} . Similarly for $\Gamma_F^{-1} = [(\mathbf{A}_F^T \otimes \mathbf{I}_p) \Gamma_H (\mathbf{A}_F \otimes \mathbf{I}_p)]^{-1}$, we only need to estimate $\Gamma_H = \text{cov}(\text{vec}(\Sigma^{-1/2} \mathbf{Z} \mathbf{E}^T))$ by substituting the sample version $\hat{\mathbf{E}}_i = (\hat{\epsilon}_{1i}^{(1)}, \dots, \hat{\epsilon}_{h_1i}^{(1)}, \dots, \hat{\epsilon}_{1i}^{(K)}, \dots, \hat{\epsilon}_{h_Ki}^{(K)})^T$ and $\hat{\epsilon}_{yi}^{(j)} = J_{yi}^{(j)} - n_y^{(j)}/n - (\mathbf{X}_i - \bar{\mathbf{X}}_{..})^T \hat{\boldsymbol{\xi}}_y^{(j)} n_y^{(j)}/n$, $i = 1, \dots, n$, $y = 1, \dots, h_j$, for all $j = 1, \dots, K$. Then a consistent estimate of Γ_F^{-1} follows by noticing $\mathbf{A}_F = \text{diag}(\mathbf{A}_1, \dots, \mathbf{A}_K)$ is a block diagonal non-stochastic matrix. For the robust FIRE and robust DIRE, we apply the same estimation procedure except that $\hat{\epsilon}_{yi}^{(j)}$ was replaced by $\hat{\epsilon}_{yi}^{(j)} = J_{yi}^{(j)} - n_y^{(j)}/n$.

We discussed in Section 2 that we could guarantee the non-singularity of Γ_F by choosing slicing schemes with different slice boundaries. However, the $\sum_{j=1}^K (h_j - 1)p$ by $\sum_{j=1}^K (h_j - 1)p$ matrix $\hat{\Gamma}_F$ is possibly ill-conditioned and thus we use the Moore-Penrose generalized inverse of $\hat{\Gamma}_F$ instead of the regular inverse matrix in our algorithm. In Section 8, we discuss the possibility of using regularized estimators of $\hat{\Gamma}_F^{-1}$ to replace the generalized inverse.

5 Numerical studies

In this section we compare our fused inverse regression estimators to CUME, IRE and SIR using quantile slicing with various values of h . The simulation results based on

various settings support our theoretical conclusions. We then apply our method to a concrete compressive strength dataset (Yeh, 1998).

5.1 Inverse regression

We calculated the vector correlation coefficient q^2 (Hotelling, 1936), the trace correlation r^2 (Hooper, 1959) and angles $\boldsymbol{\theta} = (\theta_1, \dots, \theta_d)^T$ to measure the closeness between the estimated CS $\text{span}(\hat{\boldsymbol{\beta}})$ and the true CS $\mathcal{S}_{Y|\mathbf{X}} = \text{span}(\boldsymbol{\beta})$:

$$q^2(\hat{\boldsymbol{\beta}}, \boldsymbol{\beta}) = \prod_{i=1}^d \rho_i, \quad r^2(\hat{\boldsymbol{\beta}}, \boldsymbol{\beta}) = \frac{1}{d} \sum_{i=1}^d \rho_i,$$

where $\rho_1, \rho_2, \dots, \rho_d$ are ordered eigenvalues of $\boldsymbol{\beta}^T \hat{\boldsymbol{\beta}} \hat{\boldsymbol{\beta}}^T \boldsymbol{\beta}$. Both $q^2(\hat{\boldsymbol{\beta}}, \boldsymbol{\beta})$ and $r^2(\hat{\boldsymbol{\beta}}, \boldsymbol{\beta})$ range from 0 to 1, with values closer to 1 indicating that the two subspace $\text{span}(\hat{\boldsymbol{\beta}})$ and $\mathcal{S}_{Y|\mathbf{X}}$ are closer. Let $\hat{\boldsymbol{\beta}} = (\hat{\boldsymbol{\beta}}_1, \dots, \hat{\boldsymbol{\beta}}_d) \in \mathbb{R}^{p \times d}$. Then the angle θ_i is defined as the angle between $\hat{\boldsymbol{\beta}}_i$ and its projection onto $\mathcal{S}_{Y|\mathbf{X}}$. In the following simulation examples, we report the averaged values of q^2 , r^2 and $\boldsymbol{\theta}$ and their standard errors.

To imitate the asymptotic performance of the proposed two fused inverse regression estimators, we generated data with sample size $n = 600$, dimension $p = 6$ from a model with dimension $d = 2$. The model was constructed with a clear optimal choice of slicing scheme: 5 slices with equal number of observations in each slice. Y was simulated from the uniform distribution on the interval $[0, 5]$ and

$$\mathbf{X} = \boldsymbol{\Phi} \boldsymbol{\eta} \mathbf{f}(Y) + c(Y) \boldsymbol{\epsilon}, \quad (5.1)$$

where $\boldsymbol{\epsilon}$ is a multivariate normal random variable with mean 0, covariance matrix $\boldsymbol{\Sigma} > 0$, and is independent of Y . The vector valued function $\mathbf{f}(Y) \in \mathbb{R}^{r \times 1}$ was designed to consist of indicator functions of Y :

$$\mathbf{f}(Y) = (J_1, \dots, J_5)^T, \quad J_k = I(k-1 < Y \leq k), \quad k = 1, \dots, 5,$$

which is why 5 slices are optimal for model (5.1). We used $c(Y) \in \mathbb{R}^1$ to construct either homoscedastic or heteroscedastic error structure. The elements of both constant matrices

$\Phi \in \mathbb{R}^{p \times d}$ and $\eta \in \mathbb{R}^{d \times r}$ were generated as a random sample from the standard normal distribution and then standardized to be semi-orthogonal. From Cook and Forzani (2008), we know that for this setting $\mathcal{S}_{Y|\mathbf{X}} = \Sigma^{-1} \text{span}(\Phi)$.

Four different error structures were employed to study the fusing methods:

- (a) isotropic error structure $c(Y) = 0.1$ and $\Sigma = \mathbf{I}_p$;
- (b) autoregressive error structure $c(Y) = 0.1$, $(\Sigma)_{ij} = 0.5^{|i-j|}$;
- (c) heteroscedastic error structure $c(Y) = 0.01J_1 + 0.05(J_2 + J_3) + J_4 + 10J_5$, $\Sigma = \mathbf{I}_p$;
- (d) heteroscedastic error structure $c(Y) = 0.01J_1 + 0.05(J_2 + J_3) + J_4 + 10J_5$, $(\Sigma)_{ij} = 0.5^{|i-j|}$.

Table 1 shows the performances of SIR, IRE, CUME, FIRE and DIRE. For SIR and IRE, we reported two choices of numbers of slices, $h = 5$ and $h = 7$. This is because we know $h = 5$ is the optimal choice in this model, while $h = 7$ is just an arbitrary choice to represent a typical situation when we have no prior information about how to choose the slicing scheme. FIRE and DIRE were obtained from fusing a set of 13 quantile slicing schemes $H = \{3, \dots, 15\}$. This set likely covers all the sensible slicing schemes, so we have no responsibility in picking the “best” numbers of slices.

As can be seen in Table 1, in all the four error structure settings, both fused estimators outperformed CUME, SIR with $h = 7$ and IRE with $h = 7$. While at the same time, DIRE always outperformed FIRE. In a great majority of comparisons, DIRE was the best method overall. By comparing it to SIR with $h = 5$ and IRE with $h = 5$, the two classical method with the optimal number of slices, we notice that it was at least as good as if not better than IRE with $h = 5$. However, the fused methods lost to SIR with $h = 5$ in the simulation model with isotropic error structure because this simplest case is designed to best fit in the SIR context. SIR with 5 slices outperformed all other methods in the

	(a)		(b)		(c)		(d)	
	r^2	q^2	r^2	q^2	r^2	q^2	r^2	q^2
SIR h=5	0.971	0.943	0.969	0.938	0.657	0.352	0.644	0.332
SIR h=7	0.944	0.890	0.934	0.869	0.411	0.117	0.459	0.141
IRE h=5	0.947	0.894	0.958	0.916	0.957	0.916	0.936	0.874
IRE h=7	0.912	0.826	0.917	0.838	0.946	0.894	0.930	0.863
CUME	0.948	0.898	0.921	0.846	0.541	0.194	0.546	0.200
DIRE	0.961	0.922	0.968	0.936	0.949	0.900	0.942	0.886
FIRE	0.957	0.915	0.965	0.930	0.943	0.889	0.933	0.869
s.d.	0.0010	0.0020	0.0008	0.0016	0.0027	0.0052	0.0034	0.0064

Table 1: Comparisons of the fused estimators and other methods from 5-slice inverse model (5.1), with 4 different error structures (a)-(d)

isotropic error structures (a) and (b), because SIR with 5 slices gives the MLE for this problem (see Cook and Forzani, 2008), while IRE and the fused estimators had more parameters to estimate and lost some efficiency. Even with such isotropic error structure, the user still has to choose the correct number of slices for SIR to outperform our fused estimators. Although SIR with the best choice of slices could be more efficient than the fused estimators, the efficiency gain by SIR is very little. But in more complex situations, SIR as well as CUME may fail to capture the CS when IRE and fused estimators still work well. In the more challenging settings (c) and (d), where the error structures are heteroscedastic, SIR and CUME fails to capture the CS.

Figure 5.1 shows the failure of SIR and CUME with heteroscedastic error structure (d) in terms of the angles θ_1 and θ_2 , and shows the optimality of the fused estimators with respect to IRE using different numbers of slices. The dimension $d = \dim(\mathcal{S}_{Y|\mathbf{X}}) = 2$, so we used two angles to visually summarize the estimation results. In this summary plot, the horizontal axis represents the numbers of slices used in IRE and SIR. Because of the nature of this model, $h = 5$ produced the smallest angles for both IRE and SIR. As the graphs show, DIRE delivered more accurate estimates than FIRE and CUME.

To better show that the two fused estimators have smaller asymptotic variances in estimating the CS than the asymptotic variances of IRE with arbitrary number of slices

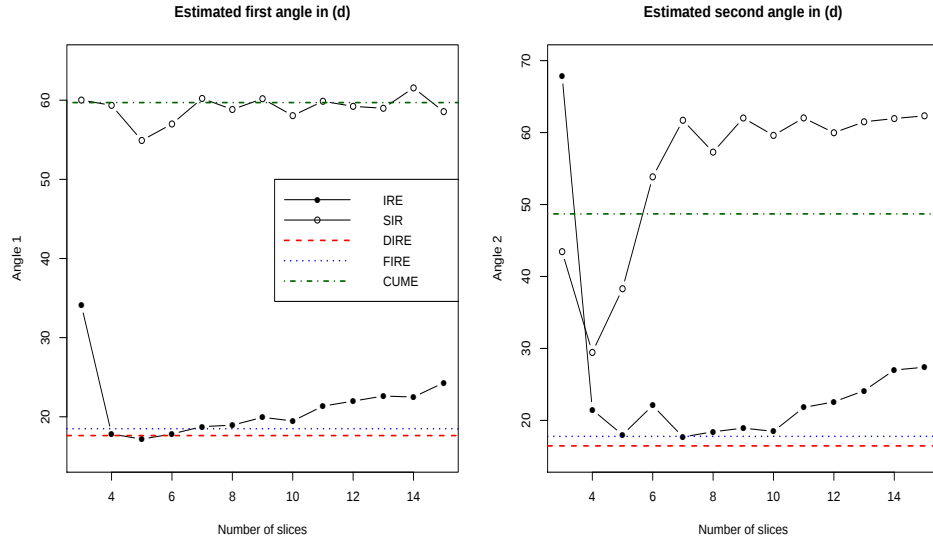


Figure 5.1: 5-slice inverse model (5.1) with heteroscedastic error structure (d).

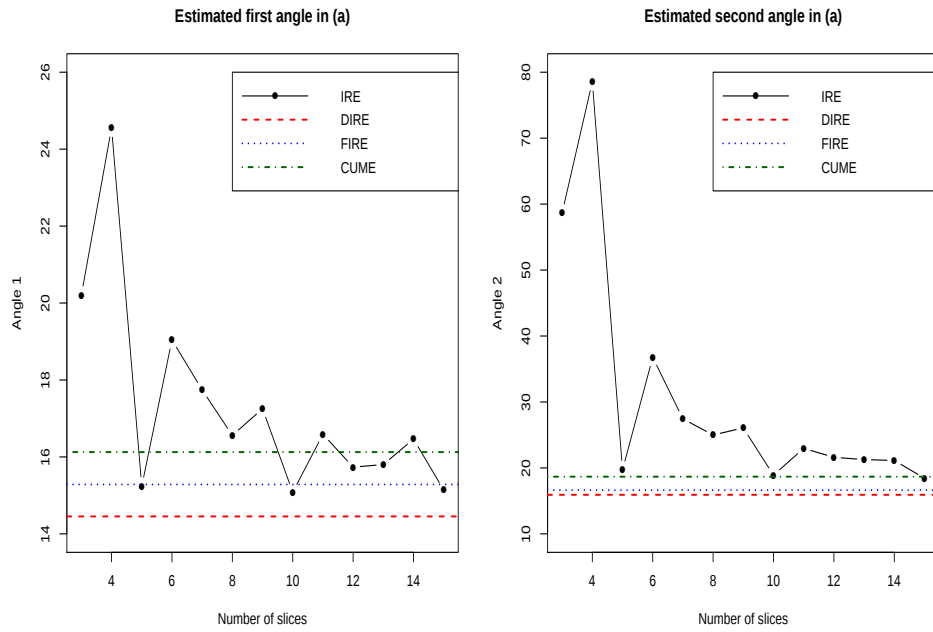


Figure 5.2: 5-slice inverse model (5.1) with isotropic error structure (a). SIR fluctuates around the line of CUME and is not included in the plot for better visualization.

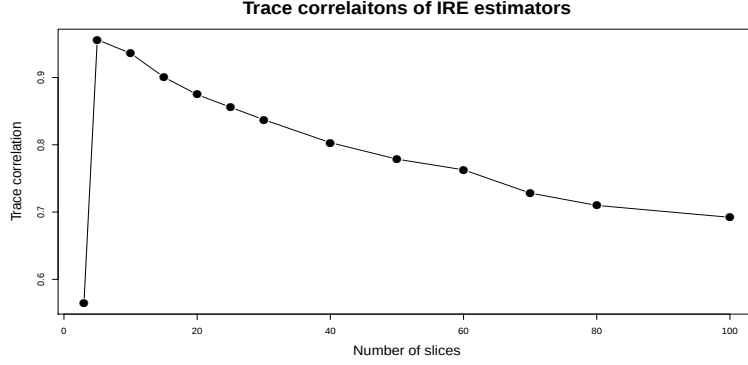


Figure 5.3: 5-slice inverse model with heteroscedastic error structure (d): trace correlations of the IRE estimators with different numbers of slices.

$h \in H = \{3, \dots, 15\}$, we plotted for the isotropic case (a) the two angles between the estimated CS and the true CS in Figure 5.2. From this plot, we can see clearly that DIRE has significantly smaller first and second angles than IRE even with 5 slices, while FIRE has similar first and second angles to IRE with 5 slices. And a small change in the numbers of slices can lead to substantial change in the output of IRE. For example, the second angle of IRE with $h = 6$ is about twice as that of $h = 5$. Fortunately, FIRE and DIRE are unaffected.

Figure 5.3 shows the behavior of the IRE estimators with h ranging from $h = 3$ to $h = 100$. In the 5-slice inverse model with the heteroscedastic error structure (d), the trace correlations of the IRE estimators are almost monotonically decreasing as the number of slices increases. This result was anticipated by Zhu and Ng (1995): too many slices with too few observations per slice is not good for the variances of estimates. In general, we find that the performance of IRE is usually monotonically decreasing in h after a moderate number. However, the fused slicing scheme of FIRE in this case has approximately 100 slices but can still perform well. Actually, FIRE and DIRE did much better than IRE with $h = 100$: From Figure 5.3 the trace correlation for IRE is about 0.70, while the trace correlations for DIRE and FIRE are from Table 1 0.942 and 0.933. As discussed in Section 3.2, fusing slicing schemes with small values of h , i.e. $h = 3, \dots, 15$, may lead

to much more accurate basis estimation than IRE with a large value $h = 100$.

5.2 Forward regression models

We have also studied many forward regression models and observed the following qualitative finite sample performance. (1) DIRE always outperforms FIRE and IRE with any h , unless the sample size is quite large and then FIRE performs better than IRE and a bit better than DIRE. We expect that this happens because the cost of estimating the FIRE inner product matrix can be substantially greater than the cost of estimating the DIRE inner product matrix. (2) When the dimension reduction subspace is easy to find, as is the case with a homoscedastic forward linear regression model, SIR with a reasonable number of slices can outperform FIRE, DIRE and IRE by a small amount. Of course, if we knew we had a simple forward model at the outset then SDR methods would be unnecessary. (3) In challenging forward regressions where the central subspace is not easy to find, FIRE and DIRE are substantially better than SIR and CUME and better than IRE by an amount that depends on the number of slices. One example of this behavior is described below; additional forward regression simulation are given in the Supplement.

Data $\{X_{i1}, \dots, X_{i15}, Y_i\}$, $i = 1, \dots, 400$, were generated as follows. First, we sampled $\mathbf{X}_i \in \mathbb{R}^{15}$ from a mixture of normal distributions

$$\mathbf{X} \sim \frac{1}{4}N(\boldsymbol{\mu}_1, \boldsymbol{\sigma}_1) + \frac{1}{2}N(\boldsymbol{\mu}_2, \boldsymbol{\sigma}_2) + \frac{1}{4}N(\boldsymbol{\mu}_3, \boldsymbol{\sigma}_3),$$

where $\boldsymbol{\mu}_1 = \boldsymbol{\mu}_3 = (1, 0, \dots, 0)^T$, $\boldsymbol{\mu}_2 = (2, 0, \dots, 0)^T$, $\boldsymbol{\sigma}_1 = \boldsymbol{\sigma}_2 = \sqrt{0.1}\mathbf{I}_{15}$ and $\boldsymbol{\sigma}_3 = \sqrt{10}\mathbf{I}_{15}$. Then the Y_i 's were generated as

$$Y = |\sin(X_1)| + 0.2\epsilon, \tag{5.2}$$

where ϵ is a uniform (0,1) random variable. Thus, $\mathcal{S}_{Y|\mathbf{X}} = \text{span}\{(1, 0, \dots, 0)^T\}$. Again, we simulated 500 datasets and summarized the angle θ between $\mathcal{S}_{Y|\mathbf{X}}$ and its estimate in the left panel of Figure 5.4, where we used $H = \{3, \dots, 15\}$ for both fused estimators. We adopted the robust version of the inner product matrices (Ni and Cook, 2007) for both

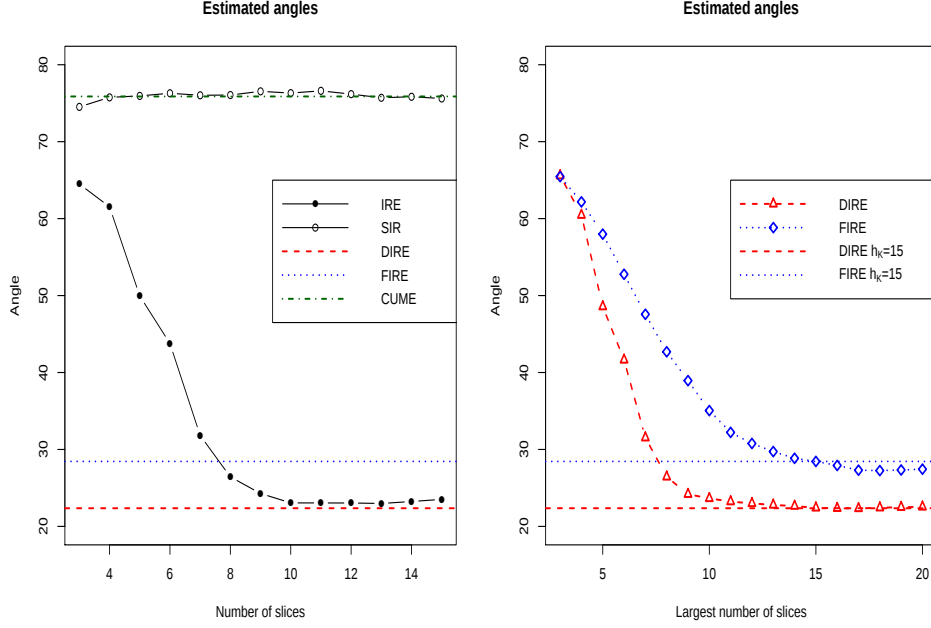


Figure 5.4: Estimated angle θ by different methods in the forward model (5.2). Left panel: θ versus h in a single quantile slicing scheme with the fused estimators and CUME superimposed. Right panel: θ versus the largest number of slices h_K in the fused estimators with $H = \{3, 4, \dots, h_K\}$.

IRE and the fused estimators. Since the intra-slice variances $\text{var}(\mathbf{X}|Y)$ are non-constant, SIR and CUME both suffered. These two methods completely failed to estimate the CS. On the other hand, FIRE and DIRE automatically adjusted weights on each slices through their inner product matrices $\hat{\mathbf{\Gamma}}_F^{-1}$ and $\hat{\mathbf{\Gamma}}_D^{-1}$, and thus estimated the CS efficiently. The left panel of Figure 5.4 shows that DIRE dominated FIRE in estimation, and outperformed IRE with any number of slices. Although it is a little difficult to see from the plot, the IRE angle starts to increase after about 10 slices, similar to the phenomenon shown in Figure 5.3.

The right panel of Figure 5.4 was constructed to illustrate the behavior of FIRE and DIRE as a function of h_K , the largest slice number in the fused estimators. Specifically, we plotted the average angle θ against h_K with $H = \{3, 4, \dots, h_K\}$. The behavior in

right panel of Figure 5.4 is typical from our experience: The fused estimators rapidly accumulate information as h_K increases through the small integers, eventually flattening with very little change thereafter. Because of such results we typically choose $10 \leq h_K \leq 20$ when fusing. We have not systematically studied the behavior of the fused estimators with $h_K > 20$.

5.3 Concrete compressive strength dataset

For this illustration (Yeh, 1998. <http://archive.ics.uci.edu/ml/datasets.html>), we want to predict the compressive strength of high-performance concrete (Mpa, Y) from seven ingredients ($X_1, \dots, X_7$) and age (days, X_8). These ingredients are cement (X_1), blast furnace slag (X_2), fly ash (X_3), water (X_4), superplasticizer (X_5), coarse aggregate (X_6), and fine aggregate (X_7) are all measured in kg/m^3 . We used all 1030 observations in estimation.

We applied SIR and IRE with various numbers of slices, and the usual dimension tests always indicated $d = 2$ with significant level 0.05. IRE dimension tests with different quantile slicing schemes typically result in the same conclusion for d . However, because of the discrete nature of the decision, it seems inevitable that the conclusion about d will vary with the slicing scheme in some datasets (see, for example, Ye and Weiss (2003)), and then additional analysis in the context of the fused estimator may be necessary to guide the final choice. A graphical study, the general permutation test proposed by Cook and Yin (2001) and some form of cross validation may all be useful in this regard, depending on application specific goals. It is also possible to develop dimension estimation based directly on the fused estimator, but that is outside the scope of this report.

We then plotted Y versus the first two orthogonal directions obtained from CUME, DIRE and FIRE in Figure 5.5, where we again chose $H = \{3, \dots, 15\}$ for DIRE and FIRE. We randomly sampled and plotted 200 out of the 1030 observations and a LOWESS smooth line based on all the 1030 observations with the same smoothing parameter. The

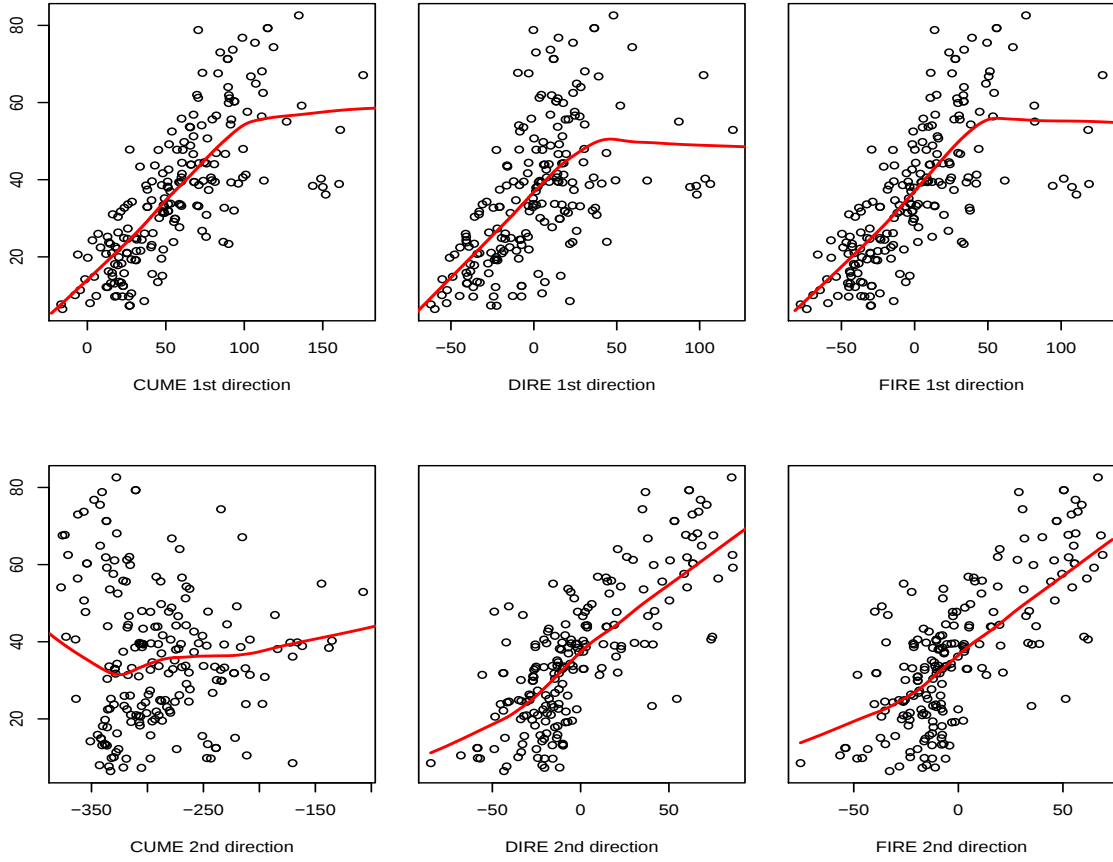


Figure 5.5: Concrete compressive strength dataset: scatterplots of the response (vertical axes) and the first two orthogonal directions (horizontal axes), with mean smooths added to aid visualization.

first directions of FIRE, DIRE and CUME are essentially the same. However, the second directions of FIRE and DIRE gave clear linear patterns while the second direction of CUME apparently provided little information regarding the response. The plots (not included) for SIR and IRE with various slicing schemes showed similar characteristics in the first direction and vague patterns in the second direction.

To further assess the estimation of the CS, we used the bootstrap criterion developed by Ye and Weiss (2003) and used by Zhu, Zhu and Feng (2010). We used the

	SIR			IRE			CUME	DIRE	FIRE
	h=5	h=10	h=15	h=5	h=10	h=15			
n=1030	0.699	0.685	0.674	0.730	0.779	0.643	0.721	0.820	0.915
n=400	0.618	0.596	0.613	0.703	0.701	0.623	0.628	0.771	0.806

Table 2: Concrete compressive strength dataset: averaged r^2 using bootstrap samples. The standard errors for the averaged r^2 are between 0.012 and 0.016 for $n = 1030$, and are between 0.010 and 0.018 for $n = 400$.

bootstrap to generate 100 dataset with sample size $n = 400$ and 100 dataset with sample size $n = 1030$. With $d = 2$ we then computed the trace correlation coefficient $r^2(\hat{\beta}, \hat{\beta}^b)$, $b = 1, \dots, 100$, between the originally estimated CS $\text{span}(\hat{\beta})$ and the bootstrap estimate $\text{span}(\hat{\beta}^b)$. The results are summarized in Table 2. From the table, we can see that FIRE has the largest r^2 , followed by DIRE. Therefore, FIRE is the best method for this dataset, which might have been anticipated since FIRE is the asymptotically optimal fused estimator.

6 Composite likelihoods and fused LAD

A good choice of slicing scheme may become more crucial in second order methods such as SAVE (Cook and Weisberg 1991), DR (Li and Wang 2007) and LAD (Cook and Forzani 2009). This is because the second order methods relies on higher moment of the conditional variable $\mathbf{X}_y \equiv (\mathbf{X} | J_y = 1)$, $y = 1, \dots, h$. In this section, we describe a fused version of LAD.

Within each slice y , the LAD model assumes conditional normality of $\mathbf{X}_y$: $\mathbf{X}_y \sim N(\boldsymbol{\mu}_y, \boldsymbol{\Delta}_y)$ with arbitrary mean function $\boldsymbol{\mu}_y$ and covariance $\boldsymbol{\Delta}_y > 0$. Let $\boldsymbol{\Delta} = E(\boldsymbol{\Delta}_Y) \equiv E\{\text{var}(\mathbf{X} | Y)\}$ and assume the existence of a d -dimensional CS. Then the MLE for the CS is obtained by maximizing the following likelihood-based objective function (Cook and Forzani 2009):

$$L_d(\boldsymbol{\beta}; h) = \frac{n}{2} \log |\boldsymbol{\beta}^T \hat{\boldsymbol{\Sigma}} \boldsymbol{\beta}| - \sum_{y=1}^h \frac{n_y}{2} \log |\boldsymbol{\beta}^T \hat{\boldsymbol{\Delta}}_y \boldsymbol{\beta}|, \quad (6.1)$$

where $\hat{\Sigma}$ is the sample covariance matrix of $\mathbf{X}$, $\hat{\Delta}_y$ is the sample covariance matrix of the data within slice y and the maximization is over the Grassmann manifold $\mathcal{G}_{(p,d)}$.

The objective function $L_d(\beta; h)$ is derived by partially maximizing the log-likelihood function $\ell(\mu_y, \Delta_y, \beta; \mathbf{X}_1, \dots, \mathbf{X}_h) = \sum_{y=1}^h \ell(\mu_y, \Delta_y, \beta; \mathbf{X}_y)$ over the nuisance parameters μ_y and Δ_y , $y = 1, \dots, h$. By the definition of minimum discrepancy function (Shapiro 1986), the likelihood-based objective function $L_d(\beta; h)$ can be reparameterized into a minimum discrepancy function (Cook and Forzani 2009; Section A.5). Analogous to FIRE and DIRE, we consider fusing a set of slicing schemes $H = \{h_1, \dots, h_K\}$ via the full likelihood of $\mathbf{X}_{yj}$, $y = 1, \dots, h_j$ and $j = 1, \dots, K$ and the composite likelihoods.

From the discussion in Section 3.2, it is not surprising to find that the full likelihood approach leads to the same objective function as LAD with the fused slicing scheme. However, unlike FIRE or IRE that uses $\hat{\zeta}_F$ in its objective function, LAD objective function (6.1) depends largely on sample covariance matrices $\tilde{\Delta}_y$ of each slice, which involves many more parameters. And ensuring well-conditioned $\tilde{\Delta}_y$'s requires a reasonable number of observations per slice. Hence the full likelihood approach of fusing LAD can be impractical.

Fortunately, we can adapt the idea of DIRE in the likelihood-based SDR context. By treating the slicing schemes independently, we write the fused objective function as $L_d^{\text{FLAD}}(\beta; H) \equiv \sum_{j=1}^K L_d(\beta; h_j)$, where $L_d(\beta; h_j)$ is the LAD objective function defined in (6.1). This is exactly the composite likelihood approach proposed by Lindsay (1988). See Varin et al. (2011) for an overview of recent development in the theory and application of composite likelihoods. More specifically, our fusing method corresponds to the block composite log-likelihood, where the whole set of variables $\{\mathbf{X}_1^{(1)}, \dots, \mathbf{X}_{h_1}^{(1)}, \dots, \mathbf{X}_{h_K}^{(K)}\}$ was divided into K subsets and the block composite log-likelihood was the sum of the log-likelihoods for every subset of variables.

The $\sqrt{n}$ -consistency of the block-composite fused LAD then follows from the theory of composite likelihoods. The simulation results for fused LAD are very encouraging, as

illustrated in the Supplementary materials.

7 Fusing to estimate the central solution space

Li and Dong (2009) introduced the notion of the central solution subspace (CSS) and developed methods based on CSS that relax the linearity condition and allow non-linear conditional means $E(\mathbf{X}|\beta^T \mathbf{X})$.

We adapted the idea of DIRE to one of the CSS methods, CSS-SIR, which is based on a quadratic objective function that depends on slicing (see Li and Dong, 2009, for details). For a fixed number of slices h , we write the objective function of CSS-SIR as $L_h(\boldsymbol{\eta})$, where $\boldsymbol{\eta}$ represents a basis for the CSS. We define the fused objective function over a slicing set $H = \{h_1, \dots, h_K\}$ as $L^F(\boldsymbol{\eta}) = \sum_{j=1}^K L_{h_j}(\boldsymbol{\eta})$. The fused estimator $\hat{\boldsymbol{\eta}}^F$ is obtained by minimizing this objective function. The $\sqrt{n}$ -consistency of $\hat{\boldsymbol{\eta}}^F$ follows from Li and Dong (2009; Corollary 6.1).

In order to illustrate the performance of this fused estimator, we conducted simulation studies based on the three simulation models used by Li and Dong (2009). The results given in the Supplement show that CSS-SIR can be quite sensitive to the number of slices. Moreover, the fused estimators accumulate information rapidly and have little variability for all $10 \leq h_K \leq 20$. Fusing provided significant improvement over the original CSS-SIR method in all three models. The gain from the fused estimator is the greatest in the most challenging model among the three models, suggesting that fusing is particularly needed when finding the central solution subspace is not easy.

8 Discussion

We proposed a practically useful solution to the important, long-standing problem of choosing a slicing scheme in sufficient dimension reduction methodology. FIRE is the asymptotically optimal methods but, owing to the cost of estimating the inner product

matrix, a large sample size may be needed to manifest this property. DIRE may generally be the method of choice. Although it is not asymptotically optimal, its performance was better than that of IRE unless the sample size was quite large and it overshadowed all other competitors.

An alternative approach is to directly obtain a regularized estimate of the inner product matrix Γ_F^{-1} . For instance, the sparse permutation invariant covariance estimator (Rothman et al, 2008) might offer improvement over DIRE, which uses only the block diagonal parts of $\hat{\Gamma}_F$.

Follow Cook and Ni (2005) and Ni and Cook (2007), we could build the theoretical foundation of tests for predictor effects and the CS dimensionality. Let $\hat{\mathcal{F}}_d(\mathbf{B}, \mathbf{C}; \Psi) = \mathcal{F}_d(\hat{\mathbf{B}}, \hat{\mathbf{C}}; \Psi)$ denote the minimized objective function value for the fused estimators with inner product matrix Ψ . Then under the same condition as Theorem 2, $n\hat{\mathcal{F}}_d(\mathbf{B}, \mathbf{C}; \Gamma_F^{-1})$ and $n\hat{\mathcal{F}}_d(\mathbf{B}, \mathbf{C}; \mathbf{G}_F^{-1})$ both follow an asymptotic χ^2 -squared distribution with degree of freedom $(p - d)(\sum_{j=1}^K (h_j - 1) - d)$. Then the asymptotic χ^2 -tests follows straightforwardly.

Recently, Ma and Zhu (2012) developed semi-parametric SDR methods that do not require the linearity condition. It is also possible to develop fused semi-parametric methods to relax linearity condition. The methods proposed by Ma and Zhu (2012) require nonparametric estimation of conditional moments, where the kernel bandwidth selection problem can be alleviated by fused estimator.

Acknowledgments

We are grateful to Yuexiao Dong and Bing Li for providing us the codes for computing the CSS estimators, and grateful to Liping Zhu, Lixing Zhu and Zhenghui Feng for the cumulative slicing estimation codes. The authors also like to thank the Editor, the Associate Editor and two referees for those helpful comments that led to significant improvements in this article. Research for this article was supported in part by grant DMS-1007547

from the U.S. National Science Foundation.

References

- [1] BATES, J.M. AND GRANGER, C.W.J. (1969), The combination of forecasts. *Operations Research Quarterly*, **20**, 451–468.
- [2] COOK, R.D. (1998), Regression Graphics: Ideas for Studying Regressions Through Graphics. New York: Wiley
- [3] COOK, R.D. AND FORZANI, L. (2008), Principal fitted components for dimension reduction in regression. *Statistical Science*. **23**, 485–501.
- [4] COOK, R.D. AND FORZANI, L. (2009), Likelihood-based sufficient dimension reduction. *Journal of the American Statistical Association*. **104**, 197–208.
- [5] COOK, R.D. AND NI, L. (2005), Sufficient dimension reduction via inverse regression: a minimum discrepancy approach, *Journal of the American Statistical Association*, **100**, 410–428.
- [6] COOK, R.D. AND WEISBERG, S. (1991), Discussion of "Sliced inverse regression for dimension reduction" by K.-C. Li. *Journal of the American Statistical Association*. **86**, 328–332.
- [7] COOK, R.D. AND YIN, X. (2001), Dimension reduction and visualization in discriminant analysis (with discussion). *Australia/New Zealand Journal of Statistics*. **43**, 147–199.
- [8] DONG, Y. AND LI, B. (2010), Dimension reduction for nonelliptically distributed predictors: second order methods, *Biometrika*, **97**, 279–294.
- [9] HALL, P. AND LI, K.C. (1993), On almost linearity of low dimensional projection from high dimensional data. *The Annals of Statistics*. **21**, 867–889.
- [10] HOOPER, J. (1959), Simultaneous equations and canonical correlation theory. *Econometrica*. **27**, 245–256.
- [11] HOTELLING, H. (1936), Relations between two sets of variates. *Biometrika*. **28**, 321–377.
- [12] HSING, T. AND CARROLL, R.J. (1992), An asymptotic theory for sliced inverse regression. *The Annals of Statistics*. **20**, 1040–1061.
- [13] LI, B. AND WANG, S. (2007), On directional regression for dimension reduction. *Journal of the American Statistical Association*. **102**, 997–1008.
- [14] LI, B. AND DONG, Y. (2009), Dimension reduction for nonelliptically distributed predictors, *The Annals of Statistics*. **37**, 1272–1298.

- [15] LI, K.C. (1991), Sliced inverse regression for dimension reduction. *Journal of the American Statistical Association*. **86**, 316–342.
- [16] LINDSAY, B. (1988), Composite likelihood methods. *Contemporary Mathematics*. **80**, 220–239.
- [17] LUCENO, A. (1999), Discrete approximations to continuous univariate distributions—and alternative to simulation. *Journal of the Royal Statistical Society B*. **61**, 345–352.
- [18] MA, Y. AND ZHU, L. (2012), A semiparametric approach to dimension reduction. *Journal of the American Statistical Association*. **107**, 168–3179.
- [19] NI, L. AND COOK, R.D. (2007), A robust inverse regression estimator *Statistics and Probability Letters*, **77**, 343–349.
- [20] ROTHMAN, A.J., BICKEL, P.J., LEVINA, E., AND ZHU, J. (2008), Sparse Permutation Invariant Covariance Estimation, *Electronic Journal of Statistics*, **2**, 494–515.
- [21] SETODJI, C. AND COOK, R. D. (2004), K-means inverse regression. *Technometrics*, **46**, 421–429.
- [22] SHAPIRO, A. (1986), Asymptotic theory of overparameterized structural models, *Journal of the American Statistical Association*, **81**, 142–149.
- [23] TIMMERMAN, A.G. (2006), Forecast Combinations. *Handbook of Economic Forecasting*, Elsevier.
- [24] VARIN, C., REID, N. AND FIRTH, D. (2011), An overview of composite likelihood methods. *Statistica Sinica*. **21**, 5–42.
- [25] YE, Z. AND WEISS, R.E. (2003), Using the Bootstrap to Select One of a New Class of Dimension Reduction Methods. *Journal of the American Statistical Association*. **98**, 968–979.
- [26] ZHU, L. AND NG, K.W. (1995), Asymptotics of sliced inverse regression. *Statistica Sinica*, **5**, 727–736.
- [27] ZHU, L., ZHU, L. AND FENG, Z. (2010), Dimension Reduction in Regressions Through Cumulative Slicing Estimation. *Journal of the American Statistical Association*. **105**, 1455–1466.

Supplement: Proofs and Additional Simulations for “Fused estimators of the central subspace in sufficient dimension reduction”

R. Dennis Cook and Xin Zhang

A Proofs

A.1 Theorem 1

We prove the more general case where $K \geq 2$ for Theorem 1 as follows. Follow the proof of Cook and Ni (2005; Theorem 1), let $\mathbf{Z}_i$ denote the i^{th} observed standardized predictor and $\boldsymbol{\epsilon}_i^{(j)} = (\epsilon_{1i}^{(j)}, \dots, \epsilon_{h_j i}^{(j)})^T$ be the i^{th} value for the random vector $\boldsymbol{\epsilon}^{(j)}$, as defined in Section 3. Then we have

$$\sqrt{n}(\text{vec}(\hat{\boldsymbol{\xi}}^{(j)} \hat{\mathbf{D}}_j) - \text{vec}(\boldsymbol{\beta} \boldsymbol{\gamma}_j \mathbf{D}_j)) = n^{-1/2} \sum_{i=1}^n \text{vec}(\boldsymbol{\Sigma}^{-1/2} \mathbf{Z}_i \boldsymbol{\epsilon}_i^{(j)T}) + O_p(n^{-1/2}),$$

where $(\mathbf{Z}_i, \boldsymbol{\epsilon}_i^{(j)})$ are iid random vectors. Thus, for the combined case, by writing

$$\text{vec}(\boldsymbol{\xi}^{(1)} \mathbf{D}_1, \dots, \boldsymbol{\xi}^{(K)} \mathbf{D}_K) = (\text{vec}^T(\boldsymbol{\xi}^{(1)} \mathbf{D}_1), \dots, \text{vec}^T(\boldsymbol{\xi}^{(K)} \mathbf{D}_K))^T,$$

we have the following results,

$$\begin{aligned} & \sqrt{n}(\text{vec}(\hat{\boldsymbol{\xi}}^{(1)} \hat{\mathbf{D}}_1, \dots, \hat{\boldsymbol{\xi}}^{(K)} \hat{\mathbf{D}}_K) - \text{vec}(\boldsymbol{\beta} \boldsymbol{\gamma}_1 \mathbf{D}_1, \dots, \boldsymbol{\beta} \boldsymbol{\gamma}_K \mathbf{D}_K)) \\ &= n^{-1/2} \sum_{i=1}^n \text{vec}(\boldsymbol{\Sigma}^{-1/2} \mathbf{Z}_i (\boldsymbol{\epsilon}_i^{(1)T}, \dots, \boldsymbol{\epsilon}_i^{(K)T})) + O_p(n^{-1/2}) \\ &= n^{-1/2} \sum_{i=1}^n \text{vec}(\boldsymbol{\Sigma}^{-1/2} \mathbf{Z}_i \mathbf{E}_i^T) + O_p(n^{-1/2}), \end{aligned}$$

Because $(\mathbf{Z}_i, \mathbf{E}_i)$ are iid, we have the asymptotic distribution

$$\sqrt{n}(\text{vec}(\hat{\boldsymbol{\xi}}^{(1)} \hat{\mathbf{D}}_1, \dots, \hat{\boldsymbol{\xi}}^{(K)} \hat{\mathbf{D}}_K) - \text{vec}(\boldsymbol{\xi}^{(1)} \mathbf{D}_1, \dots, \boldsymbol{\xi}^{(K)} \mathbf{D}_K)) \sim N(0, \boldsymbol{\Gamma}_H),$$

and $\mathbf{\Gamma}_H = \text{cov}(\text{vec}(\mathbf{\Sigma}^{-1/2} \mathbf{Z} \mathbf{E}^T))$. The resulting asymptotic distribution of $\text{vec}(\hat{\boldsymbol{\zeta}}_C)$ could be seen from the relation below

$$\begin{aligned} \text{vec}(\hat{\boldsymbol{\zeta}}_C) &= \text{vec}(\hat{\boldsymbol{\xi}}^{(1)} \mathbf{D}_1 \mathbf{A}_1, \dots, \hat{\boldsymbol{\xi}}^{(K)} \mathbf{D}_K \mathbf{A}_K) \\ &= \text{vec}((\hat{\boldsymbol{\xi}}^{(1)} \mathbf{D}_1, \dots, \hat{\boldsymbol{\xi}}^{(K)} \mathbf{D}_K) \cdot \text{diag}(\mathbf{A}_1, \dots, \mathbf{A}_K)) \\ &= \text{vec}((\hat{\boldsymbol{\xi}}^{(1)} \mathbf{D}_1, \dots, \hat{\boldsymbol{\xi}}^{(K)} \mathbf{D}_K) \mathbf{A}_H) \\ &= \mathbf{A}_H \otimes \mathbf{I}_p \cdot \text{vec}(\hat{\boldsymbol{\xi}}^{(1)} \mathbf{D}_1, \dots, \hat{\boldsymbol{\xi}}^{(K)} \mathbf{D}_K). \end{aligned}$$

A.2 Theorem 2

Following the proof of Cook and Ni (2005; Theorem 2), we define the Jacobian matrices as follows.

$$\begin{aligned} \Delta_j &= (\boldsymbol{\nu}_j^T \otimes \mathbf{I}_p, \mathbf{I}_{(h_j-1)} \otimes \boldsymbol{\beta}_j), \quad j = 1, 2, \\ \Delta_F &= (\boldsymbol{\omega}_F^T \otimes \mathbf{I}_p, \mathbf{I}_{\sum_j (h_j-1)} \otimes \boldsymbol{\beta}_F) = \begin{pmatrix} \tilde{\Delta}_1 \\ \tilde{\Delta}_2 \end{pmatrix} \\ &= \begin{pmatrix} \boldsymbol{\omega}_1^T \otimes \mathbf{I}_p & \mathbf{I}_{h_1-1} \otimes \boldsymbol{\beta}_F & 0 \\ \boldsymbol{\omega}_2^T \otimes \mathbf{I}_p & 0 & \mathbf{I}_{h_2-1} \otimes \boldsymbol{\beta}_F \end{pmatrix}. \end{aligned}$$

By applying Proposition 4.1 in Shapiro (1986), we get the $\sqrt{n}$ -consistency of the estimators and their asymptotic covariance matrices $\mathbf{V}_j = \Delta_j (\Delta_j^T \mathbf{\Gamma}_j^{-1} \Delta_j)^{-1} \Delta_j^T$, $j = 1, 2$, and

$$\mathbf{W} = \Delta_F (\Delta_F^T \boldsymbol{\Psi} \Delta_F)^{-1} \Delta_F^T \boldsymbol{\Psi} \mathbf{\Gamma}_F \boldsymbol{\Psi} \Delta_F (\Delta_F^T \boldsymbol{\Psi} \Delta_F)^{-1} \Delta_F^T = \begin{pmatrix} \mathbf{W}_1 & \mathbf{W}_{12} \\ \mathbf{W}_{12}^T & \mathbf{W}_2 \end{pmatrix},$$

where $\mathbf{W}_j$'s are the block diagonal element of $\mathbf{W}$. We next want to show that $\mathbf{W}_j \leq \mathbf{V}_j$, $j = 1, 2$, when $\boldsymbol{\Psi} = \mathbf{\Gamma}_F^{-1}$. Because of the interchangeability of two slicing schemes, we will only need to show that $\mathbf{W}_1 \leq \mathbf{V}_1$. By direct calculation,

$$\mathbf{W}_1 = \tilde{\Delta}_1 (\Delta_F^T \mathbf{\Gamma}_F^{-1} \Delta_F)^{-1} \tilde{\Delta}_1^T, \quad (\text{A.1})$$

where the Jacobian matrices $\tilde{\Delta}_1$ have the same row dimension as Δ_1 , but different column dimensions. Now we transform the asymptotic covariance matrix $\mathbf{\Gamma}_F = \begin{pmatrix} \mathbf{\Gamma}_1 & \mathbf{\Gamma}_{12} \\ \mathbf{\Gamma}_{12}^T & \mathbf{\Gamma}_2 \end{pmatrix}$

into block diagonal matrix by a non-singular transformation matrix $\mathbf{H}$,

$$\mathbf{H} = \begin{pmatrix} \mathbf{I}_{p(h_1-1)} & \mathbf{0} \\ -\mathbf{\Gamma}_{12}^T \mathbf{\Gamma}_1^{-1} & \mathbf{I}_{p(h_2-1)} \end{pmatrix},$$

so that $\mathbf{H}\mathbf{\Gamma}_F\mathbf{H}^T = \begin{pmatrix} \mathbf{\Gamma}_1 & 0 \\ 0 & \mathbf{\Gamma}_2 - \mathbf{\Gamma}_{12}^T \mathbf{\Gamma}_1^{-1} \mathbf{\Gamma}_{12} \end{pmatrix}$ substitute into (A.1), we get

$$\mathbf{W}_1 = \tilde{\mathbf{\Delta}}_1 (\mathbf{\Delta}_F^T \mathbf{\Gamma}_F^{-1} \mathbf{\Delta}_F)^{-1} \tilde{\mathbf{\Delta}}_1^T = \tilde{\mathbf{\Delta}}_1 (\mathbf{\Delta}_F^T \mathbf{H}^T (\mathbf{H}\mathbf{\Gamma}_F\mathbf{H}^T)^{-1} \mathbf{H}\mathbf{\Delta}_F)^{-1} \tilde{\mathbf{\Delta}}_1^T,$$

$$\mathbf{H}\mathbf{\Delta}_F = \begin{pmatrix} \mathbf{I}_{p(h_1-1)} & \mathbf{0} \\ -\mathbf{\Gamma}_{12}^T \mathbf{\Gamma}_1^{-1} & \mathbf{I}_{p(h_2-1)} \end{pmatrix} \begin{pmatrix} \tilde{\mathbf{\Delta}}_1 \\ \tilde{\mathbf{\Delta}}_2 \end{pmatrix} \equiv \begin{pmatrix} \tilde{\mathbf{\Delta}}_1 \\ \mathbf{D}_2 \end{pmatrix}.$$

Then we have

$$\begin{aligned} \mathbf{W}_1 &= \tilde{\mathbf{\Delta}}_1 \left\{ \begin{pmatrix} \tilde{\mathbf{\Delta}}_1^T & \mathbf{D}_2^T \end{pmatrix} \begin{pmatrix} \mathbf{\Gamma}_1^{-1} & 0 \\ 0 & (\mathbf{\Gamma}_2 - \mathbf{\Gamma}_{12}^T \mathbf{\Gamma}_1^{-1} \mathbf{\Gamma}_{12})^{-1} \end{pmatrix} \begin{pmatrix} \tilde{\mathbf{\Delta}}_1 \\ \mathbf{D}_2 \end{pmatrix} \right\}^{-1} \tilde{\mathbf{\Delta}}_1^T \\ &= \tilde{\mathbf{\Delta}}_1 \left\{ \tilde{\mathbf{\Delta}}_1^T \mathbf{\Gamma}_1^{-1} \tilde{\mathbf{\Delta}}_1 + \mathbf{D}_2^T (\mathbf{\Gamma}_2 - \mathbf{\Gamma}_{12}^T \mathbf{\Gamma}_1^{-1} \mathbf{\Gamma}_{12})^{-1} \mathbf{D}_2 \right\}^{-1} \tilde{\mathbf{\Delta}}_1^T \\ &\leq \tilde{\mathbf{\Delta}}_1 \left\{ \tilde{\mathbf{\Delta}}_1^T \mathbf{\Gamma}_1^{-1} \tilde{\mathbf{\Delta}}_1 \right\}^{-1} \tilde{\mathbf{\Delta}}_1^T. \end{aligned} \tag{A.2}$$

Finally, because of $\beta_F \omega_1 = \beta_1 \nu_1 = \zeta^{(1)}$, we can find a non-singular transformation $\mathbf{A}$ such that

$$\tilde{\mathbf{\Delta}}_1 = \begin{pmatrix} \omega_1^T \otimes \mathbf{I}_p & \mathbf{I}_{h_1-1} \otimes \beta_F & 0 \end{pmatrix} = \begin{pmatrix} \nu_1^T \otimes \mathbf{I}_p & \mathbf{I}_{h_1-1} \otimes \beta_1 & 0 \end{pmatrix} \mathbf{A} = \begin{pmatrix} \mathbf{\Delta}_1 & 0 \end{pmatrix} \mathbf{A}.$$

Then (A.2) can be written as

$$\begin{aligned} \mathbf{W}_1 &\leq \tilde{\mathbf{\Delta}}_1 \left\{ \tilde{\mathbf{\Delta}}_1^T \mathbf{\Gamma}_1^{-1} \tilde{\mathbf{\Delta}}_1 \right\}^{-1} \tilde{\mathbf{\Delta}}_1^T \\ &= \begin{pmatrix} \mathbf{\Delta}_1 & 0 \end{pmatrix} \left\{ \begin{pmatrix} \mathbf{\Delta}_1^T \\ 0 \end{pmatrix} \mathbf{\Gamma}_1^{-1} \begin{pmatrix} \mathbf{\Delta}_1 & 0 \end{pmatrix} \right\}^{-1} \begin{pmatrix} \mathbf{\Delta}_1^T \\ 0 \end{pmatrix} \\ &= \begin{pmatrix} \mathbf{\Delta}_1 & 0 \end{pmatrix} \begin{pmatrix} \mathbf{\Delta}_1^T \mathbf{\Gamma}_1^{-1} \mathbf{\Delta}_1 & 0 \\ 0 & 0 \end{pmatrix}^{-1} \begin{pmatrix} \mathbf{\Delta}_1^T \\ 0 \end{pmatrix} \\ &= \mathbf{\Delta}_1 (\mathbf{\Delta}_1^T \mathbf{\Gamma}_1^{-1} \mathbf{\Delta}_1)^{-1} \mathbf{\Delta}_1^T \\ &= \mathbf{V}_1. \end{aligned}$$

B Additional Simulations

B.1 Simulations for fused LAD

We simulated 100 datasets for each of the following settings with $n = 300$ sample each and $p = 15$.

Model 1. We simulated Y from $\text{uniform}(0, 5)$ distribution and

$$\mathbf{X}|Y = 0.5\mathbf{\Gamma}\boldsymbol{\eta}\mathbf{f}(Y) + \boldsymbol{\epsilon} + \sigma(Y)\mathbf{\Gamma}\boldsymbol{\epsilon},$$

where $(\boldsymbol{\epsilon}, \boldsymbol{\epsilon})^T \in \mathbb{R}^{p+2}$ is a multivariate normal random vector with mean 0, identity covariance, and is independent of Y . The vector valued function $\mathbf{f}(Y) \in \mathbb{R}^{r \times 1}$ was designed to consist of indicator functions of Y :

$$\mathbf{f}(Y) = (J_1, \dots, J_5)^T, \quad J_k = I(k-1 < Y \leq k), \quad k = 1, \dots, 5.$$

We used $\sigma(Y) \in \mathbb{R}^1$ to construct heteroscedastic error structure. The elements of both constant matrices $\mathbf{\Gamma} \in \mathbb{R}^{p \times d}$ and $\boldsymbol{\eta} \in \mathbb{R}^{d \times r}$ were generated as a random sample from the standard normal distribution and then standardized to be semi-orthogonal. The central subspace for this model is $\mathcal{S}_{Y|\mathbf{X}} = \text{span}(\mathbf{\Gamma})$.

Two choices of $\sigma(Y)$ were used in this setting: (a) $\sigma(Y) = Y$; and (b) $\sigma(Y) = (1, 2, 3, 4, 5) \times \mathbf{f}(Y)$. In (a), we know the optimal slicing scheme for estimating the mean function is 5-slices quantile slicing, but it may not be the best for the conditional variance part. However, in (b), 5-slices quantile slicing scheme should be the best for both the mean and variance estimation.

Model 2. We simulated Y , $\mathbf{f}(Y)$, $\mathbf{\Gamma}$ and $\boldsymbol{\eta}$ same as Mode 1, and

$$\mathbf{X}|Y = \mathbf{\Gamma}_1\boldsymbol{\eta}\mathbf{f}(Y) + \boldsymbol{\epsilon} + \sigma(Y)\mathbf{\Gamma}_2\boldsymbol{\epsilon},$$

where $\mathbf{\Gamma}_1$ and $\mathbf{\Gamma}_2$ were the two directions of $\mathbf{\Gamma} = (\mathbf{\Gamma}_1, \mathbf{\Gamma}_2)$, and $(\boldsymbol{\epsilon}, \boldsymbol{\epsilon})^T \in \mathbb{R}^{p+1}$ was multivariate normal with mean 0, identity covariance and was independent of Y . The

central subspace for this model was $\mathcal{S}_{Y|\mathbf{X}} = \text{span}(\mathbf{\Gamma})$ where one direction was in the mean and another direction in the variance.

Two choices of $\sigma(Y)$ were used in this setting: (a) $\sigma(Y) = 1, 2$ or 5 if $Y < 5/3$, $5/3 < Y < 10/3$ or $Y > 10/3$, respectively; and (b) $\sigma(Y) = (1, 2, 3, 4, 5) \times \mathbf{f}(Y)$. In (a), there was no quantile slicing scheme in $H = \{3, \dots, 15\}$ such that both the means and the covariances were constant within each slice. This means that the LAD model was misspecified for any slicing schemes. In (b), 5-slices quantile slicing scheme should be the best for both the mean and variance estimation.

In Figure B.1, we plotted the averaged angles between the CS and the estimated subspace. Since the subspaces were two-dimensional, we only meant to consider the largest angles between subspaces. The results shown in Figure B.1 were very encouraging. For all these settings, the fused estimators were robust and had little changes even the largest number of slices h_K varied from 6 to 15. The standard errors for each points in the plots were all within 0.5 to 1.5, implying that the fused estimators with $h_K = 15$ had never been significantly worse than the best single slicing schemes. Also, the improvement of fusing with $h_K = 15$ over single slicing scheme with $h > 9$ were very significant in all these four settings.

B.2 Simulations for fused CSS-SIR

In order to illustrate the performance of this fused estimator, we conducted simulation studies for $n = 200$, $p = 6$ based on the three simulation models used by Li and Dong (2009):

$$(I): \quad Y = e^{X_3} + (X_4 + 1.5)^2 + \epsilon,$$

$$(II): \quad Y = 0.4X_3^2 + 3\sin(X_4/4) + 0.5\epsilon,$$

$$(III): \quad Y = X_3/[0.5 + (1.5 + X_4)^2] + 0.1\epsilon,$$

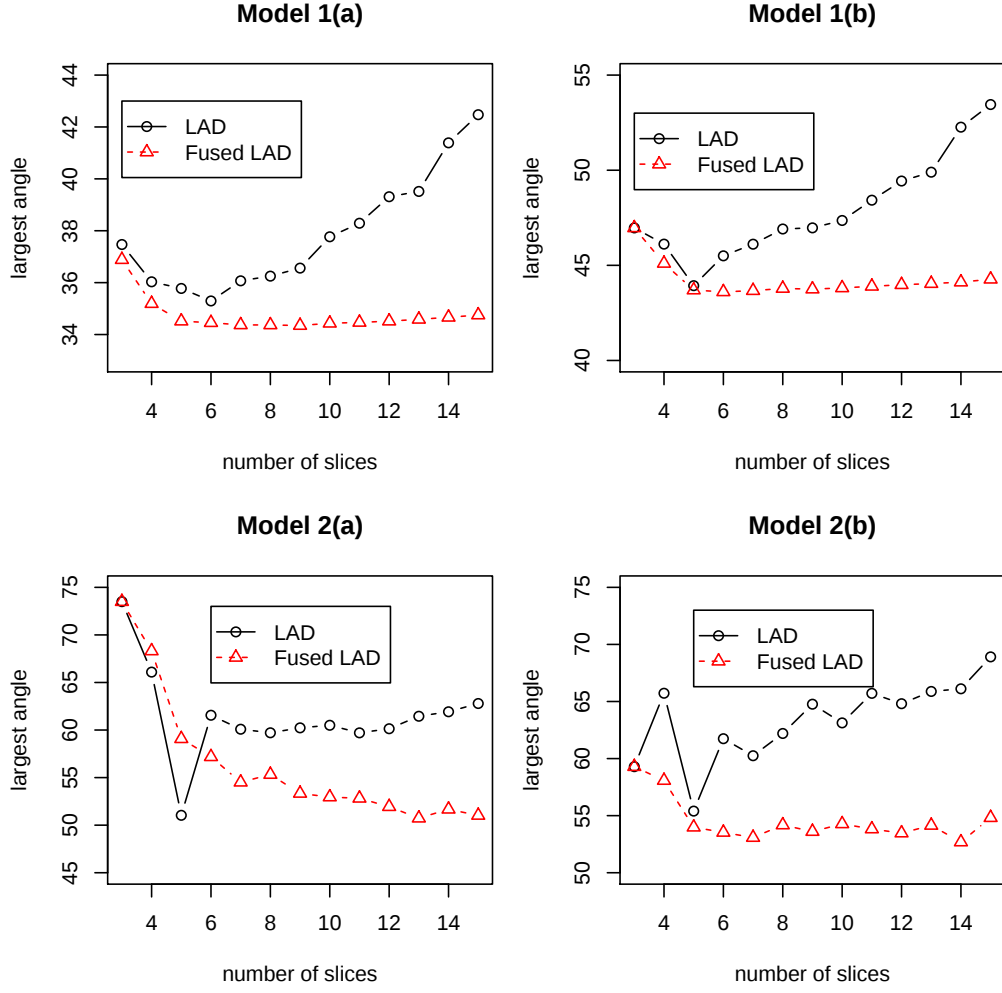


Figure B.1: Averaged angles between the CS and the estimated subspaces. The horizontal axis indicates the number of slices h for LAD or the largest number of slices h_K in $H = \{3, \dots, h_K\}$ for fused LAD.

where ϵ , X_1 , X_2 , X_5 , X_6 are iid $N(0, 1)$ random variables. To introduce nonlinearity in the predictors, Li and Dong generated X_3 and X_4 as follows

$$\begin{aligned} X_3 &= 0.2X_1 + 0.2(X_2 + 2)^2 + 0.2\delta, \\ X_4 &= 0.1 + 0.1(X_1 + X_2) + 0.3(X_1 + 1.5)^2 + 0.2\delta, \end{aligned}$$

where δ is $N(0, 1)$ random variable. We simulated 100 datasets for each model and computed the squared multiple correlation coefficient ρ^2 between estimated predictors $\hat{\beta}^T \mathbf{X}$ and the truth $\beta^T \mathbf{X} = (X_3, X_4)^T$ (Hall and Mathiason 1990; Li and Dong 2009). The squared multiple correlation coefficient takes its maximum value 2 if $\hat{\beta}^T \mathbf{X}$ and $\beta^T \mathbf{X}$ have a perfect linear relationship, takes minimum value 0 if $\hat{\beta}^T \mathbf{X}$ and $\beta^T \mathbf{X}$ are uncorrelated and takes the value 1 if one direction in $\hat{\beta}^T \mathbf{X} = (\hat{\beta}_1^T \mathbf{X}, \hat{\beta}_2^T \mathbf{X})$ has a perfect linear relationship with $\beta^T \mathbf{X}$ and the other direction in $\hat{\beta}^T \mathbf{X}$ is uncorrelated with $\beta^T \mathbf{X}$.

In Figure B.2 we plotted the average of $2 - \rho^2$, which measures the distance between the estimated predictors and the truth, against the number of slices h used by SIR and CSS-SIR, or against the largest slice number h_K in the fused CSS-SIR with $H = \{3, \dots, h_K\}$. Li and Dong used $h = 10$ in their simulations. The horizontal dashed lines are the averaged distance $2 - \rho^2$ for the fused estimators with $H = \{3, \dots, 15\}$. The results show that CSS-SIR can be quite sensitive to the number of slices. Moreover, the fused estimators accumulate information rapidly and have little variability for all $10 \leq h_K \leq 20$. Fusing provided significant improvement over the original CSS-SIR method in all three models. Model I always had the largest distances among the three models and thus is the most challenging model. The gain from the fused estimator is the greatest in model I, suggesting that fusing is particularly needed when finding the central solution subspace is not easy.

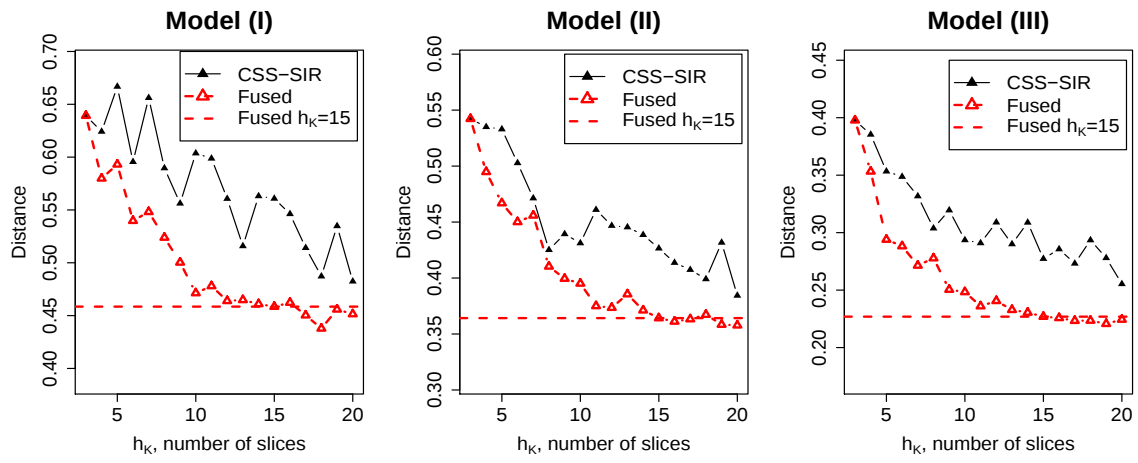


Figure B.2: Estimated distances $2 - \rho^2$ for each model. Standard errors for all estimated distances are less than 0.02.