



Estimating the Structural Dimension of Regressions via Parametric Inverse Regression

Efstathia Bura; R. Dennis Cook

Journal of the Royal Statistical Society. Series B (Statistical Methodology), Vol. 63, No. 2. (2001), pp. 393-410.

Stable URL:

<http://links.jstor.org/sici?sici=1369-7412%282001%2963%3A2%3C393%3AETSDOR%3E2.0.CO%3B2-T>

Journal of the Royal Statistical Society. Series B (Statistical Methodology) is currently published by Royal Statistical Society.

Your use of the JSTOR archive indicates your acceptance of JSTOR's Terms and Conditions of Use, available at <http://www.jstor.org/about/terms.html>. JSTOR's Terms and Conditions of Use provides, in part, that unless you have obtained prior permission, you may not download an entire issue of a journal or multiple copies of articles, and you may use content in the JSTOR archive only for your personal, non-commercial use.

Please contact the publisher regarding any further use of this work. Publisher contact information may be obtained at <http://www.jstor.org/journals/rss.html>.

Each copy of any part of a JSTOR transmission must contain the same copyright notice that appears on the screen or printed page of such transmission.

The JSTOR Archive is a trusted digital repository providing for long-term preservation and access to leading academic journals and scholarly literature from around the world. The Archive is supported by libraries, scholarly societies, publishers, and foundations. It is an initiative of JSTOR, a not-for-profit organization with a mission to help the scholarly community take advantage of advances in technology. For more information regarding JSTOR, please contact support@jstor.org.

Estimating the structural dimension of regressions via parametric inverse regression

Efstathia Bura

George Washington University, Washington DC, USA

and R. Dennis Cook

University of Minnesota, St Paul, USA

[Received November 1999. Final revision November 2000]

Summary. A new estimation method for the dimension of a regression at the outset of an analysis is proposed. A linear subspace spanned by projections of the regressor vector $\mathbf{X}$, which contains part or all of the modelling information for the regression of a vector $\mathbf{Y}$ on $\mathbf{X}$, and its dimension are estimated via the means of parametric inverse regression. Smooth parametric curves are fitted to the p inverse regressions via a multivariate linear model. No restrictions are placed on the distribution of the regressors. The estimate of the dimension of the regression is based on optimal estimation procedures. A simulation study shows the method to be more powerful than sliced inverse regression in some situations.

Keywords: Asymptotic test for dimension; Dimension reduction; Inverse regression; Parametric inverse regression; Sliced inverse regression

1. Introduction

Let $\mathbf{Y} \in \mathbb{R}^m$ and $\mathbf{X} \in \mathbb{R}^p$ with joint cumulative distribution function (CDF) $F(\mathbf{Y}, \mathbf{X})$. Regression analyses typically tend to concentrate on the study of the first two moments of the conditional CDF of $\mathbf{Y}$ given $\mathbf{X}$, $F(\mathbf{Y}|\mathbf{X})$. In general, though, the goal of regression is the study of the behaviour of $F(\mathbf{Y}|\mathbf{X})$, as the value of $\mathbf{X}$ varies in its marginal sample space. As a means of characterizing the regression structure, consider replacing $\mathbf{X}$ by $k \leq p$ linear combinations of its components, $\eta_1^T \mathbf{X}, \dots, \eta_k^T \mathbf{X}$, without losing information on $F(\mathbf{Y}|\mathbf{X})$ so that, for all values of $\mathbf{X}$,

$$\mathbf{Y} \perp\!\!\!\perp \mathbf{X} | \boldsymbol{\eta}^T \mathbf{X} \quad (1)$$

where $\boldsymbol{\eta}$ is the $p \times k$ matrix with columns η_j . The notation $U \perp\!\!\!\perp V | W$ in expression (1) means that U is independent of V given any value for W (Dawid, 1979). Expression (1), as a mathematical formulation of the dependence of $\mathbf{Y}$ on $\mathbf{X}$, was introduced by Cook (1994a). It expresses the fact that the conditional CDF of $\mathbf{Y}|\mathbf{X}$ depends on $\mathbf{X}$ only through $\boldsymbol{\eta}^T \mathbf{X}$, the coordinates of a projection of $\mathbf{X}$ onto the k -dimensional linear subspace spanned by the columns of $\boldsymbol{\eta}$. Consequently, $\boldsymbol{\eta}^T \mathbf{X}$ contains equivalent or *sufficient*, in the statistical sense, information for the regression of $\mathbf{Y}$ on $\mathbf{X}$. Most importantly, if $k < p$, then a sufficient reduction in the dimension of the regression is achieved, which in turn leads to *sufficient*

Address for correspondence: Efstathia Bura, Department of Statistics, George Washington University, 2201 G Street NW, Washington, DC 20052, USA.
E-mail: ebura@gwu.edu

summary plots of $\mathbf{Y}$ versus $\boldsymbol{\eta}^T \mathbf{X}$ as graphical displays of all the necessary modelling information for the regression of $\mathbf{Y}$ on $\mathbf{X}$. Subsequently, sufficient summary plots can guide the selection of appropriate models for $F(\mathbf{Y}|\mathbf{X})$.

For any vector or matrix $\boldsymbol{\alpha}$, let $S(\boldsymbol{\alpha})$ denote its range space and $\dim\{S(\boldsymbol{\alpha})\}$ denote its dimension. If expression (1) holds then it also holds with $\boldsymbol{\eta}$ replaced by any basis for $S(\boldsymbol{\eta})$. In this sense, expression (1) can be regarded as a statement about $S(\boldsymbol{\eta})$ rather than a statement about $\boldsymbol{\eta}$, *per se*. Thus, when expression (1) holds we follow Li (1991, 1992) and call $S(\boldsymbol{\eta})$ a *dimension reduction subspace* for $F(\mathbf{Y}|\mathbf{X})$, or for the regression of $\mathbf{Y}$ on $\mathbf{X}$.

Obviously, the smallest dimension reduction subspace provides the greatest dimension reduction in the predictor vector. There are several ways to define such a subspace. In this paper we use the *central dimension reduction subspace*, denoted $S_{\mathbf{Y}|\mathbf{X}}$ (Cook, 1994b, 1996, 1998a, b). $S_{\mathbf{Y}|\mathbf{X}}$ is the intersection of all dimension reduction subspaces for $F(\mathbf{Y}|\mathbf{X})$ and is trivially a subspace but is not necessarily a dimension reduction subspace. The existence of central subspaces can be assured by placing fairly weak restrictions on aspects of the joint distribution of $\mathbf{Y}$ and $\mathbf{X}$ (Cook, 1994a, 1996). Throughout this paper, we focus on regressions for which central dimension reduction spaces exist.

The subspace $S_{\mathbf{Y}|\mathbf{X}} = S_{\mathbf{Y}|\mathbf{X}}(\boldsymbol{\eta})$ is in effect a ‘metaparameter’ that is used to index the conditional distribution of $\mathbf{Y}$ given $\mathbf{X}$. The columns of the $p \times k$ matrix $\boldsymbol{\eta}$ will denote a basis for the central subspace $S_{\mathbf{Y}|\mathbf{X}}$, and k will be used to denote its dimension or the *structural dimension of the regression of $\mathbf{Y}$ on $\mathbf{X}$* (Cook and Weisberg, 1994). Our main objective is the estimation of $S_{\mathbf{Y}|\mathbf{X}}$.

The paper is organized as follows: existing dimension estimation methods, with emphasis on sliced inverse regression (SIR) (Li, 1991), are reviewed in Section 2. The example in Section 2 illustrates both the application of SIR and its limitations. The estimation method proposed, namely parametric inverse regression (PIR), is introduced and described in Section 3. Section 4 contains its extension to the non-constant variance case. The algorithm describing the PIR dimension reduction procedure is presented in Section 5. In Section 6, PIR is applied to the example of Section 2 and a simulation study to compare the power of the two testing methods for dimension is carried out. A concluding discussion is presented in Section 7. The lengthier proofs are given in Appendix A.

2. Background: inverse regression and sliced inverse regression

Methods are available for estimating portions of the central subspace $S_{\mathbf{Y}|\mathbf{X}}$, provided that certain conditions are placed on the marginal distribution of the predictors.

Let $S_{E(\mathbf{X}|\mathbf{Y})}$ denote the subspace spanned by $\{E(\mathbf{X}|\mathbf{Y}) - E(\mathbf{X}): \mathbf{Y} \in \Omega_{\mathbf{Y}}\}$, where $\Omega_{\mathbf{Y}} \subset \mathbb{R}^m$ is the marginal sample space of $\mathbf{Y}$. Given expression (1), assume that the marginal distribution of the predictors $\mathbf{X}$ satisfy the following condition, which henceforth will be referred to as the *linearity condition*: for all $\mathbf{b} \in \mathbb{R}^p$, $E(\mathbf{b}^T \mathbf{X} | \boldsymbol{\eta}^T \mathbf{X})$ is linear in $\boldsymbol{\eta}^T \mathbf{X}$.

Under this linearity condition on the regressor distribution, Li (1991), theorem 3.1, showed that the centred inverse regression curve $E(\mathbf{X}|\mathbf{Y}) - E(\mathbf{X})$ satisfies

$$E(\mathbf{X}|\mathbf{Y}) - E(\mathbf{X}) \in S(\boldsymbol{\Sigma}_x \boldsymbol{\eta}).$$

Equivalently,

$$S_{E(\mathbf{X}|\mathbf{Y})} \subset S(\boldsymbol{\Sigma}_x \boldsymbol{\eta}) = \boldsymbol{\Sigma}_x S_{\mathbf{Y}|\mathbf{X}} \quad (2)$$

where $\boldsymbol{\Sigma}_x = \text{cov}(\mathbf{X})$.

The linearity condition on $E(\mathbf{b}^T \mathbf{X} | \boldsymbol{\eta}^T \mathbf{X})$ is required to hold only for the basis $\boldsymbol{\eta}$ of the central subspace. Since $\boldsymbol{\eta}$ is unknown, in practice we may require that it holds for all possible $\boldsymbol{\eta}$, which is equivalent to elliptical symmetry of the distribution of $\mathbf{X}$ (Eaton, 1986). Li (1991) mentioned that the linearity condition is not a severe restriction, since most low dimensional projections of a high dimensional data cloud are close to being normal (Diaconis and Freedman, 1984; Hall and Li, 1993). In addition, there often exist transformations of the predictors that make them comply with the linearity condition. Cook and Nachtsheim (1994) suggested reweighting of the predictor vector to make it elliptically contoured.

Suppose that $\Sigma_x > 0$ and let $\mathbf{Z}$ be the standardized version of $\mathbf{X}$,

$$\mathbf{Z} = \Sigma_x^{-1/2} \{\mathbf{X} - E(\mathbf{X})\}.$$

Obviously, $E(\mathbf{Z}) = 0$ and $\text{cov}(\mathbf{Z}) = \mathbf{I}_p$. The observable sample version $\hat{\mathbf{Z}}$ is constructed by replacing Σ_x and $E(\mathbf{X})$ with their usual moment estimates. Also, since $\mathbf{Z}$ is a 1-1 and onto linear transformation of $\mathbf{X}$, $\mathbf{Y} \perp\!\!\!\perp \mathbf{X} | \boldsymbol{\eta}^T \mathbf{X}$ if and only if $\mathbf{Y} \perp\!\!\!\perp \mathbf{Z} | \boldsymbol{\beta}^T \mathbf{Z}$, where $\boldsymbol{\beta} = \Sigma_x^{1/2} \boldsymbol{\eta}$ or $\boldsymbol{\beta}_i = \Sigma_x^{1/2} \boldsymbol{\eta}_i$, $i = 1, 2, \dots, k$. By condition (2), we obtain that

$$E(\mathbf{Z} | \mathbf{Y}) \in S(\Sigma_x^{1/2} \boldsymbol{\eta}) = S(\boldsymbol{\beta}) = S_{\mathbf{Y} | \mathbf{Z}}. \quad (3)$$

The containment relationship (3) readily implies that $E(\mathbf{Z} | \mathbf{Y}) = \mathbf{P}_\beta E(\mathbf{Z} | \mathbf{Y})$, where $\mathbf{P}_\beta$ is the orthogonal projection operator for $S(\boldsymbol{\beta})$ with respect to the usual inner product. It also implies that $S_{E(\mathbf{Z} | \mathbf{Y})}$ is a subspace of $S_{\mathbf{Y} | \mathbf{Z}}$. This does not guarantee equality between $S_{E(\mathbf{Z} | \mathbf{Y})}$ and $S_{\mathbf{Y} | \mathbf{Z}}$ and, thus, inference about $S_{E(\mathbf{Z} | \mathbf{Y})}$ possibly covers only part of $S_{\mathbf{Y} | \mathbf{Z}}$. The missed part of $S_{\mathbf{Y} | \mathbf{Z}}$ might be recovered from higher order moments of the conditional distribution of $\mathbf{Z}$ given $\mathbf{Y}$ (Cook, 1998b; Cook and Weisberg, 1991; Li, 1991, 1992), but such issues are not addressed in this paper. We assume throughout that $S_{E(\mathbf{Z} | \mathbf{Y})}$ is non-trivial, in the sense that it contains non-zero directions, should such exist.

Both equation (2) and expression (3) lead to the use of inverse regression as an estimation tool for a fraction of or the entire central dimension reduction subspace. A popular such method is SIR, proposed by Li (1991). In SIR, the range of the one-dimensional variable Y is partitioned into a fixed number of slices and the p components of $\mathbf{Z}$ are regressed on $\tilde{Y}$, a discrete version of Y resulting from slicing its range, giving p one-dimensional regression problems, instead of the possibly high dimensional forward regression of Y on $\mathbf{Z}$. Then, a very simple nonparametric estimate of the inverse regression curve $E(\mathbf{Z} | Y)$ estimates the central dimension reduction subspace via estimating $\text{cov}\{E(\mathbf{Z} | Y)\}$. This can be based on the fact that $S[\text{cov}\{E(\mathbf{Z} | Y)\}] = S_{E(\mathbf{Z} | Y)}$ except on a set of measure 0 (for example, see Cook (1998a), proposition 11.1, and Eaton (1983), proposition 2.7). The SIR estimate of $\text{cov}\{E(\mathbf{Z} | Y)\}$ is given by

$$\widehat{\text{cov}}\{E(\mathbf{Z} | Y)\} = \sum_{h=1}^H \hat{p}_h \hat{\mathbf{m}}_h \hat{\mathbf{m}}_h'$$

where H is the fixed number of slices, $\hat{p}_h = n_h/n$, with n being the total sample size and n_h the number of observations in the h th slice, and $\hat{\mathbf{m}}_h$ is the p -vector of the average of $\hat{\mathbf{Z}}$ within slice h . Let $\hat{\lambda}_1 \geq \hat{\lambda}_2 \geq \dots \geq \hat{\lambda}_p$ be the ordered eigenvalues of $\widehat{\text{cov}}\{E(\mathbf{Z} | Y)\}$. Li (1991) proved that, if $d = \dim(S_{E(\mathbf{Z} | Y)})$, the statistic

$$L_d = n \sum_{j=d+1}^p \hat{\lambda}_j \quad (4)$$

has an asymptotic χ^2 -distribution with $(p - d)(H - d - 1)$ degrees of freedom, provided that the regressors are normal. The test statistic can be used to estimate the dimension of $S_{E(X|Y)}$ by performing tests of $d = j$ versus $d \geq j + 1$, $j = 0, \dots, p - 1$.

Other testing techniques based on inverse regression that use the same simple nonparametric estimation method as Li (1991) have been developed. Schott (1994) proposed a test which requires elliptically symmetric regressors, and for which the tuning constant is the number of observations c per slice as opposed to the number of slices H in Li (1991). To obtain the asymptotic distribution of his test statistic, Schott let c go to ∞ . Velilla (1998) introduced another testing method, which does not impose restrictions on the regressor distribution, where c is fixed and the number of slices H varies.

SIR is a simple and useful technique for reducing the dimension in a regression problem; nevertheless, it has limitations. Normality of the regressor vector $\mathbf{X}$ is required for the χ^2 asymptotic test for dimension to apply (Li, 1991). Requiring normality for the predictors was proved not to be necessary for the asymptotic result to hold in Bura and Cook (1999) and Cook (1998a), where it was shown that restrictions should be placed on the conditional covariance structure of the standardized version of $\mathbf{X}$ instead. These restrictions are trivially satisfied if $\mathbf{X}$ has a multivariate normal distribution, but they also contradict Li's (1991) claim that the asymptotic distribution of L_d does not depend on the constant variance assumption of the conditional distribution of $\mathbf{X}$ given Y . Most importantly, SIR can be ambiguous about the estimate of the dimension as the latter depends sometimes crucially on the choice of the number of slices. As a result, all methods that depend on a tuning constant related to the choice of the number of slices suffer from the same ambiguity in estimation (Schott, 1994; Velilla, 1998).

To illustrate some of the issues discussed above, we consider the *horse mussel data*: the data consist of a sample of 172 horse mussel measurements collected in the Marlborough Sounds, which are located off the north-east coast of New Zealand's South Island (Camden, 1989; Cook and Weisberg, 1994; Cook, 1998a). The response variable is muscle mass M , the edible portion of the mussel, in grams. The quantitative predictors are the shell width W in millimetres, the shell height L in millimetres, the shell length L in millimetres and the shell mass S in grams. The actual sampling method is unknown, but we assume that the data are independent and identically distributed (IID) observations from the overall mussel population. The regression software package Arc (Cook and Weisberg, 1999) was used for the computations.

A scatterplot matrix of the response, shell height, shell length, shell width and shell mass is presented in Fig. 1(a). It is evident that the linearity condition needed for SIR to work may be violated. The transformed variables $\log(W)$ and $\log(S)$ will be used in place of W and S respectively, so that the linearity condition is satisfied by the regressor variables.

The results of applying SIR to the regression of M on H , L , $\log(W)$ and $\log(S)$ are given in Tables 1 and 2; Table 1 contains the results when six slices were used and Table 2 when 15 slices were used. The rows of both tables summarize hypothesis tests of the form $d = j$ versus $d > j$. For example, the first row gives the statistic $L_0 = 156.68$ with $(p - d)(H - d - 1) = (4 - 0)(6 - 1) = 20$ degrees of freedom and a p -value of 0.000. As we can see from Tables 1 and 2 SIR gives contradictory results: it estimates the dimension to be 1 or 2, depending on the number of slices used.

3. Parametric inverse regression

The proposed new dimension reduction method of PIR fits smooth parametric curves on the p inverse regressions via a multivariate linear model. No distributional restrictions are

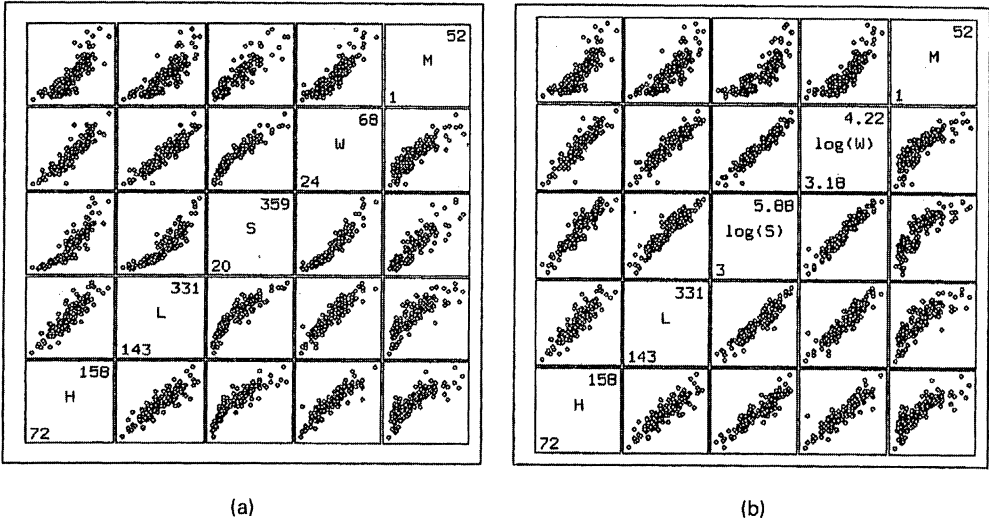


Fig. 1. Scatterplots of the mussel data: (a) untransformed predictors; (b) transformed predictors

Table 1. SIR results for $H = 6$

j	L_j	Degrees of freedom	p -value
0	156.68	20	0.000
1	22.992	12	0.028
2	9.0924	6	0.168

Table 2. SIR results for $H = 15$

j	L_j	Degrees of freedom	p -value
0	173.55	48	0.000
1	35.477	33	0.352
2	15.614	20	0.740

imposed on the regressor vector. To model the conditional expectation of $\mathbf{Z}$, the standardized version of the regressor vector $\mathbf{X}$, given $\mathbf{Y}$, a multivariate linear model is fitted with $\mathbf{Z} = (z_1, \dots, z_p)^T$ being the response and $\mathbf{Y} = (y_1, \dots, y_m)^T$ the explanatory vector. Let

$$E \begin{pmatrix} z_1 \\ \vdots \\ z_p \end{pmatrix}^T | \mathbf{Y} = (f_1(\mathbf{Y}) \dots f_q(\mathbf{Y})) \begin{pmatrix} \beta_{11} & \beta_{12} & \dots & \beta_{1p} \\ \beta_{21} & \beta_{22} & \dots & \beta_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ \beta_{q1} & \beta_{q2} & \dots & \beta_{qp} \end{pmatrix}$$

where the f_i are arbitrary, $\mathbb{R}$ -valued linearly independent known functions of $\mathbf{Y}$. Suppose that a random sample of size n is available on $(\mathbf{Y}, \mathbf{X})$ resulting in the $n \times m$ matrix $\mathcal{Y}_n$ of observations on the responses, and in the $n \times p$ matrix $\mathbf{X}_n$ of observations on the predictors. Then, including a matrix of errors $\mathbf{E}_n$, the model becomes

$$\mathbf{Z}_n|\mathcal{Y}_n = \mathbf{F}_n\mathbf{B} + \mathbf{E}_n \quad (5)$$

where $\mathbf{Z}_n = (z_{ij}) = \{\mathbf{X}_n - E(\mathbf{X}_n)\}\boldsymbol{\Sigma}_x^{-1/2}$, an $n \times p$ random matrix, $\mathbf{F}_n = (\tilde{f}_{il})$, an $n \times q$ fixed matrix with $\tilde{f}_{il} = \tilde{f}_l(\mathbf{Y}_i) = f_{il} - \sum_{i=1}^n f_{il}/n$ being the centred version of $f_{il} = f_l(\mathbf{Y}_i)$, and $\mathbf{B} = (\beta_{lj})$, the $q \times p$ matrix of coefficients. Centring is used so that the model is consistent with the fact that the expectation of the column averages of $\mathbf{Z}_n$ equals 0. The error matrix $\mathbf{E}_n$ satisfies

$$\begin{aligned} E(\mathbf{E}_n|\mathcal{Y}_n) &= 0, \\ \text{cov}\{\text{vec}(\mathbf{E}_n)|\mathcal{Y}_n\} &= \boldsymbol{\Sigma}_{z|y} \otimes \mathbf{I}_n \end{aligned} \quad (6)$$

where $\boldsymbol{\Sigma}_{z|y}$ is a $p \times p$ positive definite, unknown matrix that does not depend on $\mathbf{Y}$ and $\text{vec}(\mathbf{E}_n)$ is the vector produced by concatenating the columns of the error matrix $\mathbf{E}_n$. The symbol $\otimes$ denotes the Kronecker product. Clearly, the rank of $\mathbf{F}_n$ is q . We assume that $n \geq \max(p, q)$ to avoid trivial cases. No distributional assumptions on the errors are made except that the rows of the error matrix $\mathbf{E}_n$ are conditionally independent with mean 0 and constant covariance matrix $\boldsymbol{\Sigma}_{z|y}$, given all the responses $\mathcal{Y}_n$.

According to equation (5), $S_{E(Z|Y)}$ is the linear subspace of $S_{Y|Z}$ which is spanned by the rows of $\mathbf{F}_n\mathbf{B}$, i.e. $S(\mathbf{B}^T\mathbf{F}_n^T) = S_{E(Z|Y)}$. Therefore, since

$$\text{rank}(\mathbf{B}^T\mathbf{F}_n^T) = \text{rank}(\mathbf{F}_n\mathbf{B}),$$

$\text{rank}(\mathbf{F}_n\mathbf{B}) \leq \dim(S_{Y|Z})$, yielding that $\text{rank}(\mathbf{F}_n\mathbf{B})$ is a lower bound on the dimension of the central dimension reduction subspace.

Observe that $\text{rank}(\mathbf{F}_n\mathbf{B}) = \text{rank}(\mathbf{B}^T\mathbf{F}_n^T\mathbf{F}_n\mathbf{B}) = \text{rank}(\mathbf{B})$ since $\mathbf{F}_n^T\mathbf{F}_n$ is a positive definite matrix (see section A4.4 of Seber (1977)). In consequence, inference on the dimension of $S_{E(Z|Y)}$ can be based solely on $\mathbf{B}$ in the sense that an estimate of the rank of $\mathbf{B}$ constitutes an estimate of a lower bound on the dimension of $S_{Y|Z}$.

The estimate of $\mathbf{B}$ to be used for inference on the rank of $\mathbf{B}$ is the ordinary least squares estimate, given by

$$\hat{\mathbf{B}}_n = (\mathbf{F}_n^T\mathbf{F}_n)^{-1}\mathbf{F}_n^T\mathbf{Z}_n. \quad (7)$$

Since $\mathbf{Z}_n$ is unobservable, $\hat{\mathbf{B}}_n$ is unobservable and thus in practice it is necessary to use the sample version of $\mathbf{Z}_n$ when computing $\hat{\mathbf{B}}_n$. However, it is sufficient to work in terms of $\mathbf{Z}_n$ for deriving the asymptotic distribution of the test statistic Λ_d described in theorem 1 for the rank of $\mathbf{B}$. To find the asymptotic distribution of Λ_d , we first need to find the asymptotic distribution of a standardized version of $\hat{\mathbf{B}}_n$.

Let $\mathbf{H}_n$ denote the covariance matrix of $\sqrt{n} \text{vec}(\hat{\mathbf{B}}_n - \mathbf{B})$, i.e.

$$\mathbf{H}_n = \boldsymbol{\Sigma}_{z|y} \otimes (\mathbf{F}_n^T\mathbf{F}_n/n)^{-1}.$$

If $\mathbf{H}_n$ has a positive definite limit matrix $\mathbf{H}$, then

$$\sqrt{n} \text{vec}(\hat{\mathbf{B}}_n - \mathbf{B}) \xrightarrow{\mathcal{D}} N_{qp}(0, \mathbf{H}) \quad (8)$$

provided that certain conditions are satisfied (see lemma 1 in Appendix A).

Assume that there is a $q \times q$ positive definite matrix $\mathbf{G}$ so that

$$(\mathbf{F}_n^T\mathbf{F}_n/n)^{-1} \xrightarrow{n \rightarrow \infty} \mathbf{G}. \quad (9)$$

Also, assume that a consistent estimate $\hat{\boldsymbol{\Sigma}}_{z|y}$ is available, as $\boldsymbol{\Sigma}_{z|y}$ is usually unknown. For example, $\hat{\boldsymbol{\Sigma}}_{z|y}$ can be taken to be $(n - q)^{-1}(\mathbf{Z}_n - \mathbf{F}_n\hat{\mathbf{B}}_n)^T(\mathbf{Z}_n - \mathbf{F}_n\hat{\mathbf{B}}_n)$, which is also unbiased for $\boldsymbol{\Sigma}_{z|y}$. Let

$$\hat{\mathbf{H}}_n = \hat{\Sigma}_{z|y} \otimes (\mathbf{F}_n^T \mathbf{F}_n / n)^{-1}. \quad (10)$$

Then, if assumption (9) holds,

$$\hat{\mathbf{H}}_n \xrightarrow[n \rightarrow \infty]{} \mathbf{H} \quad \text{in probability.} \quad (11)$$

The convergence in expression (11) is a direct application of the triangle inequality and the fact that continuous functions of consistent estimates are themselves consistent. The remarks above in conjunction with a direct application of the multivariate version of Slutsky's theorem (see expression [A4.19] in Bunke and Bunke (1986)) and expression (8) obtain the following:

$$\sqrt{n} \hat{\mathbf{H}}_n^{-1/2} \text{vec}(\hat{\mathbf{B}}_n - \mathbf{B}) \xrightarrow{D} N(0, \mathbf{I}_{pq}) = N(0, \mathbf{I}_p \otimes \mathbf{I}_q). \quad (12)$$

Let $d = \dim(S_{E(\mathbf{Z}|\mathbf{Y})})$ and $\mathbf{G}_n^{-1} = \mathbf{F}_n^T \mathbf{F}_n / n$. We have shown that $d = \text{rank}(\mathbf{B})$ and thus, since $\text{rank}(\mathbf{B}) = \text{rank}(\mathbf{G}_n^{-1/2} \mathbf{B} \Sigma_{z|y}^{-1/2})$, we use the standardized matrix

$$\hat{\mathbf{B}}_{\text{std}} = \mathbf{G}_n^{-1/2} \hat{\mathbf{B}}_n \hat{\Sigma}_{z|y}^{-1/2}$$

to estimate d where

$$\tilde{\mathbf{B}}_n = (\mathbf{F}_n^T \mathbf{F}_n)^{-1} \mathbf{F}_n^T \hat{\mathbf{Z}}_n \quad (13)$$

is now the estimate of $\mathbf{B}$ based on the sample version $\hat{\mathbf{Z}}_n = (\mathbf{X}_n - \bar{\mathbf{X}}_n) \hat{\Sigma}_x^{-1/2}$ of the standardized predictor matrix $\mathbf{Z}_n$, where $\bar{\mathbf{X}}_n$ and $\hat{\Sigma}_x$ are the usual sample moment estimates of $E(\mathbf{X}_n)$ and Σ_x respectively. We give a test statistic (Λ_d) for $d = \text{rank}(\mathbf{B}) = \dim(S_{E(\mathbf{Z}|\mathbf{Y})})$ in theorem 1. The proof is given in Appendix A.

Theorem 1. Assume that model (5) holds, that $\mathbf{G}_n$ converges pointwise to a positive definite limit and that $\hat{\Sigma}_{z|y}$ is a consistent estimate of $\Sigma_{z|y}$. Let $\hat{\phi}_j$, $j = 1, \dots, \min(q, p)$, be the singular values of $\hat{\mathbf{B}}_{\text{std}}$. Then

$$\Lambda_d = n \sum_{j=d+1}^{\min(q,p)} \hat{\phi}_j^2 \quad (14)$$

is asymptotically distributed as a $\chi_{(q-d)(p-d)}^2$ random variable.

We use Λ_d as a test statistic for the rank of $\mathbf{B}$. For example, to test the hypothesis that $d = 1$ compare Λ_1 with the percentage points of a χ^2 -distribution with $(q-1)(p-1)$ degrees of freedom. As an aside, it is easy to see that the asymptotic test in theorem 1 coincides with the usual F -test for testing $d = 0$, i.e. that all the coefficients are 0, when $p = 1$.

4. The non-constant covariance case

In Sections 2 and 3 the parametric models that were used assumed that the covariance structure of the error matrix given $\mathbf{Y}$ was constant, i.e. independent of $\mathbf{Y}$. Occasionally, this assumption may be seriously violated, resulting in dimension estimation errors. In this section the non-constant error covariance structure case is addressed.

We assume that the regression model (5) holds but now $\text{cov}(\mathbf{Z}|\mathbf{Y})$ is a function of $\mathbf{Y}$:

$$\text{cov}(\mathbf{Z}|\mathbf{Y}) = \Sigma_{z|y}(\mathbf{Y}) = (\sigma_{ij}(\mathbf{Y}))_{i,j=1}^p.$$

In consequence, the covariance structure of the error matrix $\mathbf{E}_n$ can no longer be represented by the Kronecker product of $\Sigma_{z|y}$ and the identity $\mathbf{I}_n$, for

$$\text{cov}(Z_{ki}, Z_{kj} | \mathbf{Y} = \mathbf{Y}_k) = \sigma_{ij}(\mathbf{Y}_k)$$

for $k = 1, \dots, n, i, j = 1, \dots, p$. Hence, the covariance matrix of $\text{vec}(\mathbf{Z}_n | \mathcal{Y}_n)$, and consequently of $\text{vec}(\mathbf{E}_n)$, is an $np \times np$ symmetric matrix consisting of p^2 blocks of order $n \times n$, where the ij th block is the diagonal $n \times n$ matrix with $\sigma_{ij}(\mathbf{Y}_1), \dots, \sigma_{ij}(\mathbf{Y}_n)$ along its main diagonal for $i, j = 1, \dots, p$.

In vector form, $\hat{\mathbf{B}}_n$ can be written as

$$\text{vec}(\hat{\mathbf{B}}_n) = \text{vec}(\mathbf{W}_n \mathbf{Z}_n) = (\mathbf{I}_p \otimes \mathbf{W}_n) \text{vec}(\mathbf{Z}_n),$$

where $\mathbf{W}_n = (\mathbf{F}_n^T \mathbf{F}_n)^{-1} \mathbf{F}_n^T$ is a $q \times n$ known matrix of weights. The covariance matrix of $\text{vec}(\hat{\mathbf{B}}_n)$ equals

$$(\mathbf{I}_p \otimes \mathbf{W}_n) \text{cov}\{\text{vec}(\mathbf{Z}_n)\} (\mathbf{I}_p \otimes \mathbf{W}_n^T).$$

Thus, $\text{cov}\{\text{vec}(\hat{\mathbf{B}}_n)\}$ is a $qp \times qp$ block matrix, whose ij th block is given by

$$\mathbf{W}_n \text{diag}\{\sigma_{ij}(\mathbf{Y}_1), \dots, \sigma_{ij}(\mathbf{Y}_n)\} \mathbf{W}_n^T \quad (15)$$

for $i, j = 1, \dots, p$, and $\mathbf{H}_n = \text{cov}\{n^{1/2} \text{vec}(\mathbf{W}_n \mathbf{Z}_n - \mathbf{B})\}$ is a $qp \times qp$ block matrix, whose ij th block is given by expression (15) multiplied by n .

Assuming that $\mathbf{H}_n$ has a positive definite limit $\mathbf{H}$, we readily obtain

$$\sqrt{n} \text{vec}(\mathbf{W}_n \mathbf{Z}_n - \mathbf{B}) \xrightarrow{\mathcal{D}} N_{pq}(0, \mathbf{H}) \quad (16)$$

provided that conditions (a)–(c) of lemma 1 in Appendix A are satisfied. Additionally, expression (16) also holds if the conditional covariance of X_i and X_j given $\mathbf{Y}$ is asymptotically constant for all $i, j = 1, \dots, p$ (see lemma 2 in Appendix A).

Let $\hat{\Sigma}_{z|y}(\mathbf{Y}) = (\hat{\sigma}_{ij}(\mathbf{Y}))$ be a consistent estimate of $\Sigma_{z|y}(\mathbf{Y}) = (\sigma_{ij}(\mathbf{Y}))$, for all $i, j = 1, \dots, p$ and all $\mathbf{Y} \in \Omega_Y$. Let $\hat{\mathbf{H}}_n$ be the $qp \times qp$ matrix whose ij th block is given by expression (15) multiplied by n , with $\hat{\sigma}_{ij}(\mathbf{Y}_k)$ in place of $\sigma_{ij}(\mathbf{Y}_k)$.

Since $\hat{\mathbf{H}}_n$ is non-singular, if it were also consistent for $\mathbf{H}$, then by the multivariate version of Slutsky's theorem we would obtain

$$\sqrt{n} \hat{\mathbf{H}}_n^{-1/2} \text{vec}(\mathbf{W}_n \mathbf{Z}_n - \mathbf{B}) \xrightarrow{\mathcal{D}} N_{pq}(0, \mathbf{I}_{pq} = \mathbf{I}_p \otimes \mathbf{I}_q).$$

The dimension d is not affected by this non-singular transformation. Weak consistency, in the sense of $\hat{\mathbf{H}}_n \rightarrow \mathbf{H}$ in probability, is easy to establish given a weakly consistent estimate of $\Sigma_{z|y}(\mathbf{Y})$. Moreover, by placing further conditions on the entries of $\hat{\Sigma}_{z|y}(\mathbf{Y})$ we obtain that $\hat{\mathbf{H}}_n$ is L^2 consistent for $\mathbf{H}$ (see lemma 3 in Appendix A).

Let $\text{vec}(\hat{\mathbf{B}}_{\text{std}}) = \hat{\mathbf{H}}_n^{-1/2} \text{vec}(\mathbf{W}_n \mathbf{Z}_n)$. The $q \times p$ matrix $\hat{\mathbf{B}}_{\text{std}}$ that results from the arrangement of the vector $\text{vec}(\hat{\mathbf{B}}_{\text{std}})$ into p columns satisfies $\text{rank}(\hat{\mathbf{B}}_{\text{std}}) = \text{rank}(\mathbf{W}_n \hat{\mathbf{Z}}_n)$, since $\hat{\mathbf{H}}_n^{-1/2}$ is non-singular. Also, let

$$\Lambda_d = n \sum_{j=d+1}^{\min(p,q)} \hat{\phi}_j^2, \quad (17)$$

where $\hat{\phi}_j$, $j = 1, \dots, \min(p, q)$, denote the ordered singular values of $\hat{\mathbf{B}}_{\text{std}}$. The following theorem states the conditions under which the asymptotic distribution of Λ_d is χ^2 .

Theorem 2. Assume that conditions (a)–(c) of lemma 1 are satisfied and that $\hat{\mathbf{H}}_n$ is consistent or L^2 consistent for a positive definite matrix $\mathbf{H}$. If $d = \text{rank}(\mathbf{B}) = \dim(S_{E(Z|Y)})$,

then Λ_d as defined in equation (17) is asymptotically distributed as a $\chi^2_{(p-d)(q-d)}$ random variable.

Proof. The proof is analogous to the proof of theorem 1 in Section 3.

The inferential procedure on d is the same as in the constant covariance case, provided that $\Sigma_{z|y}(\mathbf{Y})$ can be estimated consistently. The computation of a consistent estimate of $\Sigma_{z|y}$ is presented next.

Assume that $E(\mathbf{Z}_j^2) < \infty$, for all $j = 1, \dots, p$. Let

$$\begin{aligned}\hat{\sigma}_{ij}(\mathbf{Y}) &= \widehat{\text{cov}}_n(\mathbf{Z}_i, \mathbf{Z}_j | \mathbf{Y}) \\ &= \hat{E}_n(\mathbf{Z}_i \mathbf{Z}_j | \mathbf{Y}) - \hat{E}_n(\mathbf{Z}_i | \mathbf{Y}) \hat{E}_n(\mathbf{Z}_j | \mathbf{Y})\end{aligned}\quad (18)$$

for $i, j \in \{1, 2, \dots, p\}$, where $\hat{E}_n(\cdot | \mathbf{Y})$ denotes the least squares estimate of $E(\cdot | \mathbf{Y})$ from regressing the argument in $(\cdot)$ on $\mathbf{Y}$. The choice of the regression model to be fitted on the argument in $(\cdot)$ is guided by the data. Under well-known conditions, least squares estimates are consistent and thus

$$\hat{\sigma}_{ij}(\mathbf{Y}) \rightarrow \sigma_{ij}(\mathbf{Y}) \quad (19)$$

in probability, for all $i, j \in \{1, 2, \dots, p\}$, and all $\mathbf{Y} \in \Omega_{\mathbf{Y}}$.

Let $\hat{\Sigma}_{z|y}(\mathbf{Y}_k)$ be the $p \times p$ matrix with entries given by equation (18), computed at $\mathbf{Y}_k$ for $k = 1, \dots, n$. Then, expression (19) implies that $\hat{\Sigma}_{z|y}(\mathbf{Y}_k)$ is a consistent estimate of $\Sigma_{z|y}(\mathbf{Y}_k)$, for all $k = 1, \dots, n$.

To obtain L^2 -consistency for the covariance matrix estimate $\hat{\mathbf{H}}_n$, we should also require that $\hat{\sigma}_{ij}(\mathbf{Y}_k)$ be L^2 consistent for $\sigma_{ij}(\mathbf{Y}_k)$, for all $i, j = 1, \dots, p$, $k = 1, \dots, n$, and that

$$\text{cov}\{\hat{\sigma}_{ij}(\mathbf{Y}_k), \hat{\sigma}_{ij}(\mathbf{Y}_l)\} \xrightarrow{n \rightarrow \infty} 0,$$

for all $k, l = 1, \dots, n$, $k \neq l$ (see lemma 3).

5. Parametric inverse regression algorithm

We can use the asymptotic distribution of Λ_d to estimate the rank d of $\mathbf{B}$, or equivalently the dimension of the subspace $S_{E(\mathbf{Z}|\mathbf{Y})} \subset S_{\mathbf{Y}|\mathbf{Z}}$, as follows: start with $j = 0$. To test the hypothesis that $d = j$, carry out the following steps.

Step 1: by visual inspection of the scatterplot matrix, decide what function(s) of $\mathbf{Y}$ fit the data the best. The choice of function(s) can be facilitated by dynamically overlaying curves on the scatterplots such as polynomials of different degrees. This can be easily implemented in software that supports dynamic graphics. Form the centred incidence matrix $\mathbf{F}_n$ for the p inverse regressions, and compute its rank q .

Step 2: standardize the regressor vector by letting $\hat{\mathbf{Z}}_n = (\mathbf{X}_n - \bar{\mathbf{X}}_n) \hat{\Sigma}_x^{-1/2}$, where $\bar{\mathbf{X}}_n = \mathbf{1}_n \bar{\mathbf{X}}^T$ and $\hat{\Sigma}_x$ is the moment estimate of the covariance matrix of $\mathbf{X}$.

Step 3: obtain the least squares matrix of estimated coefficients $\tilde{\mathbf{B}}_n$ from the p inverse regressions of the $\hat{\mathbf{Z}}_i$, $i = 1, 2, \dots, p$, on $\mathbf{Y}$.

Step 4: by visual inspection of the scatterplot matrix, decide whether the constant covariance assumption holds. If it does,

- let $\hat{\Sigma}_{z|y}$ be the matrix of residuals from the regression of $\mathbf{Z}$ on $\mathbf{Y}$ divided by $n - q$, and $\mathbf{G}_n = (\mathbf{F}_n^T \mathbf{F}_n / n)^{-1}$ and
- compute the standardized least squares matrix of coefficients, $\hat{\mathbf{B}}_{\text{std}} = \mathbf{G}_n^{-1/2} \tilde{\mathbf{B}}_n \hat{\Sigma}_{z|y}^{-1/2}$.

If it does not,

- (a) use least squares to estimate $\hat{\sigma}_{ij}(\mathbf{Y}_k)$ as described in equation (18), for $k = 1, \dots, n$,
- (b) let $\mathbf{W}_n = (\mathbf{F}_n^T \mathbf{F}_n)^{-1} \mathbf{F}_n^T$. Form the covariance matrix $\mathbf{H}_n$, as a $qp \times qp$ matrix, whose ij th block is given by expression (15) multiplied by n .
- (c) compute the standardized matrix of coefficients $\hat{\mathbf{B}}_{\text{std}}$ as described in the paragraph preceding equation (17).

Step 5: compute the ordered singular values $\hat{\phi}_j$ of $\hat{\mathbf{B}}_{\text{std}}$ and construct the test statistic Λ_d .

Step 6: compare Λ_j with the quantiles of a $\chi^2_{(p-j) \times (q-j)}$ distribution; if it is smaller conclude that $d = j$; if it is bigger, conclude that $d > j$, set $j = j + 1$ and repeat the procedure.

The d eigenvectors of the least squares estimate of $\mathbf{B}$, that correspond to its d largest eigenvalues, multiplied by $\mathbf{F}_n$ yield estimates of d of the basis vectors of $S_{Y|Z}$. They, in turn, can be scaled back to estimates of basis vectors of the central dimension reduction subspace for the non-standardized $\mathbf{X}$, by multiplication with $\hat{\Sigma}_x^{-1/2}$ on the left.

6. Applications and simulations

6.1. The horse mussel data revisited

To illustrate the PIR dimension reduction method, we reconsider the horse mussel data. The scatterplot matrix in Fig. 1(b) visually suggests fitting quadratic curves on all three inverse regression plots. The quadratic fit is also supported by overlaying polynomials of degree 2 on the scatterplots of the regressors *versus* the response. The results of the analysis are given in Table 3.

The test indicates a one-dimensional structure. Thus, one linear combination of the regressors is sufficient to characterize the behaviour of the conditional CDF of M given $H, L, \log(S)$ and $\log(W)$:

$$0.028H - 0.029L - 0.593 \log(S) + 0.804 \log(W).$$
(20)

The coefficients of the transformed regressors in expression (20) are the elements of the eigenvector corresponding to the largest eigenvalue from the eigenanalysis in step 6 of the PIR algorithm of Section 5.

6.2. Power comparison of sliced inverse regression and parametric inverse regression

This section contains a simulation-based comparison of the power of the two dimension estimation methods SIR and PIR. For all regression models to be considered in this section, three sample sizes are used: $n = 50, 100, 250$. For each sample size and each distribution of the regressor vector $\mathbf{X}$, the p -values corresponding to the test statistics for selected dimensions for both tests were collected over 1000 replications. For the PIR method, only polynomials in Y were fitted to the inverse regressions. The degree of the polynomial was decided on by a visual inspection of the scatterplots of a few initial simulated data, in conjunction with the

Table 3. Results for the mussel data

j	Λ_j	Degrees of freedom	p -value
0	701.82	8	0.000
1	5.8895	3	0.117

interactive superimposition of curves of different degrees. The latter can be easily carried out in Arc (Cook and Weisberg, 1999).

The first model to be studied has structural dimension 1. The response Y is generated according to the model

$$y = x_1 + x_2 + x_4 + 0.5\epsilon \quad (21)$$

where ϵ is a standard normal variate. Two distributions are considered for the regressor vector $\mathbf{X}$:

- (a) $\mathbf{X} \sim N(0, \mathbf{I}_4)$ and
- (b) $\mathbf{X}$ distributed as Pearson type II with parameters $m = -0.5$ and $\Sigma = \mathbf{I}_4$ (Johnson, 1987).

The Pearson type II distribution belongs to the family of elliptically contoured distributions and therefore satisfies the linearity condition that is necessary for the theory to apply.

The numerical entries of the rows of Tables 4 and 5 corresponding to the test statistics indexed by 0 are empirical estimates of the power of the corresponding test. They represent the proportion of times that the null hypothesis of dimension 0 is rejected, when the nominal significance level is 0.05. The numbers in parentheses are the analogous proportions for a significance level of 0.01. The entries of the rows of Tables 4 and 5 corresponding to test

Table 4. Empirical power and size of SIR applied to model (21)

	Results for normal $\mathbf{X}$			Results for Pearson II $\mathbf{X}$		
	$H = 5$	$H = 10$	$H = 15$	$H = 5$	$H = 10$	$H = 15$
$n = 50$						
L_0	1.0 (1.0)	1.0 (0.975)	0.964 (0.640)	0.993 (0.968)	0.941 (0.687)	0.756 (0.350)
L_1	0.053 (0.009)	0.036 (0.008)	0.025 (0.003)	0.063 (0.010)	0.041 (0.006)	0.031 (0.004)
$n = 100$						
L_0	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	0.999 (0.999)	0.997 (0.996)	1.0 (0.993)
L_1	0.053 (0.009)	0.053 (0.008)	0.036 (0.007)	0.068 (0.013)	0.052 (0.005)	0.045 (0.009)
$n = 250$						
L_0	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (0.999)	0.999 (0.999)
L_1	0.053 (0.011)	0.042 (0.005)	0.047 (0.008)	0.064 (0.015)	0.053 (0.008)	0.055 (0.014)

Table 5. Empirical power and size of PIR for model (21)

	Results for normal $\mathbf{X}$				Results for Pearson II $\mathbf{X}$	
	Degree 1	Degree 2	Degree 3	Degree 4	Degree 1	Degree 2
$n = 50$						
Λ_0	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (0.999)	0.999 (0.999)
Λ_1		0.087 (0.022)	0.025 (0.007)	0.061 (0.022)		0.037 (0.013)
$n = 100$						
Λ_0	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)
Λ_1		0.068 (0.017)	0.025 (0.005)	0.024 (0.003)		0.033 (0.003)
$n = 250$						
Λ_0	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)
Λ_1		0.052 (0.005)	0.059 (0.014)	0.061 (0.016)		0.02 (0)

statistics indexed by 1 are empirical estimates of the size of the test. The missing entries in Table 5 correspond to fitting a first-degree polynomial, thus allowing only the test of $d = 0$ versus $d = 1$ because $q = 1$. The row entries of the tables throughout this section are to be interpreted in the same or an analogous way. The symbol H stands for the number of slices, and degree stands for the degree of the fitted polynomial.

Apparently, the SIR and PIR tests are roughly equally powerful for both normally and Pearson-II-distributed X s, with PIR performing uniformly slightly better. The estimated power entries also seem to indicate that the performance of SIR deteriorates as the number of slices increases, especially when X is Pearson II distributed and the sample size is small. PIR's performance, in contrast, is consistently more powerful even when we overfit the inverse mean functions as exhibited in Table 5 where polynomials of degree 2, 3 and 4 were fitted (for normal X).

Next, we consider the following model of structural dimension 2:

$$y = x_1(x_2 + x_4 + 1) + 0.5\epsilon \tag{22}$$

where ϵ is a standard normal variate. The regressor vector $\mathbf{X}$ is assumed to have a four-variate standard normal distribution.

From Table 6 we can see that SIR would fail to detect dimension 2 with probability as high as 0.94 (when $n = 50$ and $H = 15$). For SIR to give correct results, the sample size would have to be significantly increased. In our simulations, SIR performs well when $n = 250$ and the number of slices is small. However, PIR is significantly more powerful for all sample sizes, even for a sample size of 50, and both degrees 3 and 4 (Table 7).

The last model that we study is based on model C in Velilla (1998), page 1094. The regressor vector $\mathbf{X}^T = (X_1, \dots, X_5)$ was generated so that $X_1 = W_1$, $X_2 = V_1 + W_2/2$, $X_3 = -V_1 + W_2/2$, $X_4 = V_2 + V_3$ and $X_5 = V_2 - V_3$. The only restriction placed on $\mathbf{V}$ and $\mathbf{W}$ is that they be independent. Here, V_1 , V_2 and V_4 are IID $t_{(4)}$ random variables, $V_3 \sim t_{(3)}$ and $V_5 \sim t_{(5)}$, and W_1 and W_2 are IID gamma(0.25) random variables. The dependent variable y was generated according to the model

$$y = (4 + x_1)(2 + x_2 + x_3) + 0.5\epsilon \tag{23}$$

where ϵ is a standard normal variate.

Table 6. Empirical power and size for SIR applied to model (22)

Results for the following values of H :				
	5	10	15	30
$n = 50$				
L_0	0.677 (0.486)	0.589 (0.319)	0.403 (0.161)	
L_1	0.163 (0.049)	0.123 (0.028)	0.062 (0.009)	
L_2	0.005 (0.001)	0.010 (0)	0.005 (0)	
$n = 100$				
L_0	0.951 (0.885)	0.919 (0.819)	0.879 (0.743)	
L_1	0.444 (0.244)	0.357 (0.164)	0.305 (0.121)	
L_2	0.018 (0.004)	0.023 (0.004)	0.013 (0.003)	
$n = 250$				
L_0	1.0 (0.999)	1.0 (1.0)	1.0 (0.998)	0.998 (0.99)
L_1	0.934 (0.832)	0.922 (0.821)	0.844 (0.697)	0.705 (0.477)
L_2	0.038 (0.004)	0.036 (0.001)	0.038 (0.005)	0.026 (0.002)

Table 7. Empirical power and size of PIR applied to model (22)

Results for the following values of n and degrees:						
$n = 50$		$n = 100$		$n = 250$		
Degree 3	Degree 4	Degree 3	Degree 4	Degree 3	Degree 4	
Λ_0	0.972 (0.935)	0.974 (0.938)	0.998 (0.994)	0.998 (0.993)	1.0 (1.0)	1.0 (1.0)
Λ_1	0.575 (0.388)	0.577 (0.392)	0.810 (0.668)	0.797 (0.643)	0.988 (0.969)	0.985 (0.966)
Λ_2	0.032 (0.006)	0.024 (0.007)	0.047 (0.008)	0.046 (0.009)	0.057 (0.008)	0.049 (0.007)

Model (23) has structural dimension 2. The regressor distribution satisfies the linearity condition by construction (see Velilla (1998), pages 1092–1093), even though it does not have an elliptically contoured distribution.

Table 8 indicates that SIR estimates the structural dimension to be 1, across sample sizes and choices of the number of slices. However, from Table 9 we can see that PIR is performing much better with power ranging from 56.2%, when $n = 50$, the level is 0.01 and a cubic function is fitted, to 90.4% when $n = 250$, the level is 0.05 and a fourth-degree polynomial is used.

All power calculations performed for the models presented in this section, as well as a number of power calculations for other models and/or other regressor distributions not reported in this paper, indicate that PIR is at least as powerful as SIR for models with normal

Table 8. Empirical power and size of SIR applied to model (23)

Results for the following values of H :				
	5	10	15	30
$n = 50$				
L_0	0.996 (0.996)	0.951 (0.718)	0.831 (0.5)	
L_1	0.039 (0.009)	0.075 (0.014)	0.085 (0.019)	
L_2	0.001 (0)	0.002 (0)	0.007 (0.001)	
$n = 100$				
L_0	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	
L_1	0.032 (0.008)	0.079 (0.026)	0.119 (0.035)	
L_2	0.001 (0)	0 (0)	0.007 (0)	
$n = 250$				
L_0	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)
L_1	0.049 (0.014)	0.13 (0.064)	0.164 (0.071)	0.183 (0.9)
L_2	0.002 (0)	0.004 (0)	0.007 (0)	0.011 (0.003)

Table 9. Empirical power and size of PIR applied to model (23)

Results for the following values of n and degrees:						
$n = 50$		$n = 100$		$n = 250$		
Degree 3	Degree 4	Degree 3	Degree 4	Degree 3	Degree 4	
Λ_0	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	1.0 (1.0)	
Λ_1	0.726 (0.562)	0.808 (0.668)	0.786 (0.664)	0.84 (0.732)	0.871 (0.78)	0.904 (0.836)
Λ_2	0.113 (0.021)	0.182 (0.053)	0.087 (0.02)	0.149 (0.051)	0.092 (0.027)	0.137 (0.048)

or normal-like regressor distributions. More importantly, the simulation study shows that PIR, along with its non-constant variance version, is significantly more powerful when the models are more ‘complex’, and/or the regressor distributions are far from normal, as in the case of model (23). Nevertheless, our simulations also tend to indicate that the non-constant PIR method should be employed only when the violation of the constant variance assumption is serious as it requires yet an additional estimation step.

7. Discussion

We have presented the new method of PIR for dimension reduction in a regression context. PIR fits linear models to the inverse regressions of $\mathbf{X}$ on $\mathbf{Y}$. An asymptotic χ^2 -test for the dimension d of $S_{E(\mathbf{X}|\mathbf{Y})}$ is obtained as a result of the asymptotic normality of the least squares estimate of $\mathbf{B}$. The estimated dimension is an estimate of a lower bound for the dimension of $S_{\mathbf{Y}|\mathbf{X}}$.

The results were also extended to the non-constant covariance structure case, under certain conditions. The technique of PIR does not impose any restrictions on the distribution of $\mathbf{X}$. Even though it requires choosing the parametric function to be fitted, this is guided by and based on the observed data. In contrast, SIR and its variants (Schott, 1994; Velilla, 1998) require the subjective *a priori* choice of a tuning constant, such as the number of slices, which is not suggested by the observed data but is rather arbitrarily selected. In consequence, PIR does not suffer from the ambiguity of the dimension estimation of the aforementioned methods.

The power of the PIR χ^2 -test is shown to be higher than the SIR χ^2 -test’s via a simulation study. This is not surprising as

- (a) we are fitting continuous curves based on all the data, instead of grouping the data in slices, and
- (b) the fitting method is least squares, a method yielding estimates with optimal properties.

As we confined the study only to fitting polynomials, we would expect even higher power gain by including other parametric curves that appear to reflect the data pattern more accurately. Nevertheless, PIR is desirable when marginal plots of the predictors *versus* the response give a good idea about a functional form for the inverse mean functions. When there is doubt about these mean functions, SIR may be the better choice.

PIR can be easily implemented as the computer code is distributed freely. The code can be downloaded from <http://gwis2.circ.gwu.edu/~ebura/publications.html>. It requires obtaining Arc first (Cook and Weisberg, 1999) from www.stat.umn.edu/arc.

Acknowledgements

The authors wish to thank the Joint Editor, the Associate Editor and two referees for their careful and constructive comments. Their suggestions helped to improve the paper.

Appendix A

Lemma 1. Let $\mathcal{M}_{pq}^>$ be the space of all $pq \times pq$ positive definite matrices and let $\mathcal{F}$ be the space of distributions of the errors $\mathbf{E}_n$. If

$$\mathbf{H}_n \xrightarrow{n \rightarrow \infty} \mathbf{H} \in \mathcal{M}_{pq}^> \quad (24)$$

then

$$\sqrt{n} \operatorname{vec}(\hat{\mathbf{B}}_n - \mathbf{B}) \xrightarrow{D} N_{qp}(0, \mathbf{H})$$

provided that the following three conditions are satisfied:

- (a) $\|(\mathbf{F}_n^T \mathbf{F}_n)^{-1} \mathbf{F}_n^T\|_{\max} = o(n^{-1/2})$,
- (b) $\sup_{F \in \mathcal{F}} \{\int_{\|x\| > c} \|x\|^2 dF(x)\} \rightarrow 0$ as $c \rightarrow \infty$ and
- (c) $\inf_{\Sigma \in \mathcal{M}(\mathcal{F})} \{\lambda_{\min}(\Sigma)\} \geq r > 0$,

where $\mathcal{M}(\mathcal{F}) = \{\int_{\mathbb{R}^p} xx^T dF(x): F \in \mathcal{F}\} \subset \mathcal{M}_p^+$. The notation $\|\cdot\|_{\max}$ identifies the norm on the vector space of matrices defined by $\|(a_{ij})\|_{\max} = \max_{i,j} |a_{ij}|$, for a matrix $\mathbf{A} = (a_{ij})$. Also, $\lambda_{\min}(\Sigma)$ is the smallest eigenvalue of Σ and r is some positive real value. The error distributions that are usually considered satisfy conditions (b) and (c).

Proof. The lemma follows readily from theorem 2.4.3 of Bunke and Bunke (1986), and the multivariate version of Slutsky's theorem (see expression [A 4.19] of Bunke and Bunke (1986)).

A.1. Proof of theorem 1

Consider the singular value decomposition of $\mathbf{G}^{-1/2} \mathbf{B} \Sigma_{z|y}^{-1/2}$, where $\mathbf{G}$ is the positive definite limit matrix of $n(\mathbf{F}_n^T \mathbf{F}_n)^{-1}$,

$$\mathbf{G}^{-1/2} \mathbf{B} \Sigma_{z|y}^{-1/2} = \Gamma_1^T \begin{pmatrix} \mathbf{D} & 0 \\ 0 & 0 \end{pmatrix} \Gamma_2^T,$$

$\mathbf{D}$ is a $d \times d$ diagonal matrix with the positive singular values of $\mathbf{G}^{-1/2} \mathbf{B} \Sigma_{z|y}^{-1/2}$ along its diagonal. Partition $\Gamma_1^T = (\Gamma_{11}, \Gamma_{12})$ ($q \times q$), Γ_{11} ($q \times d$), Γ_{12} ($q \times (q-d)$) and $\Gamma_2^T = (\Gamma_{21}, \Gamma_{22})^T$ ($p \times p$), where Γ_{21}^T is $d \times p$ and Γ_{22}^T is $(p-d) \times p$. By the Eaton-Tyler result (Eaton and Tyler, 1994) about the asymptotic distribution of the singular values of a matrix, the limiting distribution of the smallest $\min(q-d, p-d)$ singular values of

$$\sqrt{n}(\mathbf{G}_n^{-1/2} \hat{\mathbf{B}}_n \hat{\Sigma}_{z|y}^{-1/2})$$

is the same as the limiting distribution of the singular values of the $(q-d) \times (p-d)$ matrix

$$\sqrt{n} \mathbf{B}_n = \sqrt{n}(\Gamma_{12}^T \mathbf{G}_n^{-1/2} \hat{\mathbf{B}}_n \hat{\Sigma}_{z|y}^{-1/2} \Gamma_{22}).$$

By expression (12), we have

$$\sqrt{n} \operatorname{vec}(\Gamma_{12}^T \mathbf{G}_n^{-1/2} \hat{\mathbf{B}}_n \hat{\Sigma}_{z|y}^{-1/2} \Gamma_{22}) \xrightarrow{D} N_{(p-d)(q-d)}(0, \mathbf{I}_{p-d} \otimes \mathbf{I}_{q-d}) \quad (25)$$

where $\hat{\mathbf{B}}_n$ is the ordinary least squares estimator of $\mathbf{B}$ given in equation (7). Observe that, since $\mathbf{F}_n^T \mathbf{1}_n = 0$, $\tilde{\mathbf{B}}_n = \hat{\mathbf{B}}_n \Sigma_x^{1/2} \hat{\Sigma}_x^{-1/2}$ and thus

$$\sqrt{n}(\tilde{\mathbf{B}}_n - \mathbf{B}) = \sqrt{n}(\hat{\mathbf{B}}_n - \mathbf{B}) \Sigma_x^{1/2} \hat{\Sigma}_x^{-1/2} + \sqrt{n} \mathbf{B}(\Sigma_x^{1/2} \hat{\Sigma}_x^{-1/2} - \mathbf{I}). \quad (26)$$

Then

$$\begin{aligned} \sqrt{n} \operatorname{vec}(\Gamma_{12}^T \mathbf{G}_n^{-1/2} \tilde{\mathbf{B}}_n \hat{\Sigma}_{z|y}^{-1/2} \Gamma_{22}) &= \sqrt{n} \operatorname{vec}\{\Gamma_{12}^T \mathbf{G}_n^{-1/2} (\tilde{\mathbf{B}}_n - \mathbf{B}) \hat{\Sigma}_{z|y}^{-1/2} \Gamma_{22}\} + o_p(1) \\ &= \sqrt{n} \operatorname{vec}\{\Gamma_{12}^T \mathbf{G}_n^{-1/2} (\hat{\mathbf{B}}_n - \mathbf{B}) \Sigma_x^{-1/2} \hat{\Sigma}_x^{-1/2} \hat{\Sigma}_{z|y}^{-1/2} \Gamma_{22}\} \\ &\quad + \sqrt{n} \operatorname{vec}\{\Gamma_{12}^T \mathbf{G}_n^{-1/2} \mathbf{B}(\Sigma_x^{1/2} \hat{\Sigma}_x^{-1/2} - \mathbf{I}) \hat{\Sigma}_{z|y}^{-1/2} \Gamma_{22}\} + o_p(1). \end{aligned}$$

The first equality follows from substituting $\tilde{\mathbf{B}}_n - \mathbf{B} + \mathbf{B}$ for $\tilde{\mathbf{B}}_n$, expanding and using the fact that $\mathbf{B} \Sigma_{z|y}^{-1/2} \Gamma_{22} = 0$. The second equality follows by substituting equation (26). Then the first term on the right-hand side of the second equality is going to the distribution indicated in expression (25). The second term is going to 0 by Slutsky's theorem because

$$\Gamma_{12}^T \mathbf{G}_n^{-1/2} \mathbf{B} \rightarrow \Gamma_{12}^T \mathbf{G}^{-1/2} \mathbf{B} = 0.$$

Consequently, Λ_d has the same asymptotic distribution as the sum of the squares of the singular values of $\sqrt{n}(\mathbf{F}_{12}^T \mathbf{G}_n^{-1/2} \hat{\mathbf{B}}_n \hat{\Sigma}_{z|y}^{-1/2} \mathbf{T}_{22})$ which is $\chi_{(p-d) \times (q-d)}^2$ by expression (25).

Lemma 2. Suppose that $\Sigma(\mathbf{Y}_n) \rightarrow \Sigma$, as $n \rightarrow \infty$, for all $\mathbf{Y}_n$ in the $\mathbf{Y}$ -sample space, where Σ is non-singular. If $n\mathbf{W}_n\mathbf{W}_n^T = (\mathbf{F}_n^T \mathbf{F}_n/n)^{-1}$ has a positive definite limit matrix $\mathbf{G} = (g_{lm})$, then $\mathbf{H}_n$ has a positive definite limit matrix $\mathbf{H}$, and

$$\sqrt{n} \text{vec}(\mathbf{W}_n \mathbf{Z}_n - \mathbf{B}) \xrightarrow{D} N_{pq}(0, \mathbf{H})$$

provided that

$$\|\mathbf{W}_n\|_{\max} = \|(\mathbf{F}_n^T \mathbf{F}_n)^{-1} \mathbf{F}_n^T\|_{\max} = o(n^{-1/2}) \quad (27)$$

and conditions (b) and (c) of lemma 1 hold.

Proof. Since $\sigma_{ij}(\mathbf{Y}_k) \rightarrow \sigma_{ij}$ as $k \rightarrow \infty$, we have that for all $\epsilon > 0$ there exists k_0 such that for all $k \geq k_0$

$$\sigma_{ij} - \epsilon < \sigma_{ij}(\mathbf{Y}_k) < \sigma_{ij} + \epsilon.$$

Without loss of generality, we can assume that $\sum_{k=1}^{\infty} nW_{lk}W_{mk} = g_{lm} > 0$ (the development is analogous for the case $g_{lm} < 0$: when $g_{lm} = 0$ or when $\sigma_{ij} = 0$ the limit is 0). Then, for a sufficiently large $k_1 \geq k_0$ we have

$$(\sigma_{ij} - \epsilon) \sum_{k \geq k_1} nW_{lk}W_{mk} < \sum_{k \geq k_1} \sigma_{ij}(\mathbf{Y}_k) nW_{lk}W_{mk} < (\sigma_{ij} + \epsilon) \sum_{k \geq k_1} nW_{lk}W_{mk}.$$

There are two cases:

- (a) $\sigma_{ij} > 0$ and
- (b) $\sigma_{ij} < 0$.

Case (a) is equivalent to $\sigma_{ij} - \epsilon > 0$ and similarly case (b) is equivalent to $\sigma_{ij} + \epsilon < 0$, for sufficiently small ϵ . Therefore, for case (a) and for all $\epsilon > 0$ with $\min(g_{lm}, \sigma_{ij}) > \epsilon$ there exists $k_2 \geq k_1$ such that

$$(\sigma_{ij} - \epsilon)(g_{lm} - \epsilon) \leq \sum_{k \geq k_2} \sigma_{ij}(\mathbf{Y}_k) nW_{lk}W_{mk} \leq (\sigma_{ij} + \epsilon)(g_{lm} + \epsilon) \quad (28)$$

or, equivalently (since all products of ϵ are negligible) $\sum_k \sigma_{ij}(\mathbf{Y}_k) nW_{lk}W_{mk} \rightarrow \sigma_{ij}g_{lm}$, as $n \rightarrow \infty$. Case (b) is analogous to case (a) with the inequalities in expression (28) reversed. Therefore, the ij th block of $\mathbf{H}_n$ satisfies

$$n\mathbf{W}_n \text{diag}\{\sigma_{ij}(\mathbf{Y}_1), \dots, \sigma_{ij}(\mathbf{Y}_n)\} \mathbf{W}_n^T \rightarrow \sigma_{ij} \mathbf{G},$$

which yields that $\mathbf{H}_n \rightarrow \Sigma \otimes \mathbf{G} = \mathbf{H}$. Since the Kronecker product of two positive definite matrices is positive definite (see Harville (1997), example 7, page 369), $\mathbf{H}$ is positive definite. By a direct application of Slutsky's theorem (expression [A 4.19] of Bunke and Bunke (1986)), we obtain that

$$\mathbf{H}_n^{-1/2} \text{vec}(\mathbf{W}_n \mathbf{Z}_n - \mathbf{B}) \xrightarrow{D} N_{pq}(0, \mathbf{I}_{pq} = \mathbf{I}_p \otimes \mathbf{I}_q)$$

if and only if

$$n^{1/2} \text{vec}(\mathbf{W}_n \mathbf{Z}_n - \mathbf{B}) \xrightarrow{D} N_{pq}(0, \mathbf{H}). \quad (29)$$

Now, a sufficient condition for expression (29) is equation (27) (see Bura (1996), chapter 5).

Lemma 3. Suppose that

$$n\mathbf{W}_n\mathbf{W}_n^T \xrightarrow{n \rightarrow \infty} \mathbf{G} \in \mathcal{M}_q^>. \quad (30)$$

Suppose that $\hat{\sigma}_{ij}(\mathbf{Y})$ converges to $\sigma_{ij}(\mathbf{Y})$ in quadratic mean, for all $i, j = 1, \dots, p$, and all $\mathbf{Y}$ in the relevant sample space. Also, suppose that $\text{cov}\{\hat{\sigma}_{ij}(\mathbf{Y}_k), \hat{\sigma}_{ij}(\mathbf{Y}_l)\} \rightarrow 0$ as $n \rightarrow \infty$, for $k, l = 1, \dots, n, k \neq l$. Then, $\hat{\mathbf{H}}_n$ is an L^2 -consistent estimate of $\mathbf{H}$.

Proof. Since $\mathbf{H}$ is the limit matrix of $\mathbf{H}_n$, it suffices to show that

$$\hat{\mathbf{H}}_n - \mathbf{H}_n \xrightarrow{n \rightarrow \infty} 0 \quad \text{in } L^2. \quad (31)$$

Then, from the triangle inequality, it follows that $\hat{\mathbf{H}}_n$ is an L^2 -consistent estimate of $\mathbf{H}$. Consider the ij th blocks of $\hat{\mathbf{H}}_n$ and $\mathbf{H}_n$. The lm th entry of the ij th block of $\mathbf{H}_n$ is

$$\sum_k^n \sigma_{ij}(\mathbf{Y}_k) n W_{lk} W_{mk} \quad (32)$$

and of $\hat{\mathbf{H}}_n$ is given by expression (32) with $\hat{\sigma}_{ij}$ in place of σ_{ij} , for all $m, n = 1, \dots, q$, and all $i, j = 1, \dots, p$. Then, expression (31) is true if and only if

$$\sum_k^n \hat{\sigma}_{ij}(\mathbf{Y}_k) n W_{lk} W_{mk} - \sum_k^n \sigma_{ij}(\mathbf{Y}_k) n W_{lk} W_{mk} \xrightarrow{n \rightarrow \infty} 0$$

in L^2 , for all $m, n = 1, \dots, q$, and $i, j = 1, \dots, p$. Now,

$$\begin{aligned} E \left\{ \sum_{k=1}^n \hat{\sigma}_{ij}(\mathbf{Y}_k) n W_{lk} W_{mk} - \sum_{k=1}^n \sigma_{ij}(\mathbf{Y}_k) n W_{lk} W_{mk} \right\}^2 &= E \left[\sum_{k=1}^n \{ \hat{\sigma}_{ij}(\mathbf{Y}_k) - \sigma_{ij}(\mathbf{Y}_k) \}^2 (n W_{lk} W_{mk})^2 \right] \\ &+ E \left[\sum_{k=1}^n \sum_{\substack{r=1 \\ r \neq k}}^n \{ \hat{\sigma}_{ij}(\mathbf{Y}_k) - \sigma_{ij}(\mathbf{Y}_k) \} \{ \hat{\sigma}_{ij}(\mathbf{Y}_r) - \sigma_{ij}(\mathbf{Y}_r) \} n W_{lk} W_{mk} n W_{lr} W_{mr} \right] \\ &= \sum_{k=1}^n E \{ \hat{\sigma}_{ij}(\mathbf{Y}_k) - \sigma_{ij}(\mathbf{Y}_k) \}^2 (n W_{lk} W_{mk})^2 \end{aligned} \quad (33)$$

$$+ \sum_{k=1}^n \sum_{\substack{r=1 \\ r \neq k}}^n E [\{ \hat{\sigma}_{ij}(\mathbf{Y}_k) - \sigma_{ij}(\mathbf{Y}_k) \} \{ \hat{\sigma}_{ij}(\mathbf{Y}_r) - \sigma_{ij}(\mathbf{Y}_r) \} n W_{lk} W_{mk} n W_{lr} W_{mr}]. \quad (34)$$

The integration can be brought inside the sum by the bounded convergence theorem (see Billingsley (1986), page 214), since expression (30) holds by assumption and $\hat{\sigma}_{ij}(\mathbf{Y}_k)$ is consistent in quadratic mean for $\sigma_{ij}(\mathbf{Y}_k)$; therefore $\hat{\sigma}_{ij}(\mathbf{Y}_k) - \sigma_{ij}(\mathbf{Y}_k)$ is L^2 bounded, for all $k = 1, \dots, n$. But then, we also have that term (33) vanishes by the L^2 -consistency of $\hat{\sigma}_{ij}(\mathbf{Y}_k)$. Furthermore, term (34) goes to 0 by assumption. Hence,

$$E \left[\sum_{k=1}^n \{ \hat{\sigma}_{ij}(\mathbf{Y}_k) - \sigma_{ij}(\mathbf{Y}_k) \} n W_{lk} W_{mk} \right]^2 \xrightarrow{n \rightarrow \infty} 0$$

for all $l, m = 1, \dots, q$, $i, j = 1, \dots, p$.

References

- Billingsley, P. (1986) *Probability and Measure*. New York: Wiley.
- Bunke, H. and Bunke, O. (eds) (1986) *Statistical Inference in Linear Models*, vol. I, *Statistical Methods of Model Building*. Chichester: Wiley.
- Bura, E. (1996) Dimension reduction via inverse regression. *PhD Dissertation*. School of Statistics, University of Minnesota, St Paul.
- Bura, E. and Cook, R. D. (1999) Extending SIR: the weighted chi-square test. *J. Am. Statist. Ass.*, to be published.
- Camden, M. (1989) *The Data Bundle*. Wellington: New Zealand Statistical Association.
- Cook, R. D. (1994a) On the interpretation of regression plots. *J. Am. Statist. Ass.*, **89**, 177–189.
- (1994b) Using dimension-reduction subspaces to identify important inputs in models of physical systems. *Proc. Phys. Engng Sci. Sect. Am. Statist. Ass.*, 18–25.
- (1996) Graphics for regressions with a binary response. *J. Am. Statist. Ass.*, **91**, 983–992.
- (1998a) *Regression Graphics: Ideas for Studying Regressions through Graphics*. New York: Wiley.
- (1998b) Principal Hessian directions revisited (with discussion). *J. Am. Statist. Ass.*, **93**, 84–100.
- Cook, R. D. and Nachtshiem, C. J. (1994) Re-weighting to achieve elliptically contoured covariates in regression. *J. Am. Statist. Ass.*, **89**, 592–599.
- Cook, R. D. and Weisberg, S. (1991) Discussion on ‘Sliced inverse regression’ (by K. C. Li). *J. Am. Statist. Ass.*, **86**, 328–332.
- (1994) *An Introduction to Regression Graphics*. New York: Wiley.
- (1999) *Applied Regression including Computing and Graphics*. New York: Wiley.

- Dawid, A. P. (1979) Conditional independence in statistical theory (with discussion). *J. R. Statist. Soc. B*, **41**, 1–31.
- Diaconis, P. and Freedman, D. (1984) Asymptotics of graphical projection pursuit. *Ann. Statist.*, **12**, 793–815.
- Eaton, M. L. (1983) *Multivariate Statistics: a Vector Space Approach*. New York: Wiley.
- (1986) A characterization of spherical distributions. *J. Multiv. Anal.*, **20**, 272–276.
- Eaton, M. L. and Tyler, D. E. (1994) The asymptotic distribution of singular values with applications to canonical correlations and correspondence analysis. *J. Multiv. Anal.*, **34**, 439–446.
- Hall, P. and Li, K.-C. (1993) On almost linearity of low dimensional projections from high dimensional data. *Ann. Statist.*, **21**, 867–889.
- Harville, D. A. (1997) *Matrix Algebra from a Statistician's Perspective*. New York: Springer.
- Johnson, M. E. (1987) *Multivariate Statistical Simulation*. New York: Wiley.
- Li, K.-C. (1991) Sliced inverse regression for dimension reduction. *J. Am. Statist. Ass.*, **86**, 316–342.
- (1992) On principal Hessian directions for data visualization and dimension reduction: another application of Stein's lemma. *J. Am. Statist. Ass.*, **87**, 1025–1039.
- Schott, J. R. (1994) Determining the dimensionality in sliced inverse regression. *J. Am. Statist. Ass.*, **89**, 141–148.
- Seber, G. A. F. (1977) *Linear Regression Analysis*. New York: Wiley.
- Velilla, S. (1998) Assessing the number of linear components in a general regression problem. *J. Am. Statist. Ass.*, **93**, 1088–1098.