

High Dimensional Dimensional Reduction

David J. Olive *

Southern Illinois University

November 20, 2025

Abstract

Consider a regression or classification model with response variable Y that depends on the predictors $\mathbf{x} = (x_1, \dots, x_p)^T$ through the sufficient predictor $SP = \alpha + \mathbf{x}^T \boldsymbol{\beta}$. Let n be the number of cases. For a high dimensional model, n/p is small.

KEY WORDS: marginal maximum likelihood estimator, PCA, PLS.

1 INTRODUCTION

High dimensional statistics are used when $n < 5p$ where n is the sample size and p is the number of variables. Such a model is *overfitting*: the model does not have enough data to estimate p parameters accurately. Then n tends to not be large enough for the classical statistical method to be useful. An alternative (but less general) definition of high dimensional statistics is that p is large. Sometimes $p > Kn$ with $K \geq 10$ is called ultrahigh dimensional statistics.

In high dimensions, it is very difficult to estimate a $p \times 1$ vector $\boldsymbol{\theta}$. This result is a form of “the curse of dimensionality.” If a $\sqrt{n}$ consistent estimator of $\boldsymbol{\theta}$ is available, then the squared norm

$$\|\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}\|^2 = \sum_{i=1}^p (\hat{\theta}_i - \theta_i)^2 \propto p/n. \quad (1)$$

When p is fixed, $p/n \rightarrow 0$ as $n \rightarrow \infty$ and $\hat{\boldsymbol{\theta}}$ is a consistent estimator of $\boldsymbol{\theta}$. In high dimensions, often the estimator has not been shown to be consistent, except under very strong regularity conditions.

Some important statistical methods include regression, multivariate statistics, and classification. These methods are important for statistical learning $\approx$ machine learning, an important part of artificial intelligence. Let predictor variables for regression or multivariate statistics be $\mathbf{x} = (x_1, \dots, x_p)^T$. Let Y be a response variable for regression or

*David J. Olive is Professor, School of Mathematical & Statistical Sciences, Southern Illinois University, Carbondale, IL 62901, USA.

classification. Important regression models include generalized linear models, nonlinear regression, nonparametric regression, and survival regression models. There are n cases $(Y_i, \mathbf{x}_i^T)^T$, and for some important models, Y depends on $\mathbf{x}$ through the sufficient predictor $SP = \alpha + \mathbf{x}^T \boldsymbol{\beta}$. Some important classification models include binary regression, linear discriminant analysis, and quadratic discriminant analysis.

A binary regression model is $Y = Y|SP \sim \text{binomial}(1, \rho(SP))$ where $\rho(SP) = P(Y = 1|SP)$. There are many binary regression models, including binary logistic regression, binary probit regression, and support vector machines (SVMs) (with $Z_i = 2Y_i - 1$).

Let the multiple linear regression (MLR) model

$$Y_i = \alpha + x_{i,1}\beta_1 + \cdots + x_{i,p}\beta_p + e_i = \alpha + \mathbf{x}_i^T \boldsymbol{\beta} + e_i \quad (2)$$

for $i = 1, \dots, n$. In matrix form, this model is $\mathbf{Y} = \mathbf{X}\boldsymbol{\delta} + \mathbf{e}$, where $\mathbf{Y}$ is an $n \times 1$ vector of dependent variables, $\mathbf{X}$ is an $n \times (p+1)$ matrix of predictors, $\boldsymbol{\delta} = (\alpha, \boldsymbol{\beta}^T)^T$ is a $(p+1) \times 1$ vector of unknown coefficients, and $\mathbf{e}$ is an $n \times 1$ vector of unknown errors. Assume that the e_i are independent and identically distributed (iid) with expected value $E(e_i) = 0$ and variance $V(e_i) = \sigma^2$. A multiple linear regression model with heterogeneity has the zero mean e_i independent with $V(e_i) = \sigma_i^2$.

For estimation with ordinary least squares, let the covariance matrix of $\mathbf{x}$ be $\text{Cov}(\mathbf{x}) = \boldsymbol{\Sigma}_{\mathbf{x}} = E[(\mathbf{x} - E(\mathbf{x}))(\mathbf{x} - E(\mathbf{x}))^T]$ and the $p \times 1$ vector $\boldsymbol{\eta} = \text{Cov}(\mathbf{x}, Y) = \boldsymbol{\Sigma}_{\mathbf{x}Y} = E[(\mathbf{x} - E(\mathbf{x}))(Y - E(Y))] = (\text{Cov}(x_1, Y), \dots, \text{Cov}(x_p, Y))^T$. Let the sample covariance matrix be

$$\hat{\boldsymbol{\Sigma}}_{\mathbf{x}} = \frac{1}{n-1} \sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{x}})(\mathbf{x}_i - \bar{\mathbf{x}})^T.$$

Let

$$\hat{\boldsymbol{\eta}} = \hat{\boldsymbol{\eta}}_n = \hat{\boldsymbol{\Sigma}}_{\mathbf{x}Y} = \mathbf{S}_{\mathbf{x}Y} = \frac{1}{n-1} \sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{x}})(Y_i - \bar{Y})$$

and

$$\tilde{\boldsymbol{\eta}} = \tilde{\boldsymbol{\eta}}_n = \tilde{\boldsymbol{\Sigma}}_{\mathbf{x}Y} = \frac{1}{n} \sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{x}})(Y_i - \bar{Y}).$$

Then the OLS estimators for model (2) are $\hat{\boldsymbol{\phi}}_{OLS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y}$, $\hat{\alpha}_{OLS} = \bar{Y} - \hat{\boldsymbol{\beta}}_{OLS}^T \bar{\mathbf{x}}$, and

$$\hat{\boldsymbol{\beta}}_{OLS} = \tilde{\boldsymbol{\Sigma}}_{\mathbf{x}}^{-1} \tilde{\boldsymbol{\Sigma}}_{\mathbf{x}Y} = \hat{\boldsymbol{\Sigma}}_{\mathbf{x}}^{-1} \hat{\boldsymbol{\Sigma}}_{\mathbf{x}Y} = \hat{\boldsymbol{\Sigma}}_{\mathbf{x}}^{-1} \hat{\boldsymbol{\eta}}.$$

For a multiple linear regression model with iid cases, $\hat{\boldsymbol{\beta}}_{OLS}$ is a consistent estimator of $\boldsymbol{\beta}_{OLS} = \boldsymbol{\Sigma}_{\mathbf{x}}^{-1} \boldsymbol{\Sigma}_{\mathbf{x}Y}$ under mild regularity conditions, while $\hat{\alpha}_{OLS}$ is a consistent estimator of $E(Y) - \boldsymbol{\beta}_{OLS}^T E(\mathbf{x})$.

Let the population correlation $\rho_{ij} = \rho_{x_i, x_j} = \text{Cor}(x_i, x_j)$ and the sample correlation $r_{ij} = r_{x_i, x_j} = \text{cor}(x_i, x_j)$. Let the population correlation matrices $\text{Cor}(\mathbf{x}) = \boldsymbol{\rho}_{\mathbf{x}} = (\rho_{ij})$ and $\text{Cor}(\mathbf{x}, Y) = \boldsymbol{\rho}_{\mathbf{x}Y} = (\rho_{x_1, Y}, \dots, \rho_{x_p, Y})^T$. Let the sample covariance matrices be $\mathbf{R}_{\mathbf{x}} = (r_{ij})$ and $\mathbf{r}_{\mathbf{x}Y} = (r_{x_1, Y}, \dots, r_{x_p, Y})^T$. Then $\hat{\boldsymbol{\Sigma}}_{\mathbf{x}}$ and $\mathbf{R}$ are dispersion estimators, and $(\bar{\mathbf{x}}, \hat{\boldsymbol{\Sigma}}_{\mathbf{x}})$ is an estimator of multivariate location and dispersion.

Suppose the positive semidefinite dispersion matrix $\boldsymbol{\Sigma}$ has eigenvalue eigenvector pairs $(\lambda_1, \mathbf{d}_1), \dots, (\lambda_p, \mathbf{d}_p)$ where $\lambda_1 \geq \lambda_2 \geq \cdots \geq \lambda_p$. Let the eigenvalue eigenvector pairs of $\hat{\boldsymbol{\Sigma}}$

be $(\hat{\lambda}_1, \hat{\mathbf{d}}_1), \dots, (\hat{\lambda}_p, \hat{\mathbf{d}}_p)$ where $\hat{\lambda}_1 \geq \hat{\lambda}_2 \geq \dots \geq \hat{\lambda}_p$. These vectors are important quantities for principal component analysis (PCA).

Principal components regression (PCR), partial least squares (PLS), and several other dimension reduction models use p linear combinations $\gamma_1^T \mathbf{x}, \dots, \gamma_p^T \mathbf{x}$. Estimating the γ_i and performing the ordinary least squares (OLS) regression of Y on $(\hat{\gamma}_1^T \mathbf{x}, \hat{\gamma}_2^T \mathbf{x}, \dots, \hat{\gamma}_k^T \mathbf{x})$ and a constant gives the k -component estimator, e.g. the k -component PLS estimator or the k -component PCR estimator, for $k = 1, \dots, J$ where $J \leq p$ and the p -component estimator is the OLS estimator $\hat{\beta}_{OLS}$. Let $\gamma_i(PCR) = \mathbf{d}_i$ and $\gamma_i = \gamma_i(PLS)$. The model selection estimator chooses one of the k -component estimators, e.g. using cross validation, and will be denoted by $\hat{\beta}_{MSPLS}$ or $\hat{\beta}_{MSPCR}$.

The k -component partial least squares estimator can be found by regressing Y on a constant and on $W_i = \hat{\gamma}_i^T \mathbf{x}$ for $i = 1, \dots, k$ where $\hat{\gamma}_i = \hat{\Sigma}_{\mathbf{x}}^{i-1} \hat{\Sigma}_{\mathbf{x}Y}$ for $i = 1, \dots, k$. See Helland (1990). Let $\mathbf{X} = [\mathbf{1} \ \mathbf{X}_1]$. Chun and Keleş (2010) noted that one way to formulate PLS is to solve an optimization problem by forming $\mathbf{b}_j = \hat{\gamma}_j$ iteratively where

$$\mathbf{b}_k = \arg \max_{\mathbf{b}} \{[\text{Cor}(\mathbf{Y}, \mathbf{X}_1 \mathbf{b})]^2 V(\mathbf{X}_1 \mathbf{b})\} \quad (3)$$

subject to $\mathbf{b}^T \mathbf{b} = 1$ and $\mathbf{b}^T \Sigma_{\mathbf{x}} \mathbf{b}_j = 0$ for $j = 1, \dots, k-1$. Here V stands for the variance. So PLS is a model free way to get predictors $\hat{\gamma}_i^T \mathbf{x}$ that are fairly highly correlated with the response, and the absolute correlations tend to decrease quickly. Brown (1993, pp. 71-72) shows that an equivalent way to compute the k -component PLS estimator is to maximize $\hat{\gamma}^T \hat{\Sigma}_{\mathbf{x}Y}$ under some constraints. If the predictors are standardized to have unit sample variance, then this method becomes a correlation vector optimization problem.

The marginal maximum likelihood estimator (MMLE) is due to Fan and Lv (2008) and Fan and Song (2010). This estimator computes the marginal regression, such as the binary logistic regression, of Y on x_i resulting in the estimator $(\hat{\alpha}_{i,M}, \hat{\beta}_{i,M})$ for $i = 1, \dots, p$. Then $\hat{\beta}_{MMLE} = (\hat{\beta}_{1,M}, \dots, \hat{\beta}_{p,M})^T$.

2 What Are Some Dimension Reduction Estimators Estimating?

Several dimension reduction methods use p linear combinations $\hat{\gamma}_1^T \mathbf{x}, \dots, \hat{\gamma}_p^T \mathbf{x}$. PCA and PLS are interesting since these two methods can be used with high dimensional data. Let $W_i = \hat{\gamma}_i^T \mathbf{x}$ for $i = 1, \dots, p$. For PLS, let $\hat{\gamma}_i = \hat{\Sigma}_{\mathbf{x}}^{j-1} \hat{\Sigma}_{\mathbf{x}Y}$ with $\hat{\Sigma}_{\mathbf{x}}^0 = \Sigma_{\mathbf{x}}^0 = \mathbf{I}_p$. For PCA, let $\hat{\mathbf{d}}_i$ be orthogonal eigenvectors of $\hat{\Sigma}_{\mathbf{x}}$ where the $\mathbf{d}_i$ are orthogonal eigenvectors of $\Sigma_{\mathbf{x}}$. Then $\hat{\gamma}_i = \hat{\mathbf{d}}_i$. In low dimensions, envelope methods have some optimality properties, and can be used with more than one response variable. See, for example, Cook and Forzani (2024).

In low dimensions with one response variable, canonical correlation analysis (CCA) also has some optimality properties. For CCA, if $(Y_i, \mathbf{x}_i^T)^T$ are iid with $V(Y) = \Sigma_Y$, then

$$M = \max_{\gamma \neq \mathbf{0}} \text{Cor}(\gamma^T \mathbf{x}, Y) = \max_{\gamma \neq \mathbf{0}} \frac{\gamma^T \Sigma_{\mathbf{x}Y}}{\sqrt{\Sigma_Y} \sqrt{\gamma^T \Sigma_{\mathbf{x}} \gamma}}.$$

This optimization problem is equivalent to maximizing

$$\Sigma_Y M^2 = \max_{\gamma \neq \mathbf{0}} \frac{\gamma^T \Sigma_{\mathbf{x}Y} \Sigma_{\mathbf{x}Y}^T \gamma}{\gamma^T \Sigma_{\mathbf{x}} \gamma}$$

which has a maximum at $\gamma = \Sigma_{\mathbf{x}}^{-1} \Sigma_{\mathbf{x}Y} = \beta_{OLS}$. Hence $\hat{\gamma}_1 = \hat{\beta}_{OLS}$ for CCA of Y and $x_1, \dots, x_p$. See Mardia, Kent, and Bibby (1979, pp. 168, 282). Hence PLS is a lot like CCA for $(Y_i, \mathbf{x}_i^T)^T$ but with more constraints, and PLS can be computed in high dimensions. From the dimension reduction literature, if Y depends on $\mathbf{x}$ only through $\alpha + \beta^T \mathbf{x}$, then under the assumption of “linearly related predictors,” $\hat{\beta}_{OLS}$ estimates $\beta_{OLS} = c\beta$ for some constant c which is often nonzero. See, for example, Cook and Weisberg (1999, p. 432). Note that in high dimensions, $\hat{\gamma}_1 = \hat{\beta}_{OLS}$ can be replaced by $\hat{\gamma}_1 = \hat{\beta}$, where $\hat{\beta}$ is a high dimensional multiple linear regression estimator, such as lasso.

Instead of using the response variable Y and the predictors $X_1, \dots, X_p$, the regression model or classification model can use Y and the predictors $W_1, \dots, W_k$. Then the k -component estimator $(\hat{\alpha}_k, \hat{\beta}_k)$ is obtained by fitting the working model

$$WSP = \alpha_k + \theta_1 W_1 + \dots + \theta_k W_k = \alpha_k + \boldsymbol{\theta}_k^T \mathbf{w}$$

where $\boldsymbol{\theta}_k = (\theta_1, \dots, \theta_k)^T$ and $\mathbf{w} = (W_1, \dots, W_k)^T$. For the k -component estimator, assume the $k \times p$ matrix

$$\hat{\mathbf{A}}_{k,n} = \hat{\mathbf{A}}_k = \begin{pmatrix} \hat{\gamma}_1^T \\ \vdots \\ \hat{\gamma}_k^T \end{pmatrix} \xrightarrow{P} \mathbf{A}_k = \begin{pmatrix} \gamma_1^T \\ \vdots \\ \gamma_k^T \end{pmatrix}.$$

Then $\hat{\mathbf{A}}_k \mathbf{x} = \mathbf{w} = (W_1, \dots, W_k)^T$, and $ESP(\mathbf{w}) = \hat{\alpha}_k + \hat{\boldsymbol{\theta}}_k^T \mathbf{w} = \hat{\alpha}_k + \hat{\boldsymbol{\theta}}_k^T \hat{\mathbf{A}}_k \mathbf{x} = \hat{\alpha}_k + \hat{\beta}_k^T \mathbf{x} = ESP(\mathbf{x})$ with $\hat{\beta}_k = \hat{\mathbf{A}}_k^T \hat{\boldsymbol{\theta}}_k$. Assume $\hat{\boldsymbol{\theta}}_k \xrightarrow{P} \boldsymbol{\theta}_k$.

For example, fit a GLM, logistic regression, a support vector machine, multiple linear regression, et cetera. The θ_i depend on the method used to fit the working model. (Also, using θ_{ki} instead of θ_i is more accurate, but suppressing the subscript k is convenient.) Parts e), f), and g) of Theorem 1 and part b) of Theorem 2 are new.

Theorem 1. Consider the above notation.

a) $ESP = \hat{\alpha}_k + \hat{\beta}_k^T \mathbf{x} = \hat{\alpha}_k + (\sum_{j=1}^k \hat{\theta}_j \hat{\gamma}_j^T) \mathbf{x}$.

b) $SP = \alpha_k + \beta_k^T \mathbf{x} = \alpha_k + (\sum_{j=1}^k \theta_j \gamma_j^T) \mathbf{x}$.

c) $\hat{\beta}_k = \sum_{j=1}^k \hat{\theta}_j \hat{\gamma}_j = \hat{\mathbf{A}}_k^T \hat{\boldsymbol{\theta}}_k$.

d) $\beta_k = \sum_{j=1}^k \theta_j \gamma_j = \mathbf{A}_k^T \boldsymbol{\theta}_k$.

e) $\hat{\beta}_{kPLS} = (\sum_{j=1}^k \hat{\theta}_j \hat{\Sigma}^{j-1}) \hat{\Sigma} \mathbf{x}_Y$

f) Under iid cases, $\beta_{kPLS} = (\sum_{j=1}^k \theta_j \Sigma^{j-1}) \Sigma \mathbf{x}_Y$.

g) If $\hat{\Sigma} \mathbf{x} \xrightarrow{P} \mathbf{V} \mathbf{x}$ and $\hat{\Sigma} \mathbf{x}_Y \xrightarrow{P} \mathbf{V} \mathbf{x}_Y$, then $\hat{\beta}_{kPLS} \xrightarrow{P} \beta_{kPLS} = (\sum_{j=1}^k \theta_j \mathbf{V}^{j-1}) \mathbf{V} \mathbf{x}_Y$.

Proof. Fit WSP to get the $ESP = \hat{\alpha}_k + \hat{\theta}_1 W_1 + \dots + \hat{\theta}_k W_k = \hat{\alpha}_k + \hat{\theta}_1 \gamma_1^T \mathbf{x} + \dots + \hat{\theta}_k \gamma_k^T \mathbf{x} = \hat{\alpha}_k + \hat{\beta}_k^T \mathbf{x}$. Equating terms gives the result.

When the cases are not iid, $\hat{\beta} = \hat{\Sigma}_{\mathbf{x}}^{-1} \hat{\Sigma}_{\mathbf{x}Y}$ may be estimating $\beta = \beta_{OLS} \neq \Sigma_{\mathbf{x}}^{-1} \Sigma_{\mathbf{x}Y}$. When the errors e_i are iid, a common assumption for OLS MLR theory is

$$n(\mathbf{X}^T \mathbf{X})^{-1} = \hat{\mathbf{V}} = \begin{pmatrix} \hat{\mathbf{V}}_{11} & \hat{\mathbf{V}}_{12} \\ \hat{\mathbf{V}}_{21} & \hat{\mathbf{V}}_{22} = n \hat{\Sigma}_{\mathbf{x}}^{-1} / (n-1) \end{pmatrix} \xrightarrow{P} \mathbf{V} = \begin{pmatrix} \mathbf{V}_{11} & \mathbf{V}_{12} \\ \mathbf{V}_{21} & \mathbf{V}_{22} \end{pmatrix}.$$

Thus $\hat{\Sigma}_{\mathbf{x}}^{-1} \xrightarrow{P} \mathbf{V}_{22}$, $\hat{\Sigma}_{\mathbf{x}} \xrightarrow{P} \mathbf{V}_{22}^{-1}$, and $\hat{\Sigma}_{\mathbf{x}Y} \xrightarrow{P} \mathbf{V}_{22}^{-1} \beta$ since $\hat{\beta} = \hat{\Sigma}_{\mathbf{x}}^{-1} \hat{\Sigma}_{\mathbf{x}Y} \xrightarrow{P} \beta$.

Remark 1. The following result is useful for several multiple linear regression estimators. Let $\mathbf{w}_i = \mathbf{A}_n \mathbf{x}_i$ for $i = 1, \dots, n$ where $\mathbf{A}_n$ is a full rank $k \times p$ matrix with $1 \leq k \leq p$.

a) Let Σ^* be $\hat{\Sigma}$ or $\tilde{\Sigma}$. Then $\Sigma_{\mathbf{w}}^* = \mathbf{A}_n \Sigma_{\mathbf{x}}^* \mathbf{A}_n^T$ and $\Sigma_{\mathbf{w}Y}^* = \mathbf{A}_n \Sigma_{\mathbf{x}Y}^*$.

b) If $\mathbf{A}_n$ is a constant matrix, then $\Sigma_{\mathbf{w}} = \mathbf{A}_n \Sigma_{\mathbf{x}} \mathbf{A}_n^T$ and $\Sigma_{\mathbf{w}Y} = \mathbf{A}_n \Sigma_{\mathbf{x}Y}$.

The following result is known. Using the notation above Theorem 1, let $\hat{\mathbf{A}}_k \mathbf{x} = \mathbf{w}$. For multiple linear regression with OLS, let $Y = \alpha + \mathbf{x}^T \beta + e$. Let the working model be $Y = \alpha_k + \boldsymbol{\theta}_k^T \mathbf{w} + \epsilon$. Then the OLS estimator $\hat{\boldsymbol{\theta}}_k = \hat{\Sigma}_{\mathbf{w}}^{-1} \hat{\Sigma}_{\mathbf{w}Y}$. By Remark 1, the k -component estimator

$$\hat{\beta}_k = \hat{\mathbf{A}}_k^T \hat{\boldsymbol{\theta}}_k = \hat{\mathbf{A}}_k^T (\hat{\mathbf{A}}_k \hat{\Sigma}_{\mathbf{x}} \hat{\mathbf{A}}_k^T)^{-1} \hat{\mathbf{A}}_k \hat{\Sigma}_{\mathbf{x}Y}.$$

Suppose $k = p$ and $\hat{\mathbf{A}}_p^{-1}$ exists. Then $\hat{\beta}_p = \hat{\beta}_{OLS} = \hat{\Sigma}_{\mathbf{x}}^{-1} \hat{\Sigma}_{\mathbf{x}Y}$ since $\hat{\beta}_p =$

$$\hat{\mathbf{A}}_p^T (\hat{\mathbf{A}}_p \hat{\Sigma}_{\mathbf{x}} \hat{\mathbf{A}}_p^T)^{-1} \hat{\mathbf{A}}_p \hat{\Sigma}_{\mathbf{x}Y} = \hat{\mathbf{A}}_p^T (\hat{\mathbf{A}}_p^T)^{-1} \hat{\Sigma}_{\mathbf{x}}^{-1} (\hat{\mathbf{A}}_p)^{-1} \hat{\mathbf{A}}_p \hat{\Sigma}_{\mathbf{x}Y} = \hat{\Sigma}_{\mathbf{x}}^{-1} \hat{\Sigma}_{\mathbf{x}Y}.$$

The following theorem shows that if the p components W_i are plugged into a model that uses maximum likelihood estimation, such as a GLM, the p -component estimator $\hat{\beta}_p = \hat{\beta}_{\mathbf{x}}$, the MLE. Similar theory holds for other maximization or minimization problems, such as quasi-likelihood and partial likelihood. The profile likelihood function $L_p(\beta_{\mathbf{x}} | \mathbf{x}) = L(\beta_{\mathbf{x}}, \hat{\boldsymbol{\eta}} | \mathbf{x})$ where L is the likelihood function of all of the parameters $(\beta_{\mathbf{x}}, \boldsymbol{\eta})$ and $\hat{\boldsymbol{\eta}}$ is the MLE of $\boldsymbol{\eta}$. As above, use $\hat{\boldsymbol{\theta}} = \hat{\boldsymbol{\theta}}_{\mathbf{w}}$ to denote the MLE with $\mathbf{w}$ instead of $\hat{\beta}_{\mathbf{w}}$.

Theorem 2. Suppose the profile likelihood function $L_p(\beta_{\mathbf{x}} | \mathbf{x}) = \prod_{i=1}^n f(\mathbf{x}_i | \beta_{\mathbf{x}}) = \prod_{i=1}^n g(\mathbf{x}_i^T \beta_{\mathbf{x}})$ depends on $\mathbf{x}$ and $\beta_{\mathbf{x}}$ only through $\mathbf{x}^T \beta_{\mathbf{x}}$. a) If the maximum likelihood estimator is computed using $\mathbf{w} = \hat{\mathbf{A}}_p \mathbf{x}$ instead of $\mathbf{x}$, then $\hat{\beta}_{\mathbf{x}} = \hat{\mathbf{A}}_p^T \hat{\boldsymbol{\theta}} = \hat{\beta}_p$ provided that $\hat{\mathbf{A}}_p$ is nonsingular. b) Thus $\hat{\beta}_{\mathbf{x}} = \hat{\beta}_{pPLS} = (\sum_{j=1}^p \hat{\theta}_j \hat{\Sigma}_{\mathbf{x}}^{j-1}) \hat{\Sigma}_{\mathbf{x}Y}$ if the PLS components are used.

Proof. a)

$$L_p(\boldsymbol{\theta} | \mathbf{w}) = \prod_{i=1}^n g(\mathbf{w}_i^T \boldsymbol{\theta}) = \prod_{i=1}^n g(\mathbf{x}_i^T \hat{\mathbf{A}}^T \boldsymbol{\theta}) = \prod_{i=1}^n g(\mathbf{x}_i^T \beta^*).$$

Since the second to last term is maximized by $\hat{\mathbf{A}}^T \hat{\boldsymbol{\theta}} = \hat{\beta}_p$ and the last term is maximized by $\beta^* = \hat{\beta}_{\mathbf{x}}$, it follows that $\hat{\beta}_{\mathbf{x}} = \hat{\mathbf{A}}^T \hat{\boldsymbol{\theta}} = \hat{\beta}_p$, and $\hat{\boldsymbol{\theta}} = (\hat{\mathbf{A}}^T)^{-1} \hat{\beta}_{\mathbf{x}}$. Nonsingularity was used so that β^* varies through $\mathbb{R}^p$ as $\beta_{\mathbf{x}}$ varies through $\mathbb{R}^p$.

b) Plug in $\hat{\beta}_{\mathbf{x}} = \hat{\beta}_p = \hat{\beta}_{pPLS}$ from Theorem 1 e). $\square$

A useful high dimensional technique is to use PCA for dimension reduction. Let $U_1, \dots, U_p$ be the PCA linear combinations ($U_i = \hat{\gamma}_i^T \mathbf{x}$) ordered with respect to the largest eigenvalues. Then use $U_1, \dots, U_k$ in the regression or classification model where k is chosen in some manner. This method can be used for models with m response variables $Y_1, \dots, Y_m$. See, for example, Artigue and Smith (2019), Cook (2007, 2018), and Zhang and Chen (2020).

Consider a low or high dimensional regression or classification method with a univariate response variable Y . Let $W_1, \dots, W_p$ be the linear combinations ordered with respect to the highest squared correlations $r_1^2 \geq r_2^2 \geq \dots \geq r_p^2$ where the sample correlation $r_{i,Y} = \text{cor}(x_i, Y)$. As noted by Olive (2025), from a model selection viewpoint, using $W_1, \dots, W_k$ should work much better than using $U_1, \dots, U_k$. Also, the PLS components W_i should be used instead of the PCA W_i , since the PLS components are chosen to be fairly highly correlated with Y . Cook and Forzani (2021) used the PLS components as predictors for nonlinear regression.

3 The OPLS Estimator

The OPLS_{estimator} is the PLS estimator from Section 2 with $k = 1$. Then the ESP $= \hat{\alpha}_1 + \hat{\theta} \hat{\Sigma}_{\mathbf{xY}} \mathbf{x} = \hat{\alpha}_{OPLS} + \hat{\beta}_{OPLS} \mathbf{x}$ where $\hat{\beta}_{OPLS} = \hat{\theta} \hat{\Sigma}_{\mathbf{xY}}$. Let $\hat{\eta}_{OPLS} = \hat{\Sigma}_{\mathbf{xY}}$. Testing $H_0 : \mathbf{A} \hat{\beta}_{OPLS} = \mathbf{0}$ versus $H_1 : \mathbf{A} \hat{\beta}_{OPLS} \neq \mathbf{0}$ is equivalent to testing $H_0 : \mathbf{A} \hat{\eta} = \mathbf{0}$ versus $H_1 : \mathbf{A} \hat{\eta} \neq \mathbf{0}$ where $\mathbf{A}$ is a $k \times p$ constant matrix and $\hat{\eta} = \hat{\Sigma}_{\mathbf{xY}}$.

For multiple linear regression, Cook, Helland, and Su (2013) and Basa et al. (2024) showed that $\hat{\beta}_{OPLS} = \hat{\theta} \hat{\Sigma}_{\mathbf{xY}}$ estimates $\theta \Sigma_{\mathbf{xY}} = \beta_{OPLS}$ where

$$\theta = \frac{\Sigma_{\mathbf{xY}}^T \Sigma_{\mathbf{xY}}}{\Sigma_{\mathbf{xY}}^T \Sigma_{\mathbf{x}} \Sigma_{\mathbf{xY}}} \quad \text{and} \quad \hat{\theta} = \frac{\hat{\Sigma}_{\mathbf{xY}}^T \hat{\Sigma}_{\mathbf{xY}}}{\hat{\Sigma}_{\mathbf{xY}}^T \hat{\Sigma}_{\mathbf{x}} \hat{\Sigma}_{\mathbf{xY}}} \quad (4)$$

for $\Sigma_{\mathbf{xY}} \neq \mathbf{0}$. If $\Sigma_{\mathbf{xY}} = \mathbf{0}$, then $\beta_{OPLS} = \mathbf{0}$.

Next, some large sample theory is reviewed for $\hat{\eta}_{OPLS} = \hat{\Sigma}_{\mathbf{xY}}$ and OPLS for the multiple linear regression model, including some high dimensional tests for low dimensional quantities such as $H_0 : \beta_i = 0$ or $H_0 : \beta_i - \beta_j = 0$. These tests depended on iid cases, but not on linearity or the constant variance assumption. Hence the tests are useful for multiple linear regression with heterogeneity.

The following Olive and Zhang (2025) theorem gives the large sample theory for $\hat{\eta} = \widehat{\text{Cov}}(\mathbf{x}, Y)$. Olive et al. (2025) gave alternative proofs. This theory needs $\hat{\eta} = \hat{\Sigma}_{\mathbf{xY}}$ to exist for $\hat{\eta} = \hat{\Sigma}_{\mathbf{xY}}$ to be a consistent estimator of η . Let $\mathbf{x}_i = (x_{i1}, \dots, x_{ip})^T$ and let $\mathbf{w}_i$ and $\mathbf{z}_i$ be defined below where

$$\text{Cov}(\mathbf{w}_i) = \Sigma_{\mathbf{w}} = E[(\mathbf{x}_i - \mu_{\mathbf{x}})(\mathbf{x}_i - \mu_{\mathbf{x}})^T (Y_i - \mu_Y)^2] - \Sigma_{\mathbf{xY}} \Sigma_{\mathbf{xY}}^T.$$

Then the low order moments are needed for $\hat{\Sigma}_{\mathbf{z}}$ to be a consistent estimator of $\Sigma_{\mathbf{w}}$.

Theorem 3. Assume the cases $(\mathbf{x}_i^T, Y_i)^T$ are iid. Assume $E(x_{ij}^k Y_i^m)$ exist for $j = 1, \dots, p$ and $k, m = 0, 1, 2$. Let $\mu_{\mathbf{x}} = E(\mathbf{x})$ and $\mu_Y = E(Y)$. Let $\mathbf{w}_i = (\mathbf{x}_i - \mu_{\mathbf{x}})(Y_i - \mu_Y)$

with sample mean $\bar{\mathbf{w}}_n$. Let $\boldsymbol{\eta} = \boldsymbol{\Sigma} \mathbf{x}_Y$. Then (a)

$$\sqrt{n}(\bar{\mathbf{w}}_n - \boldsymbol{\eta}) \xrightarrow{D} N_p(\mathbf{0}, \boldsymbol{\Sigma} \mathbf{w}), \quad \sqrt{n}(\hat{\boldsymbol{\eta}}_n - \boldsymbol{\eta}) \xrightarrow{D} N_p(\mathbf{0}, \boldsymbol{\Sigma} \mathbf{w}), \quad (5)$$

$$\text{and } \sqrt{n}(\tilde{\boldsymbol{\eta}}_n - \boldsymbol{\eta}) \xrightarrow{D} N_p(\mathbf{0}, \boldsymbol{\Sigma} \mathbf{w}).$$

(b) Let $\mathbf{v}_i = (\mathbf{x}_i - \bar{\mathbf{x}}_n)(Y_i - \bar{Y}_n)$. Then $\hat{\boldsymbol{\Sigma}} \mathbf{w} = \hat{\boldsymbol{\Sigma}} \mathbf{v} + O_P(n^{-1/2})$. Hence $\tilde{\boldsymbol{\Sigma}} \mathbf{w} = \tilde{\boldsymbol{\Sigma}} \mathbf{v} + O_P(n^{-1/2})$.

(c) Let $\mathbf{A}$ be a $k \times p$ full rank constant matrix with $k \leq p$, assume $H_0 : \mathbf{A} \boldsymbol{\beta}_{OPLS} = \mathbf{0}$ is true, and assume $\hat{\theta} \xrightarrow{P} \theta \neq 0$. Then

$$\sqrt{n} \mathbf{A}(\hat{\boldsymbol{\beta}}_{OPLS} - \boldsymbol{\beta}_{OPLS}) \xrightarrow{D} N_k(\mathbf{0}, \theta^2 \mathbf{A} \boldsymbol{\Sigma} \mathbf{w} \mathbf{A}^T). \quad (6)$$

For the following theorem, consider a subset of k distinct elements from $\tilde{\boldsymbol{\Sigma}}$ or from $\hat{\boldsymbol{\Sigma}}$. Stack the elements into a vector, and let each vector have the same ordering. For example, the largest subset of distinct elements corresponds to

$$\text{vech}(\tilde{\boldsymbol{\Sigma}}) = (\tilde{\sigma}_{11}, \dots, \tilde{\sigma}_{1p}, \tilde{\sigma}_{22}, \dots, \tilde{\sigma}_{2p}, \dots, \tilde{\sigma}_{p-1,p-1}, \tilde{\sigma}_{p-1,p}, \tilde{\sigma}_{pp})^T = [\tilde{\sigma}_{jk}].$$

For random variables $x_1, \dots, x_p$, use notation such as $\bar{x}_j$ = the sample mean of the x_j , $\mu_j = E(x_j)$, and $\sigma_{jk} = \text{Cov}(x_j, x_k)$. Let

$$n \text{vech}(\tilde{\boldsymbol{\Sigma}}) = [n \tilde{\sigma}_{jk}] = \sum_{i=1}^n [(x_{ij} - \bar{x}_j)(x_{ik} - \bar{x}_k)].$$

For general vectors of elements, the ordering of the vectors will all be the same and be denoted by vectors such as $\hat{\mathbf{c}} = [\hat{\sigma}_{jk}]$, $\tilde{\mathbf{c}} = [\tilde{\sigma}_{jk}]$, $\mathbf{c} = [\sigma_{jk}]$, $\mathbf{v}_i = [(x_{ij} - \bar{x}_j)(x_{ik} - \bar{x}_k)]$, and $\mathbf{w}_i = [(x_{ij} - \mu_j)(x_{ik} - \mu_k)]$. Let $\bar{\mathbf{w}}_n = \sum_{i=1}^n \mathbf{w}_i / n$ be the sample mean of the $\mathbf{w}_i$. Assuming that $\text{Cov}(\mathbf{w}_i) = \boldsymbol{\Sigma} \mathbf{w}$ exists, then $E(\mathbf{w}_i) = E(\bar{\mathbf{w}}_n) = \mathbf{c}$.

The following Olive et al. (2025) theorem provides large sample theory for $\hat{\mathbf{c}}$ and $\tilde{\mathbf{c}}$. We use $\text{Cov}(\mathbf{w}_i) = \boldsymbol{\Sigma} \mathbf{d}$ to avoid confusion with the $\boldsymbol{\Sigma} \mathbf{w}$ used in Theorem 3. Note that $\mathbf{x}_i$ are dummy variables and could be replaced by $\mathbf{u}_i = (Y_{i1}, \dots, Y_{im}, x_{i1}, \dots, x_{ip})^T$ to get information about m response variables $Y_1, \dots, Y_m$.

Theorem 4. Assume the cases $\mathbf{x}_i$ are iid and that $\text{Cov}(\mathbf{w}_i) = \boldsymbol{\Sigma} \mathbf{d}$ exists. Using the above notation with $\mathbf{c}$ a $k \times 1$ vector,

- (i) $\sqrt{n}(\tilde{\mathbf{c}} - \mathbf{c}) \xrightarrow{D} N_k(\mathbf{0}, \boldsymbol{\Sigma} \mathbf{d})$.
- (ii) $\sqrt{n}(\hat{\mathbf{c}} - \mathbf{c}) \xrightarrow{D} N_k(\mathbf{0}, \boldsymbol{\Sigma} \mathbf{d})$.
- (iii) $\hat{\boldsymbol{\Sigma}} \mathbf{d} = \hat{\boldsymbol{\Sigma}} \mathbf{v} + O_P(n^{-1/2})$ and $\tilde{\boldsymbol{\Sigma}} \mathbf{d} = \tilde{\boldsymbol{\Sigma}} \mathbf{v} + O_P(n^{-1/2})$.

4 Large Sample Theory and Testing

Suppose the classification or regression model has a response variable Y that depends on the predictors $\mathbf{x}$ through $SP = \alpha + \boldsymbol{\beta}^T \mathbf{x}$. In low dimensions, important tests include a) $H_0 : \beta_i = 0$ (the Wald tests for MLR), b) $H_0 : \boldsymbol{\beta} = \mathbf{0}$ (the Anova F test for MLR), and c) $H_0 : (\beta_{i_1}, \dots, \beta_{i_k})^T = \mathbf{0}$ (the partial F test for MLR).

This section will derive some high dimensional analogs of the above tests.

4.1 Testing $H_0 : \beta = \mathbf{0}$

An Omnibus or Universal Test

This subsection follows Abid, Quaye, and Olive (2025) closely. Consider classification and regression models where the response variable Y only depends on the $p \times 1$ vector of predictors $\mathbf{x} = (x_1, \dots, x_p)^T$ through the sufficient predictor $SP = \alpha + \mathbf{x}^T \beta$. Let the covariance vector $\text{Cov}(\mathbf{x}, Y) = \Sigma_{\mathbf{x}Y}$. Assume the cases $(\mathbf{x}_i^T, Y_i)^T$ are iid random vectors for $i = 1, \dots, n$. Then for many such regression models, $\beta = \mathbf{0}$ if and only if $\Sigma_{\mathbf{x}Y} = \mathbf{0}$ where $\mathbf{0} = (0, \dots, 0)^T$ is the $p \times 1$ vector of zeroes.

The test of $H_0 : \Sigma_{\mathbf{x}Y} = \mathbf{0}$ versus $H_1 : \Sigma_{\mathbf{x}Y} \neq \mathbf{0}$ is equivalent to the high dimensional one sample test $H_0 : \mu = \mathbf{0}$ versus $H_A : \mu \neq \mathbf{0}$ applied to $\mathbf{w}_1, \dots, \mathbf{w}_n$ where $\mathbf{w}_i = (\mathbf{x}_i - \mu_{\mathbf{x}})(Y_i - \mu_Y)$ and the expected values $E(\mathbf{x}) = \mu_{\mathbf{x}}$ and $E(Y) = \mu_Y$. Since $\mu_{\mathbf{x}}$ and μ_Y are unknown, the test of $H_0 : \beta = \mathbf{0}$ versus $H_1 : \beta \neq \mathbf{0}$ is implemented by applying the one sample test to $\mathbf{v}_i = (\mathbf{x}_i - \bar{\mathbf{x}})(Y_i - \bar{Y})$ for $i = 1, \dots, n$.

Zhao et al. (2024) have an interesting result for the multiple linear regression model (2). Assume that the cases $(\mathbf{x}_i^T, Y_i)^T$ are iid with $E(Y) = \mu_Y$, $E(\mathbf{x}) = \mu_{\mathbf{x}}$ and nonsingular $\text{Cov}(\mathbf{x}) = \Sigma_{\mathbf{x}}$. Let $\beta = \beta_{OLS}$. Then testing $H_0 : \beta = \beta_0$ versus $H_1 : \beta \neq \beta_0$ is equivalent to testing $H_0 : \mu = \mathbf{0}$ versus $H_1 : \mu \neq \mathbf{0}$ with $\mu = E(\mathbf{w}_i) = \Sigma_{\mathbf{x}}(\beta - \beta_0)$ where $\mathbf{w}_i = (\mathbf{x}_i - \mu_{\mathbf{x}})(Y_i - \mu_Y - (\mathbf{x}_i - \mu_{\mathbf{x}})^T \beta_0)$, and a one sample test can be applied to $\mathbf{v}_i = (\mathbf{x}_i - \bar{\mathbf{x}})(Y_i - \bar{Y} - (\mathbf{x}_i - \bar{\mathbf{x}})^T \beta_0)$.

Abid, Quaye, and Olive (2025) used the above test for $\beta_0 = \mathbf{0}$. The resulting test can be used for many regression models, not just multiple linear regression. Suppose $\beta_D = \mathbf{D}^{-1} \Sigma_{\mathbf{x}Y}$ where $\mathbf{D}$ is a $p \times p$ nonsingular matrix. Then $\beta_D = \mathbf{0}$ if and only if $\Sigma_{\mathbf{x}Y} = \mathbf{0}$. Then $\mathbf{D}^{-1} = \theta \mathbf{I}$ for OPLS, $\mathbf{D}^{-1} = \Sigma_{\mathbf{x}}^{-1}$ for OLS, and $\mathbf{D}^{-1} = [\text{diag}(\Sigma_{\mathbf{x}})]^{-1}$ for the MMLE for multiple linear regression (MLR). By Theorem 1e), $\beta_{kPLS} = \mathbf{0}$ if $\Sigma_{\mathbf{x}Y} = \mathbf{0}$. Thus if the cases $(\mathbf{x}_i^T, Y_i)^T$ are iid, then using $\beta_0 = \mathbf{0}$ gives tests for $H_0 : \beta = \mathbf{0}$, $H_0 : \beta_{MMLE} = \mathbf{0}$ (for MLR), $H_0 : \Sigma_{\mathbf{x}Y} = \mathbf{0}$, $H_0 : \beta_{OPLS} = \mathbf{0}$, and $H_0 : \beta_{kPLS} = \mathbf{0}$. For multiple linear regression with heterogeneity, $\hat{\beta}_{OLS}$ is still a consistent estimator of $\beta = \beta_{OLS} = \Sigma_{\mathbf{x}}^{-1} \Sigma_{\mathbf{x}Y}$. Hence the test can be used when the constant variance assumption is violated.

Assume the cases $(\mathbf{x}_i^T, Y_i)^T$ are iid. For a generalized linear model and several other regression models that depend on the predictors $\mathbf{x}$ only through $SP = \alpha + \mathbf{x}^T \beta$, if $\beta = \mathbf{0}$, then the Y_i are iid and do not depend on $\mathbf{x}$, and thus satisfy a multiple linear regression model with $\beta_{OLS} = \mathbf{0}$. Typically, if $\beta \neq \mathbf{0}$, then $\Sigma_{\mathbf{x}Y} \neq \mathbf{0}$. Also see Theorem 2 b). An exception is when there is a lot of symmetry which rarely occurs with real data. For example, suppose $Y = m(SP) + e$ where the iid errors $e_i \sim N(0, \sigma_1^2)$ are independent of the predictors, $SP \sim N(0, \sigma_2^2)$, and the function m is symmetric about 0, e.g. $m(SP) = (SP)^2$. Then $\beta_{OLS} = \mathbf{0}$ and $\Sigma_{\mathbf{x}Y} = \mathbf{0}$ even if $\beta \neq \mathbf{0}$.

If $\beta_0 = \mathbf{0}$, then $\mathbf{w}_i = (\mathbf{x}_i - \mu_{\mathbf{x}})(Y_i - \mu_Y)$, and $E(\mathbf{w}_i) = E(\mathbf{u}_i) = E[\mathbf{x}_i(Y_i - \mu_Y)] = \Sigma_{\mathbf{x}Y}$. Then apply a high dimensional one sample test on the $\mathbf{v}_i = (\mathbf{x}_i - \bar{\mathbf{x}})(Y_i - \bar{Y})$. Note that the sample mean $\bar{\mathbf{v}} = \tilde{\Sigma}_{\mathbf{x}Y}$.

Suppose $\mathbf{x}_1, \dots, \mathbf{x}_n$ are iid random vectors with $E(\mathbf{x}) = \mu$ and covariance matrix $\text{Cov}(\mathbf{x}) = \Sigma$. Then the test $H_0 : \mu = \mathbf{0}$ versus $H_1 : \mu \neq \mathbf{0}$ is equivalent to the test

$H_0 : \boldsymbol{\mu}^T \boldsymbol{\mu} = 0$ versus $H_1 : \boldsymbol{\mu}^T \boldsymbol{\mu} \neq 0$. Let $\mathbf{S} = \hat{\boldsymbol{\Sigma}}$. A U-statistic for estimating $\boldsymbol{\mu}^T \boldsymbol{\mu}$ is

$$T_n = T_n(\mathbf{x}) = \frac{1}{n(n-1)} \sum_{i \neq j} \mathbf{x}_i^T \mathbf{x}_j = \frac{n\bar{\mathbf{x}}^T \bar{\mathbf{x}} - \text{tr}(\mathbf{S})}{n} \quad (7)$$

where $\text{tr}()$ is the trace function. See, for example, Abid, Quaye, and Olive (2025).

Let the variance $V(W) = V(W_{ij}) = V(\mathbf{x}_i^T \mathbf{x}_j) = \sigma_W^2$ for $i \neq j$. Let $m = \text{floor}(n/2) = \lfloor n/2 \rfloor$ be the integer part of $n/2$. So $\text{floor}(100/2) = \text{floor}(101/2) = 50$. Let the iid random variables $W_i = \mathbf{x}_{2i-1}^T \mathbf{x}_{2i}$ for $i = 1, \dots, m$. Hence $W_1, W_2, \dots, W_m = \mathbf{x}_1^T \mathbf{x}_2, \mathbf{x}_3^T \mathbf{x}_4, \dots, \mathbf{x}_{2m-1}^T \mathbf{x}_{2m}$. Note that $E(W_i) = \boldsymbol{\mu}^T \boldsymbol{\mu}$ and $V(W_i) = \sigma_W^2$. Let S_W^2 be the sample variance of the W_i :

$$S_W^2 = \frac{1}{m-1} \sum_{i=1}^m (W_i - \bar{W})^2. \quad (8)$$

Zhao et al. (2024, p. 2024) showed that $\sigma_W^2 = \text{tr}(\boldsymbol{\Sigma}^2) + 2\boldsymbol{\mu}^T \boldsymbol{\Sigma} \boldsymbol{\mu}$.

The following Abid, Quaye, and Olive (2025) theorem derived the variance $V(T_n)$ under simpler regularity conditions than those in the literature. The second formula in Theorem 5a) was obtained by Chen and Qin (2010).

Theorem 5. Assume $\mathbf{x}_1, \dots, \mathbf{x}_n$ are iid, $E(\mathbf{x}_i) = \boldsymbol{\mu}$, and the variance $V(\mathbf{x}_i^T \mathbf{x}_j) = \sigma_W^2$ for $i \neq j$. Let $W_{ij} = \mathbf{x}_i^T \mathbf{x}_j$ for $i \neq j$. Let $\theta = \text{Cov}(W_{ij}, W_{id}) = \boldsymbol{\mu}^T \boldsymbol{\Sigma} \boldsymbol{\mu}$ where $j \neq d, i < j$, and $i < d$. Then

$$a) V(T_n) = \frac{2\sigma_W^2}{n(n-1)} + \frac{4(n-2)\theta}{n(n-1)} = \frac{2}{n(n-1)} \text{tr}(\boldsymbol{\Sigma}^2) + \frac{4\boldsymbol{\mu}^T \boldsymbol{\Sigma} \boldsymbol{\mu}}{n}.$$

b) If $H_0 : \boldsymbol{\mu} = \mathbf{0}$ is true, then $\theta = 0$ and

$$V_0 = V(T_n) = \frac{2\sigma_W^2}{n(n-1)} = \frac{2\text{tr}(\boldsymbol{\Sigma}^2)}{n(n-1)} = \frac{2\sigma_W^2 - 4\theta}{n(n-1)}.$$

Let $\hat{V}(T_n)$ and $\hat{V}_0(T_n)$ be consistent estimators of $V(T_n)$ and $V_0(T_n)$, respectively. Then Srivastava and Du (2008), Bai and Saranadasa (1996), Chen and Qin (2010), Li (2023), and others proved that under mild regularity conditions when H_0 is true,

$$T_n / \sqrt{\hat{V}(T_n)} = T_n / \sqrt{\hat{V}_0(T_n)} \xrightarrow{D} N(0, 1).$$

Under regularity conditions when H_0 is true, Li (2023) proved that $T_n / \sqrt{\hat{V}_0(T_n)} \xrightarrow{D} t_k$ as $p \rightarrow \infty$ for fixed $n \geq 3$ where $k = 0.5n(n-1) - 1$.

A consistent estimator of $V_0(T_n)$ needs a consistent estimator of $\sigma_W^2 = 0.5n(n-1) V_0(T_n)$. Let $s_n^2 = \hat{V}_0(T_n)$. Then one estimator is $0.5n(n-1)s_n^2 = S_W^2$ from Equation (8). An estimator nearly the same as the one used by Li (2023) is

$$0.5n(n-1)s_n^2 = \hat{\sigma}_W^2 = \frac{1}{n(n-1)} \sum_{i \neq j} (\mathbf{x}_i^T \mathbf{x}_j - T_n)^2 = \frac{1}{n(n-1)} \sum_{i \neq j} (W_{ij} - T_n)^2.$$

A New Competing Test

If the parametric distribution D is known, then the iid cases assumption can be changed to independent cases. Assume $Y_i | \mathbf{x}_i^T \boldsymbol{\beta} \sim D(\tau(\alpha + \mathbf{x}_i^T \boldsymbol{\beta}), \boldsymbol{\theta})$. If $\boldsymbol{\beta} = \mathbf{0}$, then the iid $Y_i \sim D(\tau(\alpha), \boldsymbol{\theta})$. Hence testing $H_0 : \boldsymbol{\beta} = \mathbf{0}$ vs. $H_1 : \boldsymbol{\beta} \neq \mathbf{0}$ is equivalent to testing whether the Y_i are a random sample from the $D(\tau(\alpha), \boldsymbol{\theta})$ distribution. Such a test can be done with the Kolmogorov-Smirnov test, the chi-square test, the Anderson-Darling test, the Cramér-von Mises test, et cetera. For specific distributions, there are often tests. For example, the Lilliefors test can be used to test if the Y_i are iid from a $N(\mu, \sigma^2)$ distribution where μ and σ^2 are unknown. See, for example, Kellison and London (2011, pp. 455-465), Conover (1971, pp. 295-308), Zheng, Lai, and Gould (2023), and Zheng et al. (2025).

This test has great level and extreme dimension reduction since the test does not depend on the predictors $\mathbf{x}$. The power can be sometimes be very poor if the cases are iid. a) If the $(Y_i, \mathbf{x}_i^T)^T$ are iid from a multivariate normal distribution, then the Y_i are iid $N(\mu_Y, \sigma_Y^2)$ regardless of whether $\boldsymbol{\beta} = \mathbf{0}$ or $\boldsymbol{\beta} \neq \mathbf{0}$ for the multiple linear regression model $Y | (\alpha + \mathbf{x}^T \boldsymbol{\beta}) \sim N(\alpha + \mathbf{x}^T \boldsymbol{\beta}, \sigma^2)$. b) If the $(Y_i, \mathbf{x}_i^T)^T$ are iid from some distribution where the $Y_i \in \{0, 1\}$ are binary, then the Y_i are iid $\text{bin}(n = 1, \rho_Y)$ regardless of whether $\boldsymbol{\beta} = \mathbf{0}$ or $\boldsymbol{\beta} \neq \mathbf{0}$ for the binary regression model $Y | (\alpha + \mathbf{x}^T \boldsymbol{\beta}) \sim \text{bin}(n = 1, \rho(\alpha + \mathbf{x}^T \boldsymbol{\beta}))$.

The test does not depend on $\mathbf{x}$, and can thus be done after variable selection. Also, all of the predictors can have outliers and missing values.

A Test for Binary Regression or Classification

Olive (2017, pp. 396-397) gave the result for a binary response variable $Y \in \{0, 1\}$.

Theorem 6. Let $\pi_j = P(Y = j)$ for $j = 0, 1$. Let $\boldsymbol{\mu}_j = E(\mathbf{x} | Y = j)$ for $j = 0, 1$. Then a) $\tilde{\boldsymbol{\Sigma}}_{\mathbf{x}Y} = \hat{\pi}_1 \hat{\pi}_0 (\hat{\boldsymbol{\mu}}_1 - \hat{\boldsymbol{\mu}}_0)$, and b) $\boldsymbol{\Sigma}_{\mathbf{x}, Y} = \pi_1 \pi_0 (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_0)$.

Proof. Let N_i be the number of Y s that are equal to i for $i = 0, 1$ with $n = N_1 + N_2$. Then

$$\hat{\boldsymbol{\mu}}_i = \frac{1}{N_i} \sum_{j: Y_j = i} \mathbf{x}_j$$

for $i = 0, 1$ while $\hat{\pi}_i = N_i/n$ and $\hat{\pi}_1 = 1 - \hat{\pi}_0$. Hence $\hat{\boldsymbol{\mu}}_i = \bar{\mathbf{x}}_i$ is the sample mean of the $\mathbf{x}_k$ corresponding to $Y_k = j$ for $j = 0, 1$. Then

$$\tilde{\boldsymbol{\Sigma}}_{\mathbf{x}Y} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i Y_i - \bar{\mathbf{x}} \bar{Y}.$$

$$\text{Thus } \tilde{\boldsymbol{\Sigma}}_{\mathbf{x}Y} = \frac{1}{n} \left[\sum_{j: Y_j = 1} \mathbf{x}_j (1) + \sum_{j: Y_j = 0} \mathbf{x}_j (0) \right] - \bar{\mathbf{x}} \hat{\pi}_1 =$$

$$\begin{aligned} \frac{1}{n} (N_1 \hat{\boldsymbol{\mu}}_1) - \frac{1}{n} (N_1 \hat{\boldsymbol{\mu}}_1 + N_0 \hat{\boldsymbol{\mu}}_0) \hat{\pi}_1 &= \hat{\pi}_1 \hat{\boldsymbol{\mu}}_1 - \hat{\pi}_1^2 \hat{\boldsymbol{\mu}}_1 - \hat{\pi}_1 \hat{\pi}_0 \hat{\boldsymbol{\mu}}_0 = \\ \hat{\pi}_1 (1 - \hat{\pi}_1) \hat{\boldsymbol{\mu}}_1 - \hat{\pi}_1 \hat{\pi}_0 \hat{\boldsymbol{\mu}}_0 &= \hat{\pi}_1 \hat{\pi}_0 (\hat{\boldsymbol{\mu}}_1 - \hat{\boldsymbol{\mu}}_0). \end{aligned}$$

Thus $\boldsymbol{\Sigma}_{\mathbf{x}, Y} = \pi_1 \pi_0 (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_0)$. $\square$

This result means $\boldsymbol{\eta} = \boldsymbol{\Sigma}_{\mathbf{x}, Y} = \pi_1 \pi_0 (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_0)$ and $\boldsymbol{\phi} = \boldsymbol{\mu}_1 - \boldsymbol{\mu}_0$ are quantities of interest for binary regression. Note that $\mathbf{x} = (w_1, \dots, w_k, w_1 w_2, \dots, w_1 w_k, \dots, w_{k-1} w_k)^T$ could be used to include pairwise interactions of the w_i .

Theorem 2b) suggests that typically the binary regression $\hat{\beta} = \hat{\mathbf{C}}\hat{\Sigma}\mathbf{x}_Y$. If the cases $(Y_i, \mathbf{x}_i^T)^T$ are iid, then $H_0 : \beta = \mathbf{0}$ can be tested with the omnibus test for $H_0 : \Sigma\mathbf{x}_Y$. If the cases within each group are iid, if the two groups are independent, and if $N_1/(N_1 + N_2) \rightarrow \pi_1$, then $\Sigma\mathbf{x}_Y = \pi_1\pi_0(\mu_1 - \mu_0)$ by Theorem 6b). Thus $H_0 : \beta = \mathbf{0}$ can be tested with a high dimensional two sample test for $H_0 : \mu_1 = \mu_0$.

4.2 Testing $H_0 : \beta_i = 0$

4.3 Testing $H_0 : \beta_I = (\beta_{i1}, \dots, \beta_{ik})^T = \mathbf{0}$

High Dimensional Tests

Some tests when n/p is not large are simple. Testing $H_0 : \mathbf{A}\beta_{BR} = \mathbf{0}$ versus $H_1 : \mathbf{A}\beta_{BR} \neq \mathbf{0}$ is equivalent to testing $H_0 : \mathbf{A}\eta = \mathbf{0}$ versus $H_1 : \mathbf{A}\eta \neq \mathbf{0}$ where $\mathbf{A}$ is a $k \times p$ constant matrix. Let $\text{Cov}(\hat{\eta}) = \Sigma\mathbf{w}$ be the asymptotic covariance matrix of $\hat{\eta}$. In high dimensions where $n < 5p$, we can't get a good nonsingular estimator of $\text{Cov}(\hat{\eta})$, but we can get good nonsingular estimators of $\text{Cov}((\hat{\eta}_{i1}, \dots, \hat{\eta}_{ik})^T)$ with $\mathbf{u} = (x_{i1}, \dots, x_{ik})^T$ where $n \geq Jk$ with $J \geq 10$. (Values of J much larger than 10 may be needed if some of the k predictors are skewed or if a π_i is near 0 or 1.) Simply use the sample covariance matrix with $\mathbf{u}$ replacing $\mathbf{x}$. Hence we can test hypotheses like $H_0 : \beta_i - \beta_j = 0$. In particular, testing $H_0 : \beta_i = 0$ is equivalent to testing $H_0 : \eta_i = 0$.

Data splitting uses model selection (variable selection is a special case) to reduce the high dimensional problem to a low dimensional problem. The above procedure also reduces the high dimensional problem to a low dimensional problem.

5 CONCLUSIONS

Binary regression is closely related to two sample tests. Note that $\hat{\eta} = \hat{\mu}_1 - \hat{\mu}_2$ can use other multivariate location estimators than sample means. For example, sample coordinatewise medians, sample coordinatewise trimmed means, and the Olive (2017b) T_{RMVN} estimator have large sample theory given by Rupasinghe Arachchige Don and Olive (2019) and Rupasinghe Arachchige Don and Pelawa Watagoda (2018).

Some papers on binary regression include Cai, Guo, and Ma (2023), Candès and Sur (2020), Mukherjee, Pillai, and Lin (2015), Sur and Candès (2019), Sur, Chen, and Candès (2019), and Tang and Ye (2020). Empirically, often $\beta_{LR} \approx d \beta_{OLS}$. Haggstrom (1983) suggests that d is not far from $1/\text{MSE}$ for logistic regression.

These binary regression estimators also give new ways to compare multivariate location estimators from two groups. The tests using k predictors can be performed. High dimensional tests for means from two groups can also be used. The tests that make very strong assumptions, such as multivariate normality or equal covariance matrices for the two groups, should be avoided. See Feng and Sun (2015), Gregory et al. (2015), Hu and Bai (2015), Rajapaksha and Olive (2024), and Xue and Yao (2020).

Software

The R software was used in the simulations. See R Core Team (2024). Programs will be added to the Olive (2025) collections of R functions *slpack.txt*, available from

(<http://parker.ad.siu.edu/Olive/slpack.txt>).

References

- Abid, A.M., Quaye, P.A., and Olive, D.J. (2025), “A High Dimensional Omnibus Regression Test,” *Stats*, 8, 107.
- Artigue, H., and Smith, G. (2019), “The Principal Problem with Principal Components Regression,” *Cogent Mathematics & Statistics*, 6, 1622190.
- Bai, Z.D., and Saranadasa, H. (1996), “Effects of High Dimension: by an Example of a Two Sample Problem,” *Statistica Sinica*, 6, 311-329.
- Basa, J., Cook, R.D., Forzani, L., and Marcos, M. (2024), “Asymptotic Distribution of One-Component Partial Least Squares Regression Estimators in High Dimensions,” *The Canadian Journal of Statistics*, 52, 118-130.
- Brown, P.J. (1993), *Measurement, Regression, and Calibration*, Oxford University Press, New York, NY.
- Chen, S.X., and Qin, Y.L. (2010), “A Two Sample Test for High-dimensional Data with Applications to Gene-Set Testing,” *The Annals of Statistics*, 38, 808-835.
- Conover, W.J. (1971), *Practical Nonparametric Statistics*, Wiley, New York, NY.
- Cook, R.D. (2007), “Fisher Lecture: Dimension Reduction in Regression,” *Statistical Science*, (with discussion), 22, 1-26.
- Cook, R.D. (2018), “Principal Components, Sufficient Dimension Reduction, and Envelopes,” *Annual Review of Statistics and Its Application*, 5, 533-559.
- Cook, R.D., and Forzani, L. (2021), “PLS Regression Algorithms in the Presence of Nonlinearity,” *Chemometrics and Intelligent Laboratory Systems*, 213, 104307.
- Cook, R.D., and Forzani, L. (2024), *Partial Least Squares Regression: and Related Dimension Reduction Methods*, Chapman and Hall/CRC, Boca Raton, FL.
- Cook, R.D., Helland, I.S., and Su, Z. (2013), “Envelopes and Partial Least Squares Regression,” *Journal of the Royal Statistical Society, B*, 75, 851-877.
- Cook, R.D., and Weisberg, S. (1999), *Applied Regression Including Computing and Graphics*, Wiley, New York, NY.
- Chun, H., and Keleş, S. (2010), “Sparse Partial Least Squares Regression for Simultaneous Dimension Reduction and Predictor Selection,” *Journal of the Royal Statistical Society, B*, 72, 3-25.
- Fan, J., and Lv, J. (2008), “Sure Independence Screening for Ultrahigh Dimensional Feature Space,” *Journal of the Royal Statistical Society, B*, 70, 849-911.
- Fan, J., and Song, R. (2010), “Sure Independence Screening in Generalized Linear Models with np-Dimensionality,” *The Annals of Statistics*, 38, 3217-3841.
- Helland, I.S. (1990), “Partial Least Squares Regression and Statistical Models,” *Scandinavian Journal of Statistics*, 17, 97-114.
- Kellison, S.G. and London, R.L. (2011), *Risk Models and Their Estimation*, ACTEX Publications, Winsted, CT.
- Li, J. (2023), “Finite Sample t -Tests for High-dimensional Means,” *Journal of Multivariate Analysis*, 196, 105183.
- Mardia, K.V., Kent, J.T., and Bibby, J.M. (1979), *Multivariate Analysis*, Academic Press, London, UK.
- Olive, D.J. (2017), *Linear Regression*, Springer, New York, NY.

Olive, D.J. (2025), “Some Useful Techniques for High Dimensional Statistics,” *Stats*, 8, 60.

Olive, D.J., Alshammari, A.A., Pathiranage, K.G., and Hettige, L.A.W. (2025), “Testing with the One Component Partial Least Squares and the Marginal Maximum Likelihood Estimators,” *Communications in Statistics: Theory and Methods*, to appear. <https://doi.org/10.1080/03610926.2025.2527340>

Olive, D.J., and Zhang, L. (2025), “One Component Partial Least Squares, High Dimensional Regression, Data Splitting, and the Multitude of Models,” *Communications in Statistics: Theory and Methods*, 54, 130-145.

R Core Team (2024), “R: a Language and Environment for Statistical Computing,” R Foundation for Statistical Computing, Vienna, Austria, (www.R-project.org).

Srivastava, M.S., and Du, M. (2008), “A Test for the Mean Vector with Fewer Observations Than the Dimension,” *Journal of Multivariate Analysis*, 99, 386-402.

Zhang, J., and Chen, X. (2020), “Principal Envelope Model,” *Journal of Statistical Planning and Inference*, 206, 249-262.

Zhao, A., Li, C., Li, R., and Zhang, Z. (2024), “Testing High-Dimensional Regression Coefficients in Linear Models,” *The Annals of Statistics*, 52, 2034-2058.

Zheng, W., Lai, D., and Gould, K. L. (2023), “A Simulation Study of a Class of Nonparametric Test Statistics: a Close Look of Empirical Distribution Function-Based Tests,” *Communications in Statistics - Simulation and Computation*, 52, 1132-1148.

Zheng, W., Zhu, H., Lance Gould, K., and Lai, D. (2025), “Comparing Heart PET Scans: an Adjustment of Kolmogorov-Smirnov Test under Spatial Autocorrelation,” *Journal of Applied Statistics*, 52, 253-269.