

Chatterjee

● Advanced Topics in Stat. (Prof. Shirshendu Chatterjee)Textbook: Applied Multivariate Analysis

Applied Linear Statistical model by Kutner, etc.

TopicsMultiple Linear Regression and associated inference.Logistic regression, ~~principle~~ principal component analysis,
clustering and classification, some design of experiment.HW: Biweekly HW will be posted in Blackboard.

Solution of some of them will be posted.

Exams:Two Midterms (40%) ← Tentatively on the first
class of March and April.

One project (20%)

Final Exam (40%)

← Due by the last class.

Makeup Exam (strict policy)Formula sheets (3 pages for Midterms, 4 pages for Final
and calculators are allowed)

Ch. 17 from J & W.

Multiple linear regression

$Y \leftrightarrow$ dependent variable (response)

$\begin{pmatrix} Z_1 \\ \vdots \\ Z_k \end{pmatrix} \leftrightarrow$ k independent variable (predictors)

Assume linear relation: $Y = \beta_0 + \beta_1 Z_1 + \dots + \beta_k Z_k + \varepsilon$ error

Goal: Statistical inference about $\beta_0, \dots, \beta_k$

Prediction of response based on predictor variables.

We observe the response and predictors for n observations.

This is called the training data.

$$Y_1 = \beta_0 + \beta_1 Z_{11} + \beta_2 Z_{12} + \dots + \beta_k Z_{1k} + \varepsilon_1$$

$$Y_2 = \beta_0 + \beta_1 Z_{21} + \beta_2 Z_{22} + \dots + \beta_k Z_{2k} + \varepsilon_2$$

$\vdots$

$$Y_n = \beta_0 + \beta_1 Z_{n1} + \beta_2 Z_{n2} + \dots + \beta_k Z_{nk} + \varepsilon_n$$

$$\tilde{Y} = \begin{pmatrix} Y_1 \\ Y_2 \\ \vdots \\ Y_n \end{pmatrix}, \tilde{Z} = \begin{bmatrix} 1 & Z_{11} & \dots & Z_{1k} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & Z_{n1} & \dots & Z_{nk} \end{bmatrix}, \tilde{\varepsilon} = \begin{pmatrix} \varepsilon_1 \\ \vdots \\ \varepsilon_n \end{pmatrix}, \tilde{\beta} = \begin{pmatrix} \beta_0 \\ \vdots \\ \beta_k \end{pmatrix}$$

$$\tilde{Y} = \tilde{Z}\tilde{\beta} + \tilde{\varepsilon}$$

$n \times 1$ $n \times (k+1)$ $(k+1) \times 1$ $n \times 1$

Assumption on errors

Errors have mean 0 and they are uncorrelated.

Remember that the random part in the model is the errors.

$$\mathbb{E}(\varepsilon_i) = 0 \text{ for all } i.$$

$$\text{cov}(\varepsilon_i, \varepsilon_j) = 0 \text{ for all } i \neq j.$$

$$\mathbb{E}(\varepsilon_i \varepsilon_j) = \mathbb{E}(\varepsilon_i) \mathbb{E}(\varepsilon_j)$$

$$\text{var}(\varepsilon_i) = \text{cov}(\varepsilon_i, \varepsilon_i) = \sigma^2 \quad (\text{homoskedasticity assumption})$$

In this case, we can use least square method to estimate β .

$$\text{Minimize } \sum_{i=1}^n (Y_i - \beta_0 - \beta_1 Z_{i1} - \dots - \beta_k Z_{ik})^2 \quad (*)$$

$$\mathbb{E}(Y_i) = \beta_0 + \beta_1 Z_{i1} + \dots + \beta_k Z_{ik}$$

The minimizer of (*) is called the LS estimator of β .

Notation $\hat{\beta}_{LS}$

the sum of squared error of differences.

$$= \sum_{i=1}^n (Y_i - \tilde{Z}_{i,*} \tilde{\beta})^2, \quad \tilde{Z}_{i,*} \text{ is the } i\text{th row.}$$

$$= (\tilde{Y} - \tilde{Z}\tilde{\beta})'(\tilde{Y} - \tilde{Z}\tilde{\beta}) = \tilde{\varepsilon}'\tilde{\varepsilon}$$

$$\frac{\partial}{\partial \tilde{\beta}} (\tilde{Y} - \tilde{Z}\tilde{\beta})'(\tilde{Y} - \tilde{Z}\tilde{\beta}) = 0$$

$$2\tilde{Z}'(\tilde{Y} - \tilde{Z}\tilde{\beta}) = 0 \Rightarrow \tilde{Z}'\tilde{Y} = \tilde{Z}'\tilde{Z}\tilde{\beta}$$

$$\begin{aligned} \mathbb{E}(\tilde{Y}) + \tilde{Z}\tilde{\beta} &= \tilde{Y} \\ \tilde{Z}\tilde{\beta} &= \tilde{Y} - \mathbb{E}(\tilde{Y}) \\ \tilde{\beta} &= (\tilde{Z}'\tilde{Z})^{-1}\tilde{Z}'\tilde{Y} \end{aligned}$$

$$\hat{\beta}_{LS} = (\tilde{Z}'\tilde{Z})^{-1}\tilde{Z}'\tilde{Y}$$

If $\text{rank}(Z) = k+1$, then $\hat{\beta}_{LS} = (Z'Z)^{-1}Z'Y$.

$\hat{y} = \text{predicted value of } y$
 $= Z\hat{\beta} = Z(Z'Z)^{-1}Z'Y = P_Z Y$, where

P_Z is the projection matrix on the column space of Z .

Recall that projection matrices are idempotent,

i.e., $P_Z^2 = P_Z$ Also, $I - P_Z$ is idempotent.

Also, $P_Z(I - P_Z) = 0$. $P_Z \perp (I - P_Z)$. Matrix A is idempotent if $A^2 = A$.
 e.g. $A = \begin{bmatrix} 4 & -1 \\ 12 & -3 \end{bmatrix}$, $A^2 = \begin{bmatrix} 4 & -1 \\ 12 & -3 \end{bmatrix}$

Thus, column space of P_Z and $I - P_Z$ are orthogonal.

$$E(\hat{\beta}_{LS}) = E((Z'Z)^{-1}Z'Y)$$

$$= (Z'Z)^{-1}Z'E(Y)$$

$$= (Z'Z)^{-1}Z'\beta = \beta$$

$$\text{cov}(\hat{\beta}) = \text{cov}((Z'Z)^{-1}Z'Y) = (Z'Z)^{-1}Z'\text{cov}(Y)Z(Z'Z)^{-1}$$

$$= (Z'Z)^{-1}Z'(\sigma^2 I)Z(Z'Z)^{-1}$$

$$= \sigma^2 (Z'Z)^{-1}$$

Unbiased estimator for σ^2

The estimated error $\hat{\epsilon} = Y - \hat{y}$

$$= Y - P_Z Y = (I - P_Z)Y = \hat{\epsilon}$$

$$\text{So, } \sum_{i=1}^n \hat{\epsilon}_i^2 = \hat{\epsilon}'\hat{\epsilon} = Y'(I - P_Z)'(I - P_Z)Y = Y'(I - P_Z)Y$$

$$E(\hat{\epsilon}'\hat{\epsilon}) = E[Y'(I - P_Z)Y]$$

$$= E[\text{tr}(Y'(I - P_Z)Y)]$$

$$= E[\text{tr}((I - P_Z)Y Y')] = \text{tr}(E[(I - P_Z)Y Y'])$$

$$= \text{tr}[(I - P_Z)E(Y Y')]$$

$$= \text{tr}[(I - P_Z)(\sigma^2 I + Z\hat{\beta}\hat{\beta}'Z')]$$

$$= \sigma^2 \text{tr}[(I - P_Z) + 0]$$

$$= \sigma^2 \text{tr}(I - P_Z) = \sigma^2 \text{rank}(I - P_Z) = \sigma^2(n - k - 1)$$

Since, $(I - P_Z)P_Z = 0$,

$$\text{rank}(P_Z) + \text{rank}(I - P_Z) = n$$

Also, $\text{rank}(Z) = \text{rank}(P_Z) = \# \text{ columns of } Z \text{ assuming that } Z \text{ has full column rank.}$

$$\text{So, } \text{rank}(I - P_Z) = n - (k + 1).$$

* $X'AX = \text{tr}(X'AX) = \text{tr}(AXX')$
 if A is a square symmetric $(k \times k)$ matrix and X is $(k \times 1)$ vector

$$\text{Var}(Y) + E(Y)E(Y') = Z\hat{\beta}\hat{\beta}'Z' + Z\beta\beta'Z'$$

$$\text{By } (I - P_Z)Z\hat{\beta}\hat{\beta}'Z' = 0 \implies (I - P_Z)P_Z = 0$$

So, an unbiased estimator of σ^2 is

$$\frac{\hat{\epsilon}'\hat{\epsilon}}{n-k-1} = \frac{Y'(I-P_Z)Y}{n-k-1} = \hat{\sigma}^2$$

Under the "mean 0 uncorrelated" assumption on the errors, we found unbiased estimator of β and σ^2 .

Math B7800

2/6/18, Tue

$$S(\beta) = (Y - Z\beta)'(Y - Z\beta) = \sum_{i=1}^n (y_i - z_i'\beta)^2$$

$$\beta = (\beta_0, \beta_1, \dots, \beta_r)' = \sum_{i=1}^n (y_i - \beta_0 - \beta_1 z_{i1} - \dots - \beta_r z_{ir})^2$$

$$\begin{bmatrix} \frac{\partial S}{\partial \beta_0} \\ \vdots \\ \frac{\partial S}{\partial \beta_r} \end{bmatrix} = \frac{\partial S}{\partial \beta}$$

Last time, we looked at the least square estimate of β

$$\hat{\beta}_{LS} = (Z'Z)^{-1}Z'Y$$

• Decomposition of sum of squares.

$$\text{Total sum of Square} = \sum_{i=1}^n y_i^2$$

$$\text{Total SS } Y'Y = Y'(I-P_Z)Y = Y'Y - Y'P_Z Y$$

$\sum_{i=1}^n y_i^2 = \sum_{i=1}^n (y_i - \bar{y})^2 + n\bar{y}^2$ (total SS (corrected) or total SS about the mean.)

$$\begin{aligned} \sum_{i=1}^n (y_i - \bar{y})^2 &= \sum_{i=1}^n y_i^2 - 2\sum_{i=1}^n y_i \bar{y} + \sum_{i=1}^n \bar{y}^2 \\ &= \sum_{i=1}^n y_i^2 - n\bar{y}^2 \end{aligned}$$

$$\bar{y} = \frac{1}{n} 1'Y, \quad \frac{1}{n} 1' = \begin{bmatrix} 1 & 1 & \dots & 1 \end{bmatrix}$$

in terms of Matrix?

$$\begin{array}{l} \text{Model} \\ Y = Z\beta + \epsilon \\ \begin{array}{ccc} n \times 1 & \uparrow & n \times 1 \\ & n \times (r+1) & \end{array} \end{array}$$

$$\begin{pmatrix} y_1 - \bar{y} \\ \vdots \\ y_n - \bar{y} \end{pmatrix} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} - \begin{pmatrix} \bar{y} \\ \vdots \\ \bar{y} \end{pmatrix} = \mathbf{y} - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n' \mathbf{y}$$

$$P_{\mathbf{1}_n} = \frac{1}{n} \mathbf{1}_n (\mathbf{1}_n' \mathbf{1}_n)^{-1} \mathbf{1}_n' = \frac{1}{n} \mathbf{1}_n \mathbf{1}_n'$$

$$\text{So, } (I - P_{\mathbf{1}_n})^2 = (I - P_{\mathbf{1}_n})$$

$$\sum_{i=1}^n (y_i - \bar{y})^2 = \begin{pmatrix} y_1 - \bar{y} \\ \vdots \\ y_n - \bar{y} \end{pmatrix}' \begin{pmatrix} y_1 - \bar{y} \\ \vdots \\ y_n - \bar{y} \end{pmatrix}$$

Total SS about the mean

$$= \mathbf{y}' (I - P_{\mathbf{1}_n})^2 \mathbf{y} = \mathbf{y}' (I - P_{\mathbf{1}_n}) \mathbf{y}$$

idempotent

$$\text{Total SS about mean} = \mathbf{y}' (I - P_{\mathbf{1}_n}) \mathbf{y}$$

$$\sum_{i=1}^n (y_i - \bar{y})^2 = \mathbf{y}' (I - P_{\mathbf{1}_n}) \mathbf{y} = \mathbf{y}' (I - P_{\mathbf{1}_n}) \mathbf{y}$$

Residual SS

remains
if bigger → prediction is failed

Regression SS

Model is good or not
if it is bigger, good

$$\mathbf{y}' (I - P_Z) \mathbf{y} = \mathbf{y}' (I - P_Z)^2 \mathbf{y}$$

$$\text{"residual SS"} = \mathbf{y}' (I - P_Z) (I - P_Z) \mathbf{y}$$

$$(I - P_Z) \mathbf{y} = \mathbf{y} - P_Z \mathbf{y} = \mathbf{y} - \hat{\mathbf{y}} \quad (\text{because } P_Z \mathbf{y} = \mathbf{Z} \hat{\beta} = \hat{\mathbf{y}})$$

$$= \hat{\mathbf{e}}' \hat{\mathbf{e}} = \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

residual SS

$$= \sum_{i=1}^n \hat{e}_i^2 = \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 z_{i1} - \dots - \hat{\beta}_r z_{ir})^2$$

"Regression SS"

$$\star \mathbf{y}' (P_Z - P_{\mathbf{1}_n}) \mathbf{y} = ?$$

$$P_Z \mathbf{1}_n = \mathbf{1}_n, \text{ then}$$

$$P_Z P_{\mathbf{1}_n} = P_Z \frac{1}{n} \mathbf{1}_n \mathbf{1}_n' = \frac{1}{n} P_Z \mathbf{1}_n \mathbf{1}_n' = \frac{1}{n} \mathbf{1}_n \mathbf{1}_n' = P_{\mathbf{1}_n}$$

$$\text{So, } (P_Z - P_{\mathbf{1}_n})^2 = P_Z^2 + P_{\mathbf{1}_n}^2 - P_Z P_{\mathbf{1}_n} - P_{\mathbf{1}_n} P_Z$$

$$(P_Z - P_{\mathbf{1}_n}) = P_Z + P_{\mathbf{1}_n} - P_{\mathbf{1}_n} - P_{\mathbf{1}_n}$$

$$\text{"Regression SS"} = (P_Z - P_{\mathbf{1}_n}) \quad \text{"idempotent"}$$

$$\mathbf{y}' (P_Z - P_{\mathbf{1}_n}) \mathbf{y} = \mathbf{y}' (P_Z - P_{\mathbf{1}_n}) (P_Z - P_{\mathbf{1}_n}) \mathbf{y} = \sum_{i=1}^n (\hat{y}_i - \bar{y})^2$$

$$(P_Z - P_{\mathbf{1}_n}) \mathbf{y} = \hat{\mathbf{y}} - \bar{y} \mathbf{1}_n = \begin{pmatrix} \hat{y}_1 - \bar{y} \\ \vdots \\ \hat{y}_n - \bar{y} \end{pmatrix}$$

$$P_Z \mathbf{y} - P_{\mathbf{1}_n} \mathbf{y}$$

In general, they are not equal. But here they are equal. i.e. $P_Z P_{\mathbf{1}_n} = P_{\mathbf{1}_n} P_Z$ "Not commutative"

Also observe $(I - P_Z)(P_Z - P_1) = 0$.

Also, $(I - P_1)P_1 = 0$

So, the decomposition of Y into $(I - P_1)Y$ and P_1Y is an orthogonal decomposition.

Also, $(I - P_1)Y = (I - P_Z)Y + (P_Z - P_1)Y$

This is also orthogonal decomposition.

A measure of goodness of the regression is

$$R^2 = \frac{\text{Reg. SS}}{\text{Tot. SS about Mean}} = \frac{Y'(P_Z - P_1)Y}{Y'(I - P_1)Y}$$

$$= \frac{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}{\sum_{i=1}^n (y_i - \bar{y})^2} = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

$$= 1 - \frac{\text{Res. SS}}{\text{Total SS about mean}}$$

R^2 lies between 0 and 1.

not good fit

good fit

predictor has no information for Y

If R^2 is close to 1, then the regression model is a good fit.

If R^2 is close to 0, then the model is that the predictors can not predict the response well.

$R^2 = 1$ only when $\hat{\epsilon}_i = 0$ for all i .

This means $y_i = \hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 z_{i1} + \dots + \hat{\beta}_r z_{ir}$.

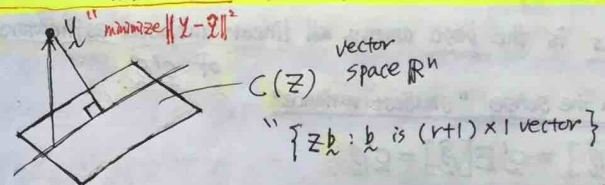
$R^2 = 0$, when $\hat{\beta}_0 = \bar{y}, \hat{\beta}_1 = \dots = \hat{\beta}_r = 0$.

$R = \sqrt{R^2}$ is called Multiple correlation coefficient.

Geometric Point of view for the LS method.

The LS problem:

minimize $\|Y - Z\beta\|^2 = (Y - Z\beta)'(Y - Z\beta)$ with respect to β



So, the minimization problem is equivalent to orthogonal projection of Y onto $C(Z)$.

So, if Zb is the closest to Y among all points of $C(Z)$, then $Y - Zb$ is orthogonal to $C(Z)$.

So, $(Y - Zb)$ is orthogonal to Zc for any $c \in C(Z)$.

So, $c'Z'(Y - Zb) = 0$ for all c .

So, $Z'(Y - Zb) = 0 \Rightarrow b = (Z'Z)^{-1}Z'Y$.

Hence, $Zb = P_Z Y$ by $Zb = Z(Z'Z)^{-1}Z'Y$.

• Best Linear Unbiased Estimator (BLUE)

For any vector $\underline{c}$, ^{column vector} suppose we want to estimate $\underline{c}'\beta$

An estimator will be called linear if it has the form $\underline{a}'Y$ for some $\underline{a}$

Thm (Gauss)

$\underline{c}'\hat{\beta}$ is the best among all linear unbiased estimators of $\underline{c}'\beta$.

Best in the sense "smallest variance"

$$E[\underline{c}'\hat{\beta}] = \underline{c}'E[\hat{\beta}] = \underline{c}'\beta$$

So, $\underline{c}'\hat{\beta}$ is unbiased for $\underline{c}'\beta$. $\underline{a}'Z\hat{\beta} = \underline{a}'\beta$ so $\underline{a}'Z = \underline{c}'$

$$\underline{c}'\hat{\beta} = \underline{c}'(Z'Z)^{-1}Z'Y = \underline{a}'Y, \text{ where } \underline{a} = Z(Z'Z)^{-1}\underline{c}$$

linear. So, $\underline{c}'\hat{\beta}$ is also linear estimator.

If $\underline{a}'Y$ is unbiased for $\underline{c}'\beta$, then $E[\underline{a}'Y] = E[\underline{a}'Z\beta + \underline{a}'\epsilon] = \underline{a}'Z\beta = \underline{c}'\beta$

$$\Rightarrow E[\underline{a}'Y] = \underline{c}'\beta \text{ for all } \beta \quad \underline{a}'Z = \underline{c}'$$

$$\text{So, } \underline{a}'Z\beta = \underline{c}'\beta \text{ for all } \beta \quad \text{So, } \underline{a}'Z = \underline{c}'$$

$$\text{by } (*) \uparrow E[Y] \text{ or } (C - A'Z)\beta = 0, \forall \beta$$

$$\text{Var}(\underline{a}'Y) = \underline{a}'\text{Var}(Y)\underline{a} = \sigma^2 \underline{a}'\underline{a}$$

$$\text{Var}(\underline{c}'\hat{\beta}) = \text{Var}(\underline{a}'Y) = \sigma^2 \underline{a}'\underline{a}$$

$$\begin{aligned} \underline{a}'\underline{a} &= (\underline{a} - \underline{a}_* + \underline{a}_*)'(\underline{a} - \underline{a}_* + \underline{a}_*) \\ &= (\underline{a} - \underline{a}_*)'(\underline{a} - \underline{a}_*) + \underline{a}_*' \underline{a}_* + 2(\underline{a} - \underline{a}_*)'\underline{a}_* \\ &= \underline{a}_*' \underline{a}_* \end{aligned}$$

Since $(\underline{a} - \underline{a}_*)'Z = \underline{c}' - \underline{c}' = 0$
 $(\underline{a} - \underline{a}_*)'\underline{a}_* = \underline{a}'\underline{a}_* - \underline{a}_*' \underline{a}_* = 0$

Math B7800

2/8/18, Thur

$$Y = (I - P_Z)Y + P_Z Y \text{ "orthogonal decomposition"}$$

BLUE (continued)

$$\underline{a}_* = Z(Z'Z)^{-1}\underline{c}$$

$$\underline{a} \text{ satisfies } \underline{a}'Z = \underline{c}'$$

$$\text{Need to check } (\underline{a} - \underline{a}_*)'\underline{a}_* = 0.$$

$$(\underline{a} - \underline{a}_*)'Z(Z'Z)^{-1}\underline{c} = 0, \text{ since } (\underline{a} - \underline{a}_*)'Z = \underline{c}' - \underline{c}' = 0$$

• Inference about Regression

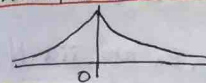
So far we didn't assume any distribution for the error ϵ_i 's.

We only assumed $E(\epsilon_i) = 0$ and $\text{cov}(\epsilon_i, \epsilon_j) = \begin{cases} 0 & \text{if } i \neq j \\ \sigma^2 & \text{if } i = j \end{cases}$

Now, we need to assume distribution of the error.

1. Suppose the errors are iid Double exponential dist.

$$\text{with pdf } f(x) = \frac{1}{2}e^{-|x|}, x \in \mathbb{R}$$



Suppose $\epsilon_1, \dots, \epsilon_n$ are iid with pdf $f(x)$. dist. of Y ?

Q: Find MLE of β, σ^2 .

$$\text{Model } Y = Z\beta + \epsilon$$

Suppose $y_1, y_2, \dots, y_n$ are independent.

$$E(y_i) = Z_i'\beta = \beta_0 + \beta_1 Z_{i1} + \dots + \beta_r Z_{ir}$$

dist. of y_i ?

PDF of y_i ?

PDF of $\varepsilon_i = y_i - z_i \beta$ is $f(x)$ & Double exponential dist.

$$f(y_i) = \frac{1}{2} e^{-|y_i - z_i \beta|}$$

So, joint pdf of $y_1, \dots, y_n$.

$$f(y_1, \dots, y_n) = \prod_{i=1}^n f(y_i) = \frac{1}{2^n} e^{-\sum_{i=1}^n |y_i - z_i \beta|} = L(\beta)$$

"Not continuous!" "likelihood function"

Maximizing $L(\beta)$ is equivalent to

Minimizing $\sum_{i=1}^n |y_i - \beta_0 - \beta_1 z_{i1} - \dots - \beta_r z_{ir}|$ w.r.t β .

The estimator that minimizes the last term is called LAD (Least absolute deviation) estimator of β .

Note: Suppose the errors are normally distributed.

$$\text{So, } \varepsilon_i \sim N(0, \sigma^2 I_n)$$

$$\text{Normal Dist. } f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}$$

So, $\varepsilon_1, \dots, \varepsilon_n$ are iid $N(0, \sigma^2)$

Q: Find MLE of β, σ^2 .

Distribution of y_i is $N(z_i \beta, \sigma^2)$

$$\text{pdf of } y_i \quad f(y_i) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2\sigma^2}(y_i - z_i \beta)^2}$$

Joint pdf of $y_1, \dots, y_n$

$$f(y_1, \dots, y_n) = \left(\frac{1}{\sqrt{2\pi}\sigma}\right)^n e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - z_i \beta)^2} = L(\beta, \sigma^2)$$

$$\ell(\beta, \sigma^2) = \log L(\beta, \sigma^2) = -\frac{n}{2} \log(2\pi) - \frac{n}{2} \log(\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - z_i \beta)^2$$

$$\frac{\partial}{\partial \beta} \ell(\beta, \sigma^2) = -\frac{1}{2\sigma^2} \frac{\partial}{\partial \beta} \sum_{i=1}^n (y_i - z_i \beta)^2 = 0$$

minimizer

$$\text{So, } \hat{\beta} = \hat{\beta}_{OLS} = (Z'Z)^{-1} Z'Y$$

MLE of β ?

$$\frac{\partial}{\partial \sigma^2} \ell(\beta, \sigma^2) = -\frac{n}{2\sigma^2} + \frac{1}{2\sigma^4} \sum_{i=1}^n (y_i - z_i \beta)^2 = 0$$

$$\text{So, } \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (y_i - z_i \hat{\beta})^2 = \frac{1}{n} \sum_{i=1}^n \hat{\varepsilon}_i^2 = \frac{1}{n} \hat{\varepsilon}' \hat{\varepsilon}$$

Hence, $\hat{\beta}, \hat{\sigma}^2$ are the MLE.

$\hat{\sigma}^2$ is MLE, but not unbiased for σ^2 .

The unbiased estimator for σ^2 is $\frac{\hat{\varepsilon}' \hat{\varepsilon}}{n-r-1} = s^2$

$\hat{\beta}$ is unbiased for β .

$$\begin{aligned} \hat{\beta} &= (Z'Z)^{-1} Z'Y \sim N_{r+1} \left(\frac{(Z'Z)^{-1} Z'Z\beta, (Z'Z)^{-1} Z'(\sigma^2 I)}{Z(Z'Z)^{-1}} \right) \\ Y &\sim N_n(Z\beta, \sigma^2 I) \\ &= N_{r+1}(\beta, \sigma^2 (Z'Z)^{-1}) \end{aligned}$$

Fact $\hat{\beta}$ and $\hat{\sigma}^2$ (or s^2) are independent?

We will show $\hat{\beta}$ and $\hat{\epsilon}$ are independent.

$$\hat{\epsilon} = (I - P_Z)Y, \quad \hat{\beta} = (Z'Z)^{-1}Z'Y$$

$$\text{cov}(\hat{\epsilon}, \hat{\beta}) = (I - P_Z)(\sigma^2 I)Z(Z'Z)^{-1} = 0.$$

Since $(I - P_Z)Z = 0$

$$\text{cov}(A_1, B_2) = A \text{cov}(Y, Z) B'$$

So, $\hat{\epsilon}$ and $\hat{\beta}$ are independent.

Since $\hat{\sigma}^2$ and s^2 are both function of $\hat{\epsilon}$ they are independent with $\hat{\beta}$.

Q: Distribution of $\hat{\sigma}^2$ and s^2 ?

Fact: $\hat{\epsilon}'\hat{\epsilon} \sim \chi^2_{n-r-1}$

$$R^2 = \frac{Y'(P_Z - P_0)Y}{Y'(I - P_Z)Y}$$

$$1 - R^2 = \frac{Y'(I - P_Z)Y}{Y'(I - P_0)Y}$$

Linear Algebra Digression

Suppose A is symmetric real matrix $n \times n$.

If $\lambda_1, \lambda_2, \dots, \lambda_n$ are the eigenvalues of A ,

then eigenvalues of A^2 are $\lambda_1^2, \dots, \lambda_n^2$.

If $A = U\Lambda U'$ is the spectral decomposition,

i.e., U is orthogonal and $\Lambda = \begin{pmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{pmatrix}$,

then $A^2 = U\Lambda U'U\Lambda U' = U\Lambda^2 U'$, $\Lambda^2 = \begin{pmatrix} \lambda_1^2 & & 0 \\ & \ddots & \\ 0 & & \lambda_n^2 \end{pmatrix}$

This shows that $\lambda_1^2, \dots, \lambda_n^2$ are the eigenvalues of A^2 .

If A is idempotent, then each eigenvalue is either 0 or 1.

because if $A = U\Lambda U'$, then $A = A^2$

$$\Rightarrow U\Lambda U' = U\Lambda^2 U'$$

$$\Rightarrow \Lambda = \Lambda^2$$

$$\Rightarrow \lambda_i = \lambda_i^2 \quad \forall i \in \mathbb{N}$$

So, λ_i is either 0 or 1.

$$\text{So, } A = U \begin{bmatrix} I & 0 \\ 0 & 0 \end{bmatrix} U', \quad U = [U_1 | U_2]$$

$n \times k \quad n \times (n-k)$

$$= U_1 U_1'$$

$n \times k \quad k \times n$, then $\text{rank}(A) = k$.

$$\hat{\epsilon}'\hat{\epsilon} = Y'(I - P_Z)Y. \quad \text{Since } (I - P_Z) \text{ is idempotent and } \text{rank}(I - P_Z) = n - r - 1$$

$$\text{So, } I - P_Z = UU', \quad \text{where } U \text{ is } n \times (n - r - 1)$$

$$\text{and } U'U = I.$$

$$\text{So, } Y'(I - P_Z)Y = Y'UU'Y$$

$\text{So, } \frac{Y'Y}{\sigma^2} \sim \chi^2_{n-r-1}$

$$\text{Now, see that } U'Y \sim N(U'Z\beta, \sigma^2 I)$$

$\frac{Y'(I - P_Z)Y}{\sigma^2} \sim \chi^2_{n-r-1}$

$$\text{Also, } UU'Z = (I - P_Z)Z = 0$$

$$\text{So, } U'UU'Z = 0$$

$$\text{So, } U'Z = 0$$

$$U'Y \sim N_{n-r-1}(0, \sigma^2 I)$$

Model: $\mathbf{y} = \mathbf{Z}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$, where $\boldsymbol{\varepsilon} \sim N_n(\mathbf{0}, I_n \sigma^2)$.

Last time $\frac{\mathbf{y}'(\mathbf{I} - \mathbf{P}_Z)\mathbf{y}}{\sigma^2} \sim \chi^2_{n-r-1}$ d.f. rank of $(\mathbf{I} - \mathbf{P}_Z)$.

Also, $\hat{\boldsymbol{\beta}}$ and $\mathbf{y}'(\mathbf{I} - \mathbf{P}_Z)\mathbf{y}$ are independent.

★ Goal: Obtain Confidence Region for $\boldsymbol{\beta}$. Then find Scheffe CI and Bonferroni CI for β_i .

Recall: $\hat{\boldsymbol{\beta}} = (\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{y} \sim N_{r+1}(\boldsymbol{\beta}, \sigma^2(\mathbf{Z}'\mathbf{Z})^{-1})$

Recall: For symmetric positive definite matrix $\mathbf{A}$, we can define

$\mathbf{Q}$: Confidence Region for $\boldsymbol{\beta}$.

$\mathbf{A}^{1/2}$

So, $(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}) \sim N_{r+1}(\mathbf{0}, \sigma^2(\mathbf{Z}'\mathbf{Z})^{-1})$.

$\frac{1}{\sigma}(\mathbf{Z}'\mathbf{Z})^{1/2}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}) \sim N_{r+1}(\mathbf{0}, \mathbf{I}_{r+1})$

So, $\left[\frac{1}{\sigma}(\mathbf{Z}'\mathbf{Z})^{1/2}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}) \right]' \frac{1}{\sigma}(\mathbf{Z}'\mathbf{Z})^{1/2}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}) = \mathbf{u}'\mathbf{u} \sim \chi^2_{r+1}$.

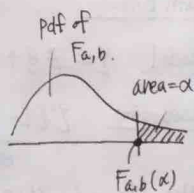
So, using independent of $\hat{\boldsymbol{\beta}}$ and $\mathbf{y}'(\mathbf{I} - \mathbf{P}_Z)\mathbf{y}$.

★ $F = \frac{\frac{1}{r+1} \frac{1}{\sigma^2} (\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})' (\mathbf{Z}'\mathbf{Z}) (\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})}{\frac{1}{n-r-1} \frac{1}{\sigma^2} \mathbf{y}'(\mathbf{I} - \mathbf{P}_Z)\mathbf{y}} \sim F_{r+1, n-r-1}$

$= \frac{1}{r+1} \frac{1}{\sigma^2} (\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})' (\mathbf{Z}'\mathbf{Z}) (\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})$

Notation: $F_{r+1, n-r-1}(\alpha)$.

$F_{a,b}(\alpha)$



So, $P(F \leq F_{r+1, n-r-1}(\alpha)) = 1 - \alpha$.

Thus, $\{\beta : (\hat{\beta} - \beta)'(Z'Z)(\hat{\beta} - \beta) \leq (r+1)s^2 F_{r+1, n-r-1}(\alpha)\}$

is a $100(1-\alpha)\%$ Confidence Region for β .

Next topic is CI for β_i : "only i, one interval."

Distribution of $\hat{\beta}_i$?

$\hat{\beta}_i \sim N(\beta_i, \sigma^2(Z'Z)^{-1}_{i+1, i+1})$ for $i=0, 1, \dots, r$.

Note that $(Z'Z)^{-1}_{i+1, i+1}$ refers to the $(i+1), (i+1)$ entry of $(Z'Z)^{-1}$.

Now, $\frac{\hat{\beta}_i - \beta_i}{\sqrt{\sigma^2(Z'Z)^{-1}_{i+1, i+1}}} \sim N(0, 1)$.

Since σ^2 is unknown, we cannot use it for a CI of β_i . We replace σ^2 by s^2 .

$\frac{\hat{\beta}_i - \beta_i}{\sqrt{s^2(Z'Z)^{-1}_{i+1, i+1}}} \sim t_{n-r-1}$.

$= \frac{(\hat{\beta}_i - \beta_i)\sqrt{s^2(Z'Z)^{-1}_{i+1, i+1}}}{\sqrt{s^2/2}} \sim N(0, 1)$

$$\frac{s^2}{\sigma^2} = \frac{1}{n-r-1} \frac{Y'(I-P_e)Y}{\sigma^2} \sim \chi^2_{n-r-1}$$

So, a $100(1-\alpha)\%$ CI for β_i is

$$\hat{\beta}_i \pm \sqrt{s^2(Z'Z)^{-1}_{i+1, i+1}} t_{n-r-1}\left(\frac{\alpha}{2}\right)$$

$$\hat{\beta}_i \pm \sqrt{\widehat{\text{var}}(\hat{\beta}_i)} t_{n-r-1}\left(\frac{\alpha}{2}\right)$$

$$\widehat{\text{var}}(\hat{\beta}_i) = \sigma^2(Z'Z)^{-1}_{i+1, i+1}$$

$$\widehat{\text{var}}(\hat{\beta}_i) = s^2(Z'Z)^{-1}_{i+1, i+1}$$

• Simultaneous CI for $\beta_i, i=0, 1, \dots, k$.

Two standard methods

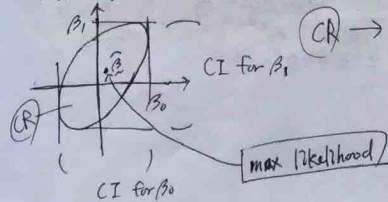
(1) Bonferroni, (2) Scheffe's method.

(1) Bonferroni simultaneous CI for $\beta_0, \beta_1, \dots, \beta_k$ with confidence $1-\alpha$ is

$$\hat{\beta}_i \pm \sqrt{\widehat{\text{var}}(\hat{\beta}_i)} t_{n-r-1}\left(\frac{\alpha}{2(k+1)}\right)$$

(2) Scheffe's method, we need to project the $(1-\alpha)100\%$ CR onto different axes.

Suppose β has 2 components.



CR $\rightarrow$ CIs for β_1, β_2 .

Diagonalization (Linear Algebra)

Suppose A is a symmetric matrix.

$$\max_{\mathbf{x}} \mathbf{x}' A \mathbf{x} = \lambda_1(A) \quad (\text{largest eigenvalue of } A)$$

$$: \|\mathbf{x}\|_2 = 1$$

$$= \max_{\mathbf{x}: \|\mathbf{x}\|_2 \leq 1} \mathbf{x}' A \mathbf{x} = \lambda_1(A).$$

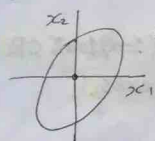
positive definite.

Suppose A is symmetric, B is symmetric and pd.

$$\max_{\mathbf{x}: \mathbf{x}' A \mathbf{x} = 1} (\mathbf{x}' A \mathbf{x}) = \max_{\mathbf{y}: \|\mathbf{y}\|_2 \leq 1} \mathbf{y}' B^{-1/2} A B^{-1/2} \mathbf{y}$$

$$\text{Define } \mathbf{y} = B^{1/2} \mathbf{x} \quad = \lambda_1(B^{-1/2} A B^{-1/2})$$

Next, for two matrices C and D for which CD, DC are both defined, then all nonzero eigenvalues of CD and DC are same.



$$\mathbf{x}' B \mathbf{x} \leq 1$$

Finding the projection onto x_1 -axis.

$$\Leftrightarrow \max_{\mathbf{x}: \mathbf{x}' B \mathbf{x} \leq 1} x_1^2$$

$$\max_{\mathbf{x}: \mathbf{x}' B \mathbf{x} \leq 1} x_1^2$$

Finding the projection onto x_1 -axis.

$$\Leftrightarrow \max_{\mathbf{x}: \mathbf{x}' B \mathbf{x} \leq 1} x_1^2$$

$$= \max_{\mathbf{x}: \mathbf{x}' B \mathbf{x} \leq 1} \mathbf{x}' A \mathbf{x}, \text{ where } A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \begin{pmatrix} 1 & 0 \end{pmatrix}$$

$$= \mathbf{e}_1 \mathbf{e}_1', \text{ where } \mathbf{e}_k \text{ is } k \text{th unit vector.}$$

$$= \lambda_1(B^{-1/2} A B^{-1/2})$$

$$= \lambda_1(B^{-1/2} A B^{-1/2})$$

$$= \lambda_1(B^{-1/2} \mathbf{e}_1 \mathbf{e}_1' B^{-1/2})$$

$$= \lambda_1(\mathbf{e}_1' B^{-1} \mathbf{e}_1) = \mathbf{e}_1' B^{-1} \mathbf{e}_1 = (B^{-1})_{1,1}$$

↑ scalar.

So, we see that

$$\max_{\mathbf{x}: \mathbf{x}' B \mathbf{x} \leq 1} x_i^2 = (B^{-1})_{i,i}$$

Recall 100(1- α)% CR for β was

$$(\hat{\beta} - \beta)' (Z'Z)^{-1} (\hat{\beta} - \beta) \leq (r+1) s^2 F_{r+1, n-r-1}(\alpha)$$

Write in the term $\mathbf{x}' B \mathbf{x} \leq 1$, where

$$\mathbf{x} = \hat{\beta} - \beta, \quad B = (Z'Z)^{-1} \frac{1}{(r+1) s^2 F_{r+1, n-r-1}(\alpha)}$$

$$\max_{\beta_i} |\beta_i - \hat{\beta}_i|^2 = (B^{-1})_{i+1, i+1}, i=0, 1, \dots, r$$

Q: Is MCR.

$$= (r+1) s^2 F_{r+1, n-r-1}(\alpha) (Z'Z)^{-1}_{i+1, i+1}$$

so, the Scheffe simultaneous CI for $\beta_0, \dots, \beta_r$ is

$$\hat{\beta}_i \pm \sqrt{(r+1) s^2 F_{r+1, n-r-1}(\alpha) (Z'Z)^{-1}_{i+1, i+1}}$$

$$= \hat{\beta}_i \pm \sqrt{(r+1) F_{r+1, n-r-1}(\alpha) \widehat{\text{var}}(\hat{\beta}_i)}$$

Scheffe simultaneous CI for all linear combinations of the components of β , namely $a'\beta$ for all vectors a is given by

$$a'\hat{\beta} \pm \sqrt{(r+1) F_{r+1, n-r-1}(\alpha) \widehat{\text{var}}(a'\hat{\beta})}, \text{ where}$$

$$\widehat{\text{var}}(a'\hat{\beta}) = a' \widehat{\text{var}}(\hat{\beta}) a$$

$$= s^2 a' (Z'Z)^{-1} a$$

$$P_Z = Z(Z'Z)^{-1}Z'$$

$$P_{Z_1} = Z_1(Z_1'Z_1)^{-1}Z_1'$$

$$Z = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$$

$$Z_1 = \begin{bmatrix} 1 & 0 \\ 3 & 0 \end{bmatrix}$$

Math B/800

2/15/18, Thur

March 06, Tue - Exam 01. (3 pages formula sheet)

Likelihood Ratio test

$$\underset{(r+1) \times 1}{\beta} = \begin{bmatrix} \beta_1 \\ \vdots \\ \beta_2 \end{bmatrix} \begin{matrix} (q+1) \times 1 \\ \\ (r-q) \times 1 \end{matrix} \quad \underset{n \times (r+1)}{Z} = \begin{bmatrix} Z_1 & | & Z_2 \end{bmatrix} \begin{matrix} n \times (q+1) & n \times (r-q) \end{matrix}$$

$$\text{Then } y = Z\beta + \varepsilon = [Z_1 | Z_2] \begin{bmatrix} \beta_1 \\ \beta_2 \end{bmatrix} + \varepsilon$$

$$\begin{bmatrix} 1 & 10 & 11 \\ 1 & 12 & 13 \\ 1 & 14 & 15 \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \\ \beta_2 \end{bmatrix}$$

$$= Z_1 \beta_1 + Z_2 \beta_2 + \varepsilon$$

$n \times 1 \quad n \times 1 \quad n \times 1$

"Test"

$$H_0: \beta_2 = 0 \text{ vs } H_1: \beta_2 \neq 0$$

Recall The likelihood ratio test (LRT).

$$\Lambda = \frac{\sup_{H_0} L(\beta, \sigma^2)}{\sup L(\beta, \sigma^2)} \text{ is the likelihood Ratio.}$$

We reject H_0 if Λ is small.

We express Λ as a monotone function of a statistic T (of the data) whose null distribution is known (under H_0).

Then $\{\Lambda < c\} \iff \{T < c' \text{ or } > c'\}$ for any constant c .

Then, for any constant c , there is another constant c' s.t.

$$\begin{cases} \Lambda < c \} \Leftrightarrow \{ T < c' \} \text{ if } \Lambda \text{ is increasing} \\ \Leftrightarrow \{ T > c' \} \text{ if } \Lambda \text{ is decreasing.} \end{cases}$$

Then we will reject H_0 if $\begin{cases} T < c' & \text{in case 1} \\ T > c' & \text{in case 2.} \end{cases}$

Since the null distribution of T will be known, we can find c' so that the "level condition" is met.

In our case,

$$L(\beta, \sigma^2) = (2\pi)^{-n/2} (\sigma^2)^{-n/2} \exp\left[-\frac{1}{2\sigma^2} (Y - Z\beta)'(Y - Z\beta)\right]$$

The maximizer of $L(\beta, \sigma^2)$ are

$$\hat{\beta} = (Z'Z)^{-1}Z'Y, \quad \hat{\sigma}^2 = \frac{(Y - Z\hat{\beta})'(Y - Z\hat{\beta})}{n}$$

$$\sup L(\beta, \sigma^2) = L(\hat{\beta}, \hat{\sigma}^2) = (2\pi)^{-n/2} (\hat{\sigma}^2)^{-n/2} e^{-n/2}$$

If H_0 is true, then $Y = Z_1\beta_1 + \varepsilon$.

$$L\left(\begin{pmatrix} \beta_1 \\ 0 \end{pmatrix}, \sigma^2\right) = (2\pi)^{-n/2} (\sigma^2)^{-n/2} \exp\left[-\frac{1}{2\sigma^2} (Y - Z_1\beta_1)'(Y - Z_1\beta_1)\right]$$

The maximizer of $L(\beta_1, \sigma^2)$ are

$$\hat{\beta}_1 = (Z_1'Z_1)^{-1}Z_1'Y, \quad \hat{\sigma}_1^2 = \frac{(Y - Z_1\hat{\beta}_1)'(Y - Z_1\hat{\beta}_1)}{n}$$

$$\text{so, } \sup_{H_0} L(\beta, \sigma^2) = L\left(\begin{pmatrix} \hat{\beta}_1 \\ 0 \end{pmatrix}, \hat{\sigma}_1^2\right)$$

$$= (2\pi)^{-n/2} (\hat{\sigma}_1^2)^{-n/2} e^{-n/2}$$

$$\text{so, } \Lambda = \left(\frac{\hat{\sigma}_1^2}{\hat{\sigma}^2}\right)^{-n/2} = \left(\frac{Y'(I - P_{Z_1})Y}{Y'(I - P_Z)Y}\right)^{-n/2}$$

"decreasing"

so, we will reject H_0 if $\frac{Y'(I - P_{Z_1})Y}{Y'(I - P_Z)Y} > c$, for some c .

Recall that $(I - P_Z), (I - P_{Z_1})$ are idempotent.

$$\text{so, } \frac{Y'(I - P_Z)Y}{\sigma^2} \sim \chi_{n-1}^2$$

"not independent"

$$\text{Under } H_0, \quad \frac{Y'(I - P_{Z_1})Y}{\sigma^2} \sim \chi_{n-q-1}^2 \quad \text{rank}(Z_1) = q+1$$

These two χ^2 will be independent if $(I - P_Z)(I - P_{Z_1}) = 0$.

$$= I - P_Z - P_{Z_1} + P_{Z_1} \neq 0$$

"Not independent" (as $P_{Z_1} \cdot Z_1 = Z_1$)

But,

$$\frac{Y'(I - P_{Z_1})Y}{Y'(I - P_Z)Y} - 1 = \frac{Y'(P_Z - P_{Z_1})Y}{Y'(I - P_Z)Y}$$

independent

$$\text{Check } (P_Z - P_{Z_1})(I - P_Z) = P_Z - P_Z^2 - P_{Z_1} + P_{Z_1} = 0.$$

Under H_0 ,

$$\frac{\frac{1}{n-q} Y'(P_Z - P_{Z_1})Y / \sigma^2}{\frac{1}{n-r-1} Y'(I - P_Z)Y / \sigma^2} \sim F_{r-q, n-r-1}$$

so, $\frac{1}{r-q} \frac{1}{s^2} Y'(P_Z - P_{Z_1})Y \sim F_{r-q, n-r-1}$

so, we reject H_0 at level of α if

$Y'(P_Z - P_{Z_1})Y > (r-q)s^2 F_{r-q, n-r-1}(\alpha)$

test Δ , decreasing, $\gg \text{finv}(1-\alpha, r-q, n-r-1)$

Confidence Region for β_2

$\beta = \begin{bmatrix} \beta_1 \\ \beta_2 \end{bmatrix}$ want a $100(1-\alpha)\%$ CR for β_2

Recall $\hat{\beta} \sim N_{r+1}(\beta, \sigma^2(Z'Z)^{-1})$. Then $\hat{\beta}_2 \sim ?$

$\hat{\beta}_2 \sim N_{r-q}(\beta_2, \sigma^2(Z_2'(I - P_{Z_1})Z_2)^{-1})$

$Z = [Z_1 | Z_2]$

$Z'Z = \begin{bmatrix} Z_1' \\ Z_2' \end{bmatrix} [Z_1, Z_2] = \begin{bmatrix} Z_1'Z_1 & Z_1'Z_2 \\ Z_2'Z_1 & Z_2'Z_2 \end{bmatrix}$

Inverse of block matrix $\begin{bmatrix} A & B \\ C & D \end{bmatrix}^{-1} = \begin{bmatrix} * & * \\ * & (D - CA^{-1}B)^{-1} \end{bmatrix}$

$(Z'Z)^{-1} = \begin{bmatrix} * & * \\ * & (Z_2'Z_2 - Z_2'Z_1(Z_1'Z_1)^{-1}Z_1'Z_2)^{-1} \end{bmatrix} (Z_2'(I - P_{Z_1})Z_2)^{-1}$

$\frac{(\hat{\beta}_2 - \beta_2)' Z_2' (I - P_{Z_1}) Z_2 (\hat{\beta}_2 - \beta_2)}{\sigma^2} \sim \chi^2_{r-q}$

Then $\frac{1}{r-q} \frac{1}{s^2} (\hat{\beta}_2 - \beta_2)' Z_2' (I - P_{Z_1}) Z_2 (\hat{\beta}_2 - \beta_2) \sim F_{r-q, n-r-1}$

so, a $100(1-\alpha)\%$ CR for β_2 is

$\{ \beta_2 : (\hat{\beta}_2 - \beta_2)' Z_2' (I - P_{Z_1}) Z_2 (\hat{\beta}_2 - \beta_2) \leq (r-q)s^2 F_{r-q, n-r-1}(\alpha) \}$

H_0 being belong to in this ellipse.
is true

Prediction using the fitted regression model

We have fitted a linear regression model on

$Y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}$, and $Z = \begin{bmatrix} 1 & z_{11} & \dots & z_{1r} \\ \vdots & \vdots & & \vdots \\ 1 & z_{n1} & \dots & z_{nr} \end{bmatrix}$

For a future observations with known predictor variables.

(say $Z_0 = [1 \ z_{01} \dots \ z_{0r}]'$) how to obtain point and interval estimates of the corresponding (1) response and (2) mean response.

Predictor

$Z_0 = \begin{bmatrix} 1 \\ z_{01} \\ \vdots \\ z_{0r} \end{bmatrix}$

If y_0 is the corresponding response, then

$$y_0 = \beta_0 + \beta_1 z_{01} + \dots + \beta_r z_{0r} + \varepsilon_0 = \mathbf{z}_0' \boldsymbol{\beta} + \varepsilon_0$$

$$= \mathbf{z}_0' \boldsymbol{\beta} + \varepsilon_0$$

$(1 \ z_{01} \ \dots \ z_{0r}) \begin{pmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_r \end{pmatrix} + \varepsilon_0$

$E[y_0] = \mathbf{z}_0' \boldsymbol{\beta}$. This is a parameter. We estimate this by

$\mathbf{z}_0' \hat{\boldsymbol{\beta}}$ (This is the BLUE). CI for $\mathbf{z}_0' \boldsymbol{\beta}$ is

$$\mathbf{z}_0' \hat{\boldsymbol{\beta}} \pm t_{n-r-1} \left(\frac{\alpha}{2} \right) \sqrt{s^2 \mathbf{z}_0' (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{z}_0}$$

this is for an average ~~est~~ prediction.

To predict y_0 , the point estimate $\mathbf{z}_0' \hat{\boldsymbol{\beta}}$ is unbiased.

Interval estimate here is called prediction interval.

$$\text{Var}(\hat{y}_0) = \text{Var}(\mathbf{z}_0' \hat{\boldsymbol{\beta}}) + \sigma^2$$

Prediction interval: $\mathbf{z}_0' \hat{\boldsymbol{\beta}} \pm t_{n-r-1} \left(\frac{\alpha}{2} \right) \sqrt{s^2 + s^2 \mathbf{z}_0' (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{z}_0}$

Math B7800

2/22/18, Thur

Prediction Interval

Set up We have our regression model $y = \mathbf{z}'\boldsymbol{\beta} + \varepsilon$

$\begin{matrix} \text{number} & & \text{parameters} \\ \downarrow & & \downarrow \\ y & = & \mathbf{z}'\boldsymbol{\beta} + \varepsilon \\ \text{response} & & \text{predictors} \quad \text{error} \end{matrix}$

We observe values of y and $\mathbf{z}$ for n observations.

We arrange them to have

$$\mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix}, \quad \mathbf{Z} = \begin{pmatrix} \mathbf{z}_1' \\ \vdots \\ \mathbf{z}_n' \end{pmatrix} = \begin{bmatrix} 1 & z_{11} & \dots & z_{1r} \\ \vdots & \vdots & & \vdots \\ 1 & z_{n1} & \dots & z_{nr} \end{bmatrix}$$

Using these we estimate $\boldsymbol{\beta}$: $\hat{\boldsymbol{\beta}} = (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{Z}'\mathbf{y}$

Now, suppose we have a future observation for which the predictor variables are known.

But, we do not know the value of the response.

$$y_0 = \mathbf{z}_0' \boldsymbol{\beta} + \varepsilon_0$$

$\begin{matrix} \text{same as earlier} \\ \uparrow \\ y_0 & = & \mathbf{z}_0' \boldsymbol{\beta} + \varepsilon_0 \\ \text{response for future obs.} & & \text{predictor for future obs.} \end{matrix}$

$E(y_0) = \mathbf{z}_0' \boldsymbol{\beta}$ This is a parameter, so we can find estimate and CI for it.

mean responses.

Prediction Interval

Goal: to find an interval that will contain y_0 (response of a future observation) with a given confidence.

Consider $y_0 - \mathbf{z}_0' \hat{\beta}$ and find its distribution.

Recall, under normality assumption, $\hat{\beta} \sim N_{p+1}(\beta, \sigma^2(\mathbf{Z}'\mathbf{Z})^{-1})$.

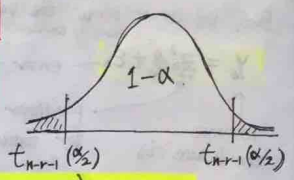
Then ind. $\mathbf{z}_0' \hat{\beta} \sim N(\mathbf{z}_0' \beta, \sigma^2 \mathbf{z}_0' (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{z}_0)$

$y_0 \sim N(\mathbf{z}_0' \beta, \sigma^2)$ variance from ϵ_0 (future error)

so, $y_0 - \mathbf{z}_0' \hat{\beta} \sim N(0, \sigma^2(1 + \mathbf{z}_0' (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{z}_0))$

$\frac{y_0 - \mathbf{z}_0' \hat{\beta}}{\sqrt{\sigma^2(1 + \mathbf{z}_0' (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{z}_0)}} \sim N(0, 1)$

$\frac{y_0 - \mathbf{z}_0' \hat{\beta}}{\sqrt{s^2(1 + \mathbf{z}_0' (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{z}_0)}} \sim t_{n-r-1}$



so, $P\left(\left| \frac{y_0 - \mathbf{z}_0' \hat{\beta}}{\sqrt{s^2(1 + \mathbf{z}_0' (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{z}_0)}} \right| \leq t_{n-r-1}\left(\frac{\alpha}{2}\right)\right) = 1 - \alpha$

so, 100(1-alpha)% PI for y_0 .

$\mathbf{z}_0' \hat{\beta} \pm \sqrt{s^2(1 + \mathbf{z}_0' (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{z}_0)} t_{n-r-1}\left(\frac{\alpha}{2}\right)$

Model verification

Recall, $\hat{\epsilon}_j = y_j - \hat{y}_j$

$\hat{\epsilon} = \begin{bmatrix} \hat{\epsilon}_1 \\ \vdots \\ \hat{\epsilon}_n \end{bmatrix} = \mathbf{y} - \hat{\mathbf{y}} = (\mathbf{I} - \mathbf{P}_\mathbf{Z}) \mathbf{y}$

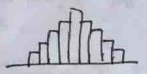
$\text{Var}(\hat{\epsilon}) = \sigma^2(\mathbf{I} - \mathbf{P}_\mathbf{Z})$

Normalized error estimate

$\hat{\epsilon}_j^x = \frac{\hat{\epsilon}_j}{\sqrt{\text{Var}(\hat{\epsilon}_j)}} = \frac{\hat{\epsilon}_j}{s^2(\mathbf{I} - \mathbf{P}_\mathbf{Z})_{jj}}$ normalized variance.

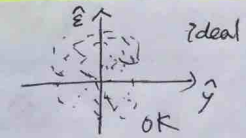
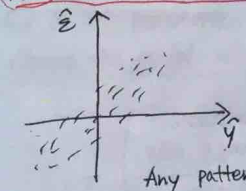
Residual Plot

If we plot the histogram of $\hat{\epsilon}_1^x, \dots, \hat{\epsilon}_n^x$, then we expect the histogram to be close to that of $N(0, 1)$



Violations indicate violations of normality assumption on

Plot $(\hat{y}_1, \hat{\epsilon}_1), \dots, (\hat{y}_n, \hat{\epsilon}_n)$



Any pattern in this plot indicates that the Lin. model assumption is not true.

Plot $\hat{\epsilon}_i$ against $Z_{*2}, \dots, Z_{*(r+1)}$

Expect no pattern if assumption are right.

But there is pattern,
that indicates

violation of the linear model
namely corresponding

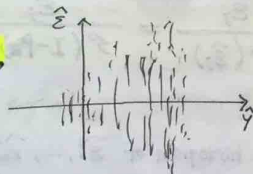
Z variable needs to be transformed before using.

~~violation plot~~

In the plot $\hat{\epsilon}_i$ vs $\hat{x}_i$,

such a pattern
violates

the "equal variance"
assumption.



Q-Q plot

Q-Q plot for the normalized residuals

To check "normality assumption".

Model selections

When the number of predictors is large, we need to select
"suitable small" subset of predictors.

First, naive approach is to consider all subsets
and choose the one which maximizes R^2 (rss)

This doesn't work, as R^2 is an increasing function of r
(# of predictors).

Adjusted R^2 , $\bar{R}^2 = 1 - (1 - R^2) \frac{n-1}{n-r-1}$

One method is to use $\bar{R}^2$ to select model.

A "better" measure is Mallows's

$$C_p = \frac{\text{residual SS using } p \text{ predictors}}{\text{residual variance } (s^2)} - (n-2p)$$

$p = \#$ variables including $\frac{1}{2}$
 $n = \#$ observations.

C_p is not monotone.

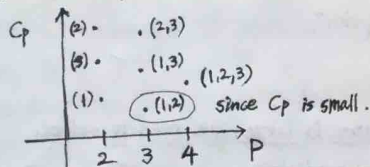
Choose the model with minimum C_p .

$$C_p = \left(\frac{\text{residual sum of squares for subset model with } p \text{ parameters, including an intercept}}{\text{residual variance for full model (residual SS)}} \right) - (n-2p)$$

$p = r+1$

A plot of the pairs (p, C_p) , one for each subset of predictors.

Ex Suppose $Y = \beta_0 + \beta_1 Z_1 + \beta_2 Z_2 + \beta_3 Z_3 + \epsilon$



These two methods work when # predictors is moderate, as we need to consider all possible subsets.

Next, we will see step-wise regression.

Math 87800

2/27/18, Tue

Predictor selection

Stepwise selection method

Recall

$$\text{If } \beta = \begin{pmatrix} \beta_{(1)} \\ \beta_{(2)} \end{pmatrix} \begin{matrix} (q+1) \times 1 \\ (r-q) \times 1 \end{matrix}$$

Test $H_0: \beta_{(2)} = 0$ vs $H_1: \beta_{(2)} \neq 0$

"likelihood ratio test"

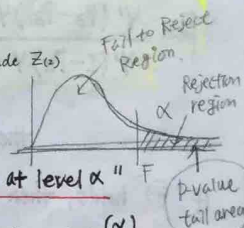
$$Z = \begin{bmatrix} Z_{(1)} & Z_{(2)} \end{bmatrix} \begin{matrix} \text{if } H_0, \text{ exclude } Z_{(2)} \\ \text{if } H_1, \text{ we need to include } Z_{(2)} \end{matrix}$$

$$\begin{matrix} n \times (q+1) & n \times (r-q) \end{matrix}$$

we use the LRT which boils down to

$$F = \frac{Y'(P_Z - P_{Z_1})Y \frac{1}{r-q}}{Y'(I - P_Z)Y \frac{1}{n-r-1}}$$

"We reject H_0 at level α " when $F > F_{r-q, n-r-1}(\alpha)$.



We will refer to the model: $Y = \beta_0 + \beta_1 Z_1 + \dots + \beta_r Z_r + \epsilon$ as "regression function".

Stepwise method

Step 1: Scan through all the variables (one at time)

and which contributes the most to the Regression SS.

$Y = \beta_0 + \beta_1 Z_1$ Find which one has R^2 value the most.

$Y = \beta_0 + \beta_2 Z_2$ compute $\frac{Y'(P_Z - P_{Z_1})Y}{Y'(I - P_Z)Y}$ * use AIC, R^2 , C_p

$$Y = \beta_0 + \beta_r Z_r \rightarrow Z = \begin{bmatrix} 1 & Z_{1r} \\ \vdots & \vdots \\ 1 & Z_{nr} \end{bmatrix}$$

Suppose k is the one which maximizes

Step 2 Test $\beta_k = \begin{bmatrix} \beta_0 \\ \beta_k \end{bmatrix}$

Test $H_0: \beta_k = 0$ vs $H_1: \beta_k \neq 0$

$$Z = \begin{bmatrix} 1 & z_{1k} \\ \vdots & \vdots \\ 1 & z_{nk} \end{bmatrix} \quad Z_1 = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix}$$

$$F = \frac{Y'(P_Z - P_{Z_1})Y \frac{1}{1-0}}{Y'(I - P_Z)Y \frac{1}{n-1-1}} > F_{1, n-2}(\alpha) \quad (*)$$

$\downarrow$ reject $H_0 \Rightarrow \beta_k \neq 0$
then include k in the model.

If $(*)$ holds, then "include" (Z_k) in the regression function.

Current model is $Y = \beta_0 + \beta_k Z_k + \epsilon$.

Step 3 We scan through the remaining variables (one at time) and find out which one contributes to the Reg. SS the most (keeping the existing regression function)

$$Y = \beta_0 + \beta_k Z_k + \beta_1 Z_1 \rightarrow Z = \begin{bmatrix} 1 & z_{1k} & z_{11} \\ \vdots & \vdots & \vdots \\ 1 & z_{nk} & z_{n1} \end{bmatrix}, \quad Z_1 = \begin{bmatrix} 1 & z_{1k} \\ \vdots & \vdots \\ 1 & z_{nk} \end{bmatrix}$$

$$Y = \beta_0 + \beta_k Z_k + \beta_{k-1} Z_{k-1}$$

$$Y = \beta_0 + \beta_k Z_k + \beta_{k+1} Z_{k+1} \rightarrow Z = \begin{bmatrix} 1 & z_{1k} & z_{1r} \\ \vdots & \vdots & \vdots \\ 1 & z_{nk} & z_{nr} \end{bmatrix}$$

$$Y = \beta_0 + \beta_k Z_k + \beta_r Z_r$$

Find out which model has "max value" of $Y'(P_Z - P_1)Y$.

Suppose L is the one that maximizes: reg. SS

$$\text{Then } \beta = \begin{bmatrix} \beta_0 \\ \beta_k \\ \beta_L \end{bmatrix}$$

$H_0: \beta_L = 0$ vs $H_1: \beta_L \neq 0$

$$F = \frac{Y'(P_Z - P_{Z_1})Y \frac{1}{2-1}}{Y'(I - P_Z)Y \frac{1}{n-2-1}} \quad \begin{matrix} v=2 \\ k=1 \\ q=1 \end{matrix}$$

$$Z = \begin{bmatrix} 1 & z_{1k} & z_{1L} \\ \vdots & \vdots & \vdots \\ 1 & z_{nk} & z_{nL} \end{bmatrix}, \quad Z_1 = \begin{bmatrix} 1 & z_{1k} \\ \vdots & \vdots \\ 1 & z_{nk} \end{bmatrix}$$

Include (Z_L) if $F > F_{1, n-3}(\alpha)$

If included, then current model is

$$Y = \beta_0 + \beta_k Z_k + \beta_L Z_L + \epsilon$$

Step 4 When a new variable is included, we need to check whether one of the existing ones fail the F test or not.

Suppose at a stage, we include Z_{im} to $(Z_{i1}, Z_{i2}, \dots, Z_{im-1})$

then test:

$$\beta_0 = \begin{bmatrix} \beta_0 \\ \beta_{i1} \\ \vdots \\ \beta_{im} \end{bmatrix}$$

$$H_0: \beta_{i1} = 0 \text{ vs } H_1: \beta_{i1} \neq 0 \quad \parallel \quad Z_k \quad Z_L \quad \dots$$

$$H_0: \beta_{im-1} = 0 \text{ vs } H_1: \beta_{im-1} \neq 0$$

"Repeat Steps 3, 4"

If in one of the tests, the H_0 is accepted, then the corresponding variable goes out.

$$\beta_k = 0$$

At the end, we will have a subset $\{i_1, \dots, i_p\}$ of $\{1, \dots, r\}$ such that no additional variable can be included,

(i.e., F test is insignificant) and no existing variable leaves (i.e., the corresponding F test is significant).

Then we stop.

"don't drop."

reject H_0

⇒ include variable.

the corresponding

$$AIC = n \ln \left(\frac{\text{residual SS for subset model with } p \text{ parameters, including an intercept}}{n} \right) + 2p$$

overall, we want to select models from those having the smaller values of AIC.

Information Criterion Method

The famous information criterion assigned to subsets of variables is the Akaike's Information on Criterion (AIC).

For a subset of size p ,

$$AIC = n \log \left(\frac{\text{Residual SS for that subset}}{n} \right) + 2p$$

decreasing increasing

one needs to minimize AIC to select variables.

← Multiplied

Multi-collinearity

So, far we have assumed that Z has full rank.

then $Z'Z$ was nonsingular.

If this condition does not hold, then we say that the data has Multi-collinearity.

In some cases, $Z\alpha = \epsilon$ for some $\alpha \neq 0$.

In this case, $Z'Z$ is singular.

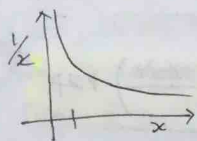
In many cases, $Z'Z$ is "nearly singular", meaning

it has some "very small" eigenvalues.

In the first case, we choose a linearly independent subset of columns of Z .

In the second case, $(Z'Z)^{-1}$ becomes numerically unstable.

e.g.



$$0.01 \rightarrow 100$$

$$0.02 \rightarrow 50$$

Then remove one of the highly correlated columns of Z .

exam 1.

① Multivariate "Multiple" Regression.

3/1/18, Thur

For each individual, we have a vector of response $(1 \times m)$

The model

$$y' = [y_1, \dots, y_m] = [1, z_1, \dots, z_r] \begin{bmatrix} \beta_{01} & \dots & \beta_{0m} \\ \vdots & \ddots & \vdots \\ \beta_{r1} & \dots & \beta_{rm} \end{bmatrix} + \varepsilon' \quad \begin{matrix} \text{"m-responses"} \\ \text{(multiple responses)} \end{matrix} \quad \begin{matrix} 1 \times (r+1) \\ \\ (r+1) \times m \end{matrix} \quad \begin{matrix} 1 \times m \end{matrix}$$

$$E[\varepsilon] = \mathbf{0}_{m \times 1}, \quad E[\varepsilon \varepsilon'] = \Sigma_{m \times m}$$

We observe the responses and predictors from n individuals.

values of

$$y'_i = z'_i \beta + \varepsilon'_i$$

$$\vdots$$

$$y'_n = z'_n \beta + \varepsilon'_n$$

We organize them:

$$Y = \begin{bmatrix} y'_1 \\ \vdots \\ y'_n \end{bmatrix}_{n \times m}, \quad Z = \begin{bmatrix} z'_1 \\ \vdots \\ z'_n \end{bmatrix}_{n \times (r+1)} = \begin{bmatrix} 1 & z_{11} & \dots & z_{1r} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & z_{n1} & \dots & z_{nr} \end{bmatrix}_{n \times (r+1)}$$

$$\beta = \begin{bmatrix} \beta_{01} & \dots & \beta_{0m} \\ \vdots & \ddots & \vdots \\ \beta_{r1} & \dots & \beta_{rm} \end{bmatrix}_{(r+1) \times m}, \quad \varepsilon = \begin{bmatrix} \varepsilon'_1 \\ \vdots \\ \varepsilon'_n \end{bmatrix}_{n \times m}$$

Multivariate multiple regression model: $Y = Z\beta + \varepsilon$

$$= \begin{bmatrix} \varepsilon_{11} & \dots & \varepsilon_{1m} \\ \vdots & \ddots & \vdots \\ \varepsilon_{n1} & \dots & \varepsilon_{nm} \end{bmatrix}$$

Each observation has m -response variables which are correlated.

$$Y' = Z' \beta + \varepsilon' \quad (\text{the model})$$

$1 \times m$ $1 \times (r+1)$ $(r+1) \times m$ $1 \times m$

$$E[\varepsilon'] = 0'$$

$$\text{cov}(\varepsilon') = \Sigma_{m \times m}$$

$$= \begin{bmatrix} \sigma_{11} & \dots & \sigma_{1m} \\ \vdots & & \vdots \\ \sigma_{m1} & \dots & \sigma_{mm} \end{bmatrix}$$

" m = # responses for each individual observation"

we observe the values of the predictors and responses from n individuals.

We arrange them into

$$Y = \begin{bmatrix} Y_1' \\ \vdots \\ Y_n' \end{bmatrix} = \begin{bmatrix} Y_{11}, \dots, Y_{1m} \\ \vdots \\ Y_{n1}, \dots, Y_{nm} \end{bmatrix}$$

$n \times m$

$$Z = \begin{bmatrix} Z_1' \\ \vdots \\ Z_n' \end{bmatrix} = \begin{bmatrix} 1 & Z_{11} & \dots & Z_{1r} \\ \vdots & \vdots & & \vdots \\ 1 & Z_{n1} & \dots & Z_{nr} \end{bmatrix}, \quad \beta = \begin{bmatrix} \beta_{01} & \dots & \beta_{0m} \\ \vdots & & \vdots \\ \beta_{r1} & \dots & \beta_{rm} \end{bmatrix}$$

$n \times (r+1)$ $(r+1) \times m$

$$\varepsilon = \begin{bmatrix} \varepsilon_1' \\ \vdots \\ \varepsilon_n' \end{bmatrix} = \begin{bmatrix} \varepsilon_{11}, \dots, \varepsilon_{1m} \\ \vdots \\ \varepsilon_{n1}, \dots, \varepsilon_{nm} \end{bmatrix}$$

$n \times m$

so, the model becomes

$$Y = Z\beta + \varepsilon$$

$\varepsilon_1', \dots, \varepsilon_n'$ are independent.

$$\begin{bmatrix} \beta_{01} \\ \vdots \\ \beta_{r1} \end{bmatrix} \dots \begin{bmatrix} \beta_{0m} \\ \vdots \\ \beta_{rm} \end{bmatrix}$$

$$Y = [Y_{(1)} | \dots | Y_{(m)}], \quad \beta = [\beta_{(1)} | \dots | \beta_{(m)}]$$

$$\varepsilon = [\varepsilon_{(1)} | \dots | \varepsilon_{(m)}]$$

$$\begin{bmatrix} Y_{11} \\ \vdots \\ Y_{n1} \end{bmatrix} \dots \begin{bmatrix} Y_{1m} \\ \vdots \\ Y_{nm} \end{bmatrix}$$

$$\begin{bmatrix} \varepsilon_{11} \\ \vdots \\ \varepsilon_{n1} \end{bmatrix} \dots \begin{bmatrix} \varepsilon_{1m} \\ \vdots \\ \varepsilon_{nm} \end{bmatrix}$$

$$\text{so, } [Y_{(1)} | \dots | Y_{(m)}] = Z [\beta_{(1)} | \dots | \beta_{(m)}] + [\varepsilon_{(1)} | \dots | \varepsilon_{(m)}]$$

$$\text{so, } Y_{(i)} = Z\beta_{(i)} + \varepsilon_{(i)}, \text{ for } i = 1, 2, \dots, m.$$

★ Goal: Estimate β , i.e., $\hat{\beta}$.

one method is to estimate $\beta_{(i)}$ by $\hat{\beta}_{(i)}$ (the least square estimate from the univariate case).

$$\hat{\beta}_{(i)} = (Z'Z)^{-1} Z' Y_{(i)}$$

Using this, $\hat{\beta} = [\hat{\beta}_{(1)} | \dots | \hat{\beta}_{(m)}]$

$$= (Z'Z)^{-1} Z' [Y_{(1)} | \dots | Y_{(m)}]$$

$$= (Z'Z)^{-1} Z' Y$$

For this $\hat{\beta}$, $\hat{Y} = Z\hat{\beta} = Z(Z'Z)^{-1} Z' Y = P_Z Y$.

and $\hat{\varepsilon} = Y - \hat{Y} = (I - P_Z) Y$

$\hat{Y} = Z(Z'Z)^{-1} Z' Y$

least square The error for an estimator B of $\hat{\beta}$ defined to least square

$$f(B) = \text{tr} \left[\underbrace{(Y - ZB)'}_{m \times n} \underbrace{(Y - ZB)}_{n \times m} \right]_{m \times m}$$

$$A = [A_{(1)} \dots A_{(m)}]$$

$$(A'A)_{ij}$$

Fact $\hat{\beta}$ minimizes $f(\cdot)$, i.e.,
 $f(\hat{\beta}) \leq f(B)$ for any B .

$$f(B) = \sum_{i=1}^m (Y_{(i)} - Z\beta_{(i)})'(Y_{(i)} - Z\beta_{(i)})$$

The i th summand is minimized at $\beta_{(i)} = \hat{\beta}_{2(i)}$ (from our knowledge about the univariate case).

$$\text{So, } f(B) \geq \sum_{i=1}^m (Y_{(i)} - Z\hat{\beta}_{2(i)})'(Y_{(i)} - Z\hat{\beta}_{2(i)}) = f(\hat{\beta})$$

Now, suppose we consider a different least square error

$$g(B) = |(Y - ZB)'(Y - ZB)|$$

Fact (from Lin. Alg.)

If C and D are symmetric non-negative definite (nnd)

matrices of same size, then

$$|C+D| \geq |C| + |D|$$

determinant = product of eigenvalue.

proof case 1: If C and D are both positive semi-definite (nnd but not pd) $\downarrow$ (some eigenvalue is zero)

then $|C| = |D| = 0$ and $|C+D| \geq 0$

case 2: one of the matrices is pd. say D is pd.

then $C+D = D^{1/2}(D^{-1/2}CD^{-1/2})D^{1/2} + D$.

$$= D^{1/2}(D^{-1/2}CD^{-1/2} + I)D^{1/2}$$

$$\text{So, } |C+D| = |D^{1/2}| |I + D^{-1/2}CD^{-1/2}| |D^{1/2}| \geq |D^{1/2}| (1 + |D^{-1/2}CD^{-1/2}|) |D^{1/2}| \geq |D| + |C|$$

$$\text{write } Y - ZB = (Y - Z\hat{\beta}) + Z(\hat{\beta} - B)$$

$$\text{so, } (Y - ZB)'(Y - ZB)$$

$$= (Y - Z\hat{\beta})'(Y - Z\hat{\beta}) + (Z(\hat{\beta} - B))'Z(\hat{\beta} - B) + (Y - Z\hat{\beta})'Z(\hat{\beta} - B) + [Z(\hat{\beta} - B)]'(Y - Z\hat{\beta})$$

$$= 0 \quad \text{by } \hat{\beta} = (Z'Z)^{-1}Z'Y \Leftrightarrow (Z'Z)\hat{\beta} = Z'Y$$

$$(Z(\hat{\beta} - B))'Z(\hat{\beta} - B) = 0$$

$$\text{so, } g(B) = |(Y - ZB)'(Y - ZB)|$$

$$= |(Y - Z\hat{\beta})'(Y - Z\hat{\beta}) + (Z\hat{\beta} - ZB)'(Z\hat{\beta} - ZB)| \geq |(Y - Z\hat{\beta})'(Y - Z\hat{\beta})| + |(Z\hat{\beta} - ZB)'(Z\hat{\beta} - ZB)|$$

$$\text{so, } g(B) \geq g(\hat{\beta}) + |(\hat{\beta} - B)'Z'Z(\hat{\beta} - B)|$$

$$\geq g(\hat{\beta}) \quad (\text{as } A'A \text{ is nnd for any } A)$$

$$\Rightarrow |A'A| \text{ is non-negative.}$$

Under the additional normality assumption, i.e., $\varepsilon \sim N_m(0, \Sigma)$.
 we want to find MLE for β, Σ .

Def of MLE, Likelihood func...

The likelihood function

$L(\beta, \Sigma) =$ joint density of Y 's

$$= \prod_{i=1}^n (\text{pdf of } Y_i)$$

Recall $Y' = Z'\beta + \varepsilon'$

If $\varepsilon' \sim N_m(0, \Sigma)$, then $Y' \sim N_m(Z'\beta, \Sigma)$.

pdf of Y : $f_Y(Y) = (2\pi)^{-nm/2} |\Sigma|^{-1/2} \exp\left(-\frac{1}{2} \underbrace{(Y - Z'\beta)' \Sigma^{-1} (Y - Z'\beta)}_{(Y - Z'\beta)'}\right)$

$$\rightarrow = \prod_{i=1}^n f_{Y_i}(y_i) = (2\pi)^{-nm/2} |\Sigma|^{-1/2} \exp\left(-\frac{1}{2} \sum_{i=1}^n (y_i - \beta' z_i)' \Sigma^{-1} (y_i - \beta' z_i)\right)$$

So, $L(\beta, \Sigma) = (2\pi)^{-nm/2} |\Sigma|^{-1/2} \exp\left(-\frac{1}{2} \text{tr}\left[\underbrace{(Y - Z'\beta)}_{n \times m} \underbrace{\Sigma^{-1}}_{m \times m} \underbrace{(Y - Z'\beta)'}_{m \times n}\right]\right)$

Fact $\text{tr}(CD) = \text{tr}(DC)$

$$= (2\pi)^{-nm/2} |\Sigma|^{-1/2} \exp\left(-\frac{1}{2} \text{tr}\left[\Sigma^{-1} (Y - Z'\beta)' (Y - Z'\beta)\right]\right)$$

We will maximize in 2 steps.

First, we wrt Σ , then wrt β .

Step 1 β is fixed. call $(Y - Z'\beta)'(Y - Z'\beta) = D$.

define $h(\Sigma) = |\Sigma|^{-1/2} \exp\left(-\frac{1}{2} \text{tr}(\Sigma^{-1} D)\right)$

$\text{tr}(\Sigma^{-1} D) = \text{tr}(\Sigma^{-1/2} D \Sigma^{-1/2}) = \sum_{i=1}^m \lambda_i$, where $\lambda_1, \dots, \lambda_m$ are the eigenvalues of $\Sigma^{-1/2} D \Sigma^{-1/2}$.

Then, $\prod_{i=1}^m \lambda_i = |\Sigma^{-1/2} D \Sigma^{-1/2}| = \frac{|D|}{|\Sigma|}$

then, $h(\Sigma) = \frac{1}{|D|^{m/2}} \left(\frac{|D|}{|\Sigma|}\right)^{m/2} \exp\left(-\frac{1}{2} \sum_{i=1}^m \lambda_i\right)$

$$= \frac{1}{|D|^{m/2}} \prod_{i=1}^m \lambda_i e^{-\frac{1}{2} \lambda_i}$$

A is nnd if $\mathbf{x}'\mathbf{A}\mathbf{x} \geq 0$ for all $\mathbf{x}$

If A is symmetric, then nnd $\Leftrightarrow$ all eigen values ≥ 0

A is p.d if $\mathbf{x}'\mathbf{A}\mathbf{x} > 0$ unless $\mathbf{x} = \mathbf{0}$.

If A is symmetric, then p.d $\Leftrightarrow$ all eig values > 0 .

positive semidefinite = $\{\text{nnd}\} - \{\text{p.d}\}$.

Recall multivariate multiple linear regression.

model
$$\mathbf{Y} = \mathbf{Z}\mathbf{\beta} + \mathbf{\varepsilon}$$

$$\begin{matrix} n \times m & n \times (r+1) & (r+1) \times m \end{matrix}$$

$$\hat{\mathbf{\beta}} = (\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{Y}$$

$$\hat{\mathbf{\Sigma}} = \frac{1}{n}\hat{\mathbf{\varepsilon}}'\hat{\mathbf{\varepsilon}}$$

Result

If the rows of $\mathbf{\varepsilon}$ are iid $N_m(\mathbf{0}, \mathbf{\Sigma})$, then the MLE of $\mathbf{\beta}$ and

$\mathbf{\Sigma}$ will be $\hat{\mathbf{\beta}} = (\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{Y}$, $\hat{\mathbf{\Sigma}} = \frac{1}{n}\hat{\mathbf{\varepsilon}}'\hat{\mathbf{\varepsilon}}$, where

$$\hat{\mathbf{\varepsilon}} = \mathbf{Y} - \hat{\mathbf{Y}} = \mathbf{Y} - \mathbf{Z}\hat{\mathbf{\beta}} = \mathbf{Y} - \mathbf{Z}(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{Y} = (\mathbf{I} - \mathbf{P}_{\mathbf{Z}})\mathbf{Y}$$

Proof Recall the likelihood function

$$L(\mathbf{\beta}, \mathbf{\Sigma}) = (2\pi)^{-\frac{nm}{2}} |\mathbf{\Sigma}|^{-n/2} \exp\left(-\frac{1}{2} \text{tr}(\mathbf{\Sigma}^{-1}\mathbf{D})\right), \text{ where}$$

$$\mathbf{D} = (\mathbf{Y} - \mathbf{Z}\mathbf{\beta})'(\mathbf{Y} - \mathbf{Z}\mathbf{\beta})$$

For fixed $\mathbf{\beta}$ (so, $\mathbf{D}$ is also fixed), we will maximize

$$f(\mathbf{\Sigma}) = |\mathbf{\Sigma}|^{-n/2} \exp\left(-\frac{1}{2} \text{tr}(\mathbf{\Sigma}^{-1}\mathbf{D})\right)$$

symmetric, then eigenvalues are real value.

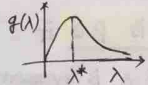
$$= \frac{1}{|\mathbf{D}|^{n/2}} \left(\frac{|\mathbf{D}|}{|\mathbf{\Sigma}|}\right)^{n/2} \exp\left[-\frac{1}{2} \text{tr}(\mathbf{\Sigma}^{-1/2} \mathbf{D} \mathbf{\Sigma}^{-1/2})\right]$$

If the eigenvalues of $\Sigma^{-1/2} D \Sigma^{-1/2}$ are $\lambda_1, \dots, \lambda_n$, then $\text{tr}(\Sigma^{-1/2} D \Sigma^{-1/2}) = \sum_{i=1}^n \lambda_i$

$$\prod_{i=1}^n \lambda_i = |\Sigma^{-1/2} D \Sigma^{-1/2}| = |\Sigma^{-1/2}| |D| |\Sigma^{-1/2}| = \frac{|D|}{|\Sigma|}$$

"determinant"

Then

$$f(\Sigma) = \frac{1}{|D|^{n/2}} \left(\prod_{i=1}^n \lambda_i \right)^{n/2} \exp\left(-\frac{1}{2} \sum_{i=1}^n \lambda_i\right)$$


$$= \frac{1}{|D|^{n/2}} \prod_{i=1}^n g(\lambda_i), \text{ where } g(\lambda) = \lambda^{n/2} e^{-\lambda/2}$$

zero positive eigenvalues

We will maximize g w.r.t. λ . If λ^* maximizes g , then $f(\Sigma)$ will be maximized by taking $\lambda_i = \lambda^*$ for all $i = 1, 2, \dots, n$.

concave function

$$\frac{d}{d\lambda} \log(g(\lambda)) = \frac{n}{2\lambda} - \frac{1}{2}$$

so, $\frac{d}{d\lambda} (\log g(\lambda)) = 0 \Leftrightarrow \lambda = n$, so $\lambda^* = n$.

To find maximizer of $f(\Sigma)$, we need to find Σ so that all eigenvalues of $\Sigma^{-1/2} D \Sigma^{-1/2}$ are equal to n .

Thus, we must have $\Sigma^{-1/2} D \Sigma^{-1/2} = nI \Rightarrow n\Sigma = D$

multiply $\Sigma^{-1/2}$ to both sides.

$$\Sigma = \frac{1}{n} D$$

so, $f(\Sigma) \leq f(\frac{1}{n} D)$

so, $L(\beta, \Sigma) \leq L(\beta, \frac{1}{n} D) = L(\beta, \frac{1}{n} (Y - Z\beta)'(Y - Z\beta))$

$$= (2\pi)^{-\frac{nm}{2}} |\frac{1}{n} D|^{-\frac{n}{2}} \exp\left[-\frac{1}{2} \text{tr}\left(\left(\frac{1}{n} D\right)^{-1} D\right)\right]$$

$$= (2\pi)^{-\frac{nm}{2}} |D|^{-\frac{n}{2}} n^{n/2} e^{-\frac{nm}{2}}$$

because we have already seen $\leq L(\hat{\beta}, \frac{1}{n} (Y - Z\hat{\beta})'(Y - Z\hat{\beta}))$ that $\beta = \hat{\beta}$ minimizes

$$h(\beta) = (Y - Z\beta)'(Y - Z\beta)$$

$\max_{\beta, \Sigma} L(\beta, \Sigma) = C |\hat{\Sigma}|^{-n/2}$, C is constant

$$(2\pi)^{-\frac{nm}{2}} n^{n/2} e^{-\frac{nm}{2}}$$

properties of $\hat{\beta}$.

$$\begin{aligned} E[\hat{\beta}] &= E[(Z'Z)^{-1} Z' Y] = \\ &= (Z'Z)^{-1} Z' E(Y) = (Z'Z)^{-1} Z' Z \beta = \beta. \end{aligned}$$

Recall $Y = [Y_{(1)} | \dots | Y_{(m)}] = Z\beta + \epsilon$

$$E[\hat{\beta}_{(i)}] = \beta_{(i)}, \text{ for } i = 1, \dots, m$$

$$= Z [\beta_{(1)} | \dots | \beta_{(m)}] + [\epsilon_{(1)} | \dots | \epsilon_{(m)}]$$

$\star \text{cov}(\hat{\beta}_{(i)}, \hat{\beta}_{(k)}) =$

fact: $(CD)_{ij} = C D_{ij}$

$$\text{cov}(Z'Z)^{-1} Z' Y_{(i)}, (Z'Z)^{-1} Z' Y_{(k)}$$

$$\hat{\beta} = (Z'Z)^{-1} Z' Y$$

so, $\hat{\beta}_{(i)} = (Z'Z)^{-1} Z' Y_{(i)}$

$$= (Z'Z)^{-1} Z' \text{cov}(Y_{(i)}, Y_{(k)}) Z (Z'Z)^{-1}$$

$$= \text{cov}(\epsilon_{(i)}, \epsilon_{(k)})$$

$$= \sigma_{jk} I_{n \times n}$$

$$\Sigma = \begin{bmatrix} \sigma_{11} & \dots & \sigma_{1m} \\ \vdots & & \vdots \\ \sigma_{m1} & \dots & \sigma_{mm} \end{bmatrix}$$

$$\epsilon_{(i)} = \begin{bmatrix} \epsilon_{i1} \\ \vdots \\ \epsilon_{in} \end{bmatrix}, \epsilon_{(k)} = \begin{bmatrix} \epsilon_{k1} \\ \vdots \\ \epsilon_{kn} \end{bmatrix}$$

max m
of responses

$$= (Z'Z)^{-1} Z' (\sigma_{jk} I_n) Z (Z'Z)^{-1}$$

$$= \sigma_{jk} (Z'Z)^{-1}$$

Also, $\hat{\gamma}' \hat{\varepsilon} = (P_Z Y)' (I - P_Z) Y = Y' P_Z (I - P_Z) Y = 0$

$\star \text{cov}(\hat{\gamma}_{(j)}, \hat{\varepsilon}_{(k)}) = 0$ for j, k .

$= \text{cov}(P_Z Y_{(j)}, (I - P_Z) Y_{(k)})$

$= P_Z \text{cov}(Y_{(j)}, Y_{(k)}) (I - P_Z) = P_Z (\sigma_{jk} I) (I - P_Z) = 0$.

Recall $\hat{\Sigma} = \frac{1}{n} \hat{\varepsilon}' \hat{\varepsilon} = \frac{1}{n} (Y - Z\hat{\beta})' (Y - Z\hat{\beta})$

$\mathbb{E}[\hat{\Sigma}] =$

to check $\uparrow$
unbiased

$\star$
Find $\mathbb{E}[\hat{\varepsilon}_{(j)}' \hat{\varepsilon}_{(k)}]$

$= \mathbb{E}[Y_{(j)}' (I - P_Z) Y_{(k)}]$

$= \mathbb{E}[Y_{(j)}' (I - P_Z) Y_{(k)}] = \mathbb{E}[\text{tr}(Y_{(j)}' (I - P_Z) Y_{(k)})]$

$= \mathbb{E}[\text{tr}((I - P_Z) Y_{(k)} Y_{(j)}')] = \text{tr} \mathbb{E}[(I - P_Z) Y_{(k)} Y_{(j)}']$

$= \text{tr}((I - P_Z) \{ \text{cov}(Y_{(k)}, Y_{(j)}) + Z\beta_{(k)} \beta_{(j)}' Z' \})$

$= \text{tr}[(I - P_Z) \sigma_{jk} I + 0] \quad \text{tr}(I - P_Z)$

$= \sigma_{jk} \text{tr}(I - P_Z) = \sigma_{jk} \text{rank}(I - P_Z)$

idempotent
matrix

$= \sigma_{jk} (n - r - 1)$

So, $\mathbb{E}[\hat{\varepsilon}' \hat{\varepsilon}] = (n - r - 1) \Sigma$

So, unbiased estimator of Σ is

$\star S = \frac{1}{n - r - 1} \hat{\varepsilon}' \hat{\varepsilon} = \frac{1}{n - r - 1} (Y - Z\hat{\beta})' (Y - Z\hat{\beta})$
 $= \frac{n}{n - r - 1} \hat{\Sigma} = \frac{1}{n - r - 1} Y' (I - P_Z) Y$

Also, $n \hat{\Sigma} = \hat{\varepsilon}' \hat{\varepsilon} \sim W_{m, n - r - 1}(\Sigma)$ "Wishart Distribution"

Also, $\hat{\Sigma}$ and $\hat{\beta}$ are independent.

under normality assumption,

$\hat{\beta}_{(j)} \sim N_{r+1}(\beta_{(j)}, \sigma_{jj} (Z'Z)^{-1})$

3/29 ~ No class

matlab

$$|\hat{\Sigma}| = \det(\hat{\Sigma})$$

$$e^{-\frac{nm}{2}} = \exp\left(-\frac{nm}{2}\right)$$

$$\begin{bmatrix} 1 & 2 \\ 3 & 4 \\ 0 & 0 \\ 0 & 0 \end{bmatrix}$$

Math B7800

3/13/18, Tue

Likelihood Ratio Test

$$\beta_{(r+1) \times m} = \begin{bmatrix} \beta_1 \\ \beta_2 \end{bmatrix} \begin{matrix} (q+1) \times m \\ (r-q) \times m \end{matrix}$$

$$\star \text{Test: } H_0: \beta_2 = 0 \text{ VS } H_1: \beta_2 \neq 0 \quad \begin{matrix} n \times (q+1) \\ n \times (r-q) \end{matrix}$$

Partition Z matrix accordingly: $Z = \begin{bmatrix} Z_1 & Z_2 \end{bmatrix}$

$$E[Y] = Z\beta = \begin{bmatrix} Z_1 & Z_2 \end{bmatrix} \begin{bmatrix} \beta_1 \\ \beta_2 \end{bmatrix} = Z_1\beta_1 + Z_2\beta_2$$

"Under $H_0: \beta_2 = 0$ "

$$E[Y] = Z_1\beta_1 \quad (H_0 \Rightarrow \text{the predictors in } Z_2 \text{ are redundant.})$$

Recall Likelihood Ratio Test.

$$\Lambda = \frac{\max_{H_0} L(\beta, \Sigma)}{\max L(\beta, \Sigma)} \quad \text{"the likelihood ratio"}$$

so, $L(\beta, \Sigma)$ is maximized at "MLE of β, Σ " $\hat{\beta}, \hat{\Sigma}$

$$\star \hat{\beta} = \hat{\beta} = (Z'Z)^{-1}Z'Y \quad \text{and} \quad \hat{\Sigma} = \frac{1}{n}(Y - Z\hat{\beta})(Y - Z\hat{\beta})'$$

$$\text{and} \star L(\hat{\beta}, \hat{\Sigma}) = (2\pi)^{-\frac{nm}{2}} |\hat{\Sigma}|^{-\frac{n}{2}} e^{-\frac{nm}{2}} = \max L(\beta, \Sigma)$$

under H_0 , the model becomes $Y = Z_1\beta_1 + \varepsilon$.

so, under H_0 , $L(\beta, \Sigma)$ is maximized at

$$\star \beta = \begin{bmatrix} (Z_1'Z_1)^{-1}Z_1'Y \\ 0 \end{bmatrix} = \hat{\beta}_1, \quad \star \hat{\Sigma}_1 = \frac{1}{n}(Y - Z_1\hat{\beta}_1)(Y - Z_1\hat{\beta}_1)'$$

$$L\left(\begin{bmatrix} \beta_1 \\ 0 \end{bmatrix}, \hat{\Sigma}_1\right) = (2\pi)^{\frac{-nm}{2}} |\hat{\Sigma}_1|^{-n/2} e^{\frac{-nm}{2}}$$

$$\text{so, } \Delta = \frac{L\left(\begin{bmatrix} \beta_1 \\ 0 \end{bmatrix}, \hat{\Sigma}_1\right)}{L(\hat{\beta}, \hat{\Sigma})} = \left(\frac{|\hat{\Sigma}_1|}{|\hat{\Sigma}|}\right)^{n/2}$$

$$\text{so, } \Delta^{2/n} = \frac{|\hat{\Sigma}|}{|\hat{\Sigma}_1|}, \text{ We will reject } H_0, \text{ when } \Delta \text{ is small or equivalently } \frac{|\hat{\Sigma}|}{|\hat{\Sigma}_1|} \text{ is small.}$$

$$U = -[n-r-1-\frac{1}{2}(m-r+q+1)] \ln \left(\frac{|\hat{\Sigma}|}{|\hat{\Sigma}_1|} \right) \underset{\substack{\text{determinant} \\ H_0}}{\sim} \chi^2_{m(r-q)} \text{ approximately.}$$

Reject H_0 at 5% level if $U > \chi^2_{m(r-q)}(1-\alpha)$

$$\text{Note: } \frac{|\hat{\Sigma}|}{|\hat{\Sigma}_1|} = \frac{|\hat{\Sigma}|}{|\hat{\Sigma} + (\hat{\Sigma}_1 - \hat{\Sigma})|}$$

since $(\hat{\Sigma}_1 - \hat{\Sigma})$ and $\hat{\Sigma}$ are independent (we proved this only in the univariate case), we need to compare $\hat{\Sigma}$ and $\hat{\Sigma}_1 - \hat{\Sigma}$, in particular, consider the ratio matrix: $(\hat{\Sigma}_1 - \hat{\Sigma})(\hat{\Sigma})^{-1}$

Let the positive eigenvalues of $(\hat{\Sigma}_1 - \hat{\Sigma})(\hat{\Sigma})^{-1}$ are $\eta_1, \eta_2, \dots, \eta_s > 0$.

$$\text{• Wilks' lambda: } \prod_{i=1}^s \frac{1}{1+\eta_i} = \frac{|\hat{\Sigma}|}{|\hat{\Sigma} + (\hat{\Sigma}_1 - \hat{\Sigma})|} = \frac{|E|}{|E+H|} \quad \begin{cases} E = \hat{\Sigma} \\ H = \hat{\Sigma}_1 - \hat{\Sigma} \end{cases} \quad (1)$$

$$\text{• Pillai's trace: } \sum_{i=1}^s \frac{\eta_i}{1+\eta_i} = \text{tr}((\hat{\Sigma}_1 - \hat{\Sigma})(\hat{\Sigma}_1 - \hat{\Sigma} + \hat{\Sigma})^{-1}) = \text{tr}((\hat{\Sigma}_1 - \hat{\Sigma})(\hat{\Sigma}_1)^{-1}) = L_2 \quad (2)$$

$$\text{• Hotelling trace: } \sum_{i=1}^s \eta_i = \text{tr}((\hat{\Sigma}_1 - \hat{\Sigma})(\hat{\Sigma})^{-1}) = L_3 \quad (3)$$

$$\text{• Roy's greatest root: } \frac{\eta_1}{1+\eta_1} = L_4 \quad (4)$$

Fact (Linear Algebra).

For two matrices E and H , (both invertible, symmetric, and same size)

$$\frac{|E|}{|E+H|} = \frac{1}{|E^{-1}| |E+H|} = \frac{1}{|I + E^{-1}H|}$$

$$= \prod_{i=1}^s \frac{1}{1+\eta_i} \text{ if the eigenvalues of } E^{-1}H \text{ are } \eta_1, \eta_2, \dots, \eta_s, 0.$$

Apply this with $E = \hat{\Sigma}$, $H = \hat{\Sigma}_1 - \hat{\Sigma}$ to obtain (1)

$$\text{For (2), } \sum_{i=1}^s \frac{\eta_i}{1+\eta_i} = \sum_{i=1}^s \frac{1}{1+\frac{1}{\eta_i}} \quad \text{m-dimension}$$

$$\text{trace}(H(E+H)^{-1}) = \text{tr}(H[(EH^{-1}+I)H]^{-1}) = \text{tr}(HH^{-1}(EH^{-1}+I)^{-1}) = \text{tr}((I+EH^{-1})^{-1})$$

If the eigenvalues of HE^{-1} are $\eta_1, \dots, \eta_s$, then the eigenvalues of $\frac{EH^{-1}}{(HE^{-1})^{-1}}$ are $\frac{1}{\eta_1}, \dots, \frac{1}{\eta_s}$, then

the eigenvalues of $(I+EH^{-1})$ are $1+\frac{1}{\eta_1}, \dots, 1+\frac{1}{\eta_s}$, then the eigenvalues of $(I+EH^{-1})^{-1}$ are $\frac{1}{1+\frac{1}{\eta_1}}, \dots, \frac{1}{1+\frac{1}{\eta_s}}$

$$\text{then } \text{tr}((I+EH^{-1})^{-1}) = \sum_{i=1}^s \frac{1}{1+\frac{1}{\eta_i}} \text{ this satisfies (2)}$$

We will reject H_0 if L_1 is small, or L_2 is large, or L_3 is large, or L_4 is large.

Fact (Linear Algebra)

Suppose A and B are two matrices (same size and squared).

A is n.n.d and B is p.d. (so, B is non singular) $\xrightarrow{\text{invertible}} \det(B) \neq 0$.

then the eigenvalues of AB^{-1} are non-negative.

(Fact) For two matrices C and D, the non-zero eigenvalues of CD and DC are same.

using this, all non-zero eigenvalues of AB^{-1} and $B^{-1/2}AB^{-1/2}$ are same.

$$B = \begin{bmatrix} \beta_{01} & \dots & \beta_{0m} \\ \beta_{11} & \dots & \beta_{1m} \\ \vdots & \ddots & \vdots \\ \beta_{r1} & \dots & \beta_{rm} \end{bmatrix} \quad (r+1) \times m$$

$$\hat{\beta} = (Z'Z)^{-1}Z'Y$$

$$\hat{\beta}' = Y'Z(Z'Z)^{-1}$$

Prediction

For a future observation with (unknown) response vector y_0 and (known) predictor values z_0 ,

$$E[y_0] = z_0'\beta, \quad E[y_0] = \beta'z_0$$

m -entries $\quad m \times 1 \quad m \times (r+1) \quad (r+1) \times 1$

$$z_0 = \begin{bmatrix} 1 \\ z_{01} \\ \vdots \\ z_{0r} \end{bmatrix} \quad (r+1) \times 1$$

$\rightarrow$ confidence region $\Rightarrow$ Prediction region

1. Point estimator of both y_0 and $E(y_0)$

$$\hat{\beta}'z_0 = Y'Z(Z'Z)^{-1}z_0$$

2. Confidence region for $E(y_0)$

y_0 is a vector not a single value. so, not confidence interval.

3. Prediction region for y_0

$$y_0 \sim N_m(\beta'z_0, \Sigma), \quad \hat{\beta}'z_0 \sim N_m(\beta'z_0, (z_0'(Z'Z)^{-1}z_0)\Sigma)$$

Recall

$$\beta = [\beta_{(1)} \mid \dots \mid \beta_{(m)}]$$

$$E[y_0] = \begin{bmatrix} z_0'\beta_{(1)} \\ \vdots \\ z_0'\beta_{(m)} \end{bmatrix}$$

$$\text{cov}(\beta_{(1)}, \beta_{(m)}) = \sigma_{jk}(Z'Z)^{-1}$$

$$\text{so, cov}(z_0'\beta_{(1)}, z_0'\beta_{(m)}) = z_0'\sigma_{jk}(Z'Z)^{-1}z_0 = (z_0'(Z'Z)^{-1}z_0)\sigma_{jk}$$

$$\text{Hotelling } T^2 = \frac{(\hat{\beta}'z_0 - \beta'z_0)'}{\sqrt{z_0'(Z'Z)^{-1}z_0}} S^{-1} \frac{(\hat{\beta}'z_0 - \beta'z_0)}{\sqrt{z_0'(Z'Z)^{-1}z_0}}$$

$$\text{, where } S = \frac{1}{n-r-1} (Y - Z\hat{\beta})'(Y - Z\hat{\beta})$$

$$100(1-\alpha)\% \text{ confidence region} = \{\beta'z_0 : T^2 \leq \frac{m(n-r-1)}{n-r-m} F_{m, n-r-m}(\alpha)\}$$

The $100(1-\alpha)\%$ simultaneous confidence intervals for $E(Y_i) = Z_0' \beta_{(i)}$ are

$$\hat{Z}_0' \hat{\beta}_{(i)} \pm \sqrt{\frac{m(n-r-1)}{n-r-m}} F_{m, n-r-m}(\alpha) \sqrt{\hat{Z}_0' (Z'Z)^{-1} \hat{Z}_0 S_{ii}}, \quad i=1, \dots, m$$

where $\hat{\beta}_{(i)}$ is the i th column of $\hat{\beta}$ and $\hat{Z}_0$ is the i th diagonal element of S .

$$Y_0 - \hat{\beta}' Z_0 = (\beta - \hat{\beta})' Z_0 + \varepsilon_0 \sim N_m(0, (1 + Z_0' (Z'Z)^{-1} Z_0) \Sigma)$$

The $100(1-\alpha)\%$ prediction ellipsoid for Y_0 becomes

$$(Y_0 - \hat{\beta}' Z_0)' S^{-1} (Y_0 - \hat{\beta}' Z_0) \leq (1 + Z_0' (Z'Z)^{-1} Z_0) \left[\frac{m(n-r-1)}{n-r-m} F_{m, n-r-m}(\alpha) \right]$$

The $100(1-\alpha)\%$ simultaneous prediction intervals for the individual responses Y_{0i} :

$$\hat{Z}_0' \hat{\beta}_{(i)} \pm \sqrt{\frac{m(n-r-1)}{n-r-m}} F_{m, n-r-m}(\alpha) \sqrt{(1 + Z_0' (Z'Z)^{-1} Z_0) S_{ii}}, \quad i=1, \dots, m$$

$\hat{Z}_0'$ Beta-hat $(:, i)$

$$Z_0 = [1 \quad Z(3, 2:end)]'$$

$$S_{ii} = S_{ii} \text{ Beta-hat}(i, 3)$$

Math B7800

3/15/18, Thur

Recall Suppose A is pd matrix ($n \times n$) and symmetric.

$$\max \{ (b'x)^2 : x'Ax \leq 1 \} = b'A^{-1}b$$

$$\text{Proof } (b'x) = \underbrace{b'A^{-1/2}}_{\text{vector}} \underbrace{A^{1/2}x}_{\text{vector}} = (A^{-1/2}b)' (A^{1/2}x)$$

$$\begin{aligned} \text{Then } (b'x)^2 &\leq \left[(A^{-1/2}b)' (A^{-1/2}b) \right] \left[(A^{1/2}x)' (A^{1/2}x) \right] \\ &\approx (b'A^{-1}b) \underbrace{(x'Ax)}_{\leq 1} \approx b'A^{-1}b \end{aligned}$$

by Cauchy-Schwarz Inequality

For two $n \times 1$ vector x and y , $(x'y) \leq (x'x)(y'y)$
Equality holds if $y = cx$ for some constant c .

Equality holds if $A^{1/2}x = cA^{-1/2}b$ for some constant c ,
and c should be such that $x'Ax = 1$.

$$* \quad x = cA^{-1}b \quad \text{so, } x'Ax = c^2 b'A^{-1}b = 1$$

$$(c b'A^{-1}b)' A (c b'A^{-1}b) \rightarrow \text{so, } c = \pm \frac{1}{\sqrt{b'A^{-1}b}}$$

$$\text{The value of } x \text{ that attains the max is } \pm \frac{1}{\sqrt{b'A^{-1}b}} A^{-1}b$$

Recall the model

$$Y = Z\beta + \varepsilon, \quad Y_{n \times m} = \begin{bmatrix} y_1' \\ \vdots \\ y_n' \end{bmatrix}, \quad Z_{n \times (r+1)} = \begin{bmatrix} z_1' \\ \vdots \\ z_n' \end{bmatrix}$$

$$\beta_{(r+1) \times m} = [\beta_{(1)} \mid \dots \mid \beta_{(m)}] \quad \begin{bmatrix} \beta_{10} & \beta_{11} \\ \beta_{20} & \beta_{21} \\ \vdots & \vdots \end{bmatrix}$$

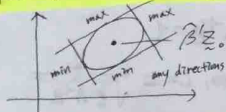
$$\varepsilon_{n \times m} = \begin{bmatrix} \varepsilon_1' \\ \vdots \\ \varepsilon_n' \end{bmatrix} \quad \text{such that } \varepsilon_1, \dots, \varepsilon_n \text{ are iid } N_m(0, \Sigma) \quad (\text{assumption}) \rightarrow \hat{\beta} = (Z'Z)^{-1}Z'Y$$

Now, we have a future observation with unknown response y_0 , but known predictors z_0 .

$$y_0 \sim N(\beta'z_0, \Sigma), \quad E[y_0] = \beta'z_0$$

We saw that $100(1-\alpha)\%$ CR for $\beta'z_0$ is

$$\{z: (z - \hat{\beta}'z_0)' S^{-1} (z - \hat{\beta}'z_0) \leq \frac{m(n-r-1)}{n-r-m} F_{m, n-r-m}(\alpha)\}$$



Q: Obtain simultaneous Scheffe CI for linear combinations of the components of $\beta'z_0$.

so, we want CI for $b'\beta'z_0$ for all $m \times 1$ vectors b .

$$\hat{\beta}'z_0 = \{z: (z - \hat{\beta}'z_0)' A (z - \hat{\beta}'z_0) \leq 1\}, \text{ where}$$

$$A = \frac{S^{-1}}{\frac{m(n-r-1)}{n-r-m} F_{m, n-r-m}(\alpha)}$$

If $y = z - \hat{\beta}'z_0$, then the constraint is $y'Ay \leq 1$ under this constraint, $(b'y)^2 \leq b'A^{-1}b$ by Cauchy-Schwarz inequality.

$$\text{so, } (b'(z - \hat{\beta}'z_0))^2 \leq b' \frac{m(n-r-1)}{n-r-m} S F_{m, n-r-m}(\alpha) b$$

so, z represents $\beta'z_0$.

so the simultaneous CI for $b'\beta'z_0$ is

$$b'\beta'z_0 \pm \sqrt{\frac{m(n-r-1)}{n-r-m} F_{m, n-r-m}(\alpha)} b'Sb$$

$P(\text{the interval } I_b \text{ contains } b'\beta'z_0 \text{ for all } b) = 1 - \alpha$
simultaneous.

• General regression function "joint distribution of Y and Z "

Suppose $(Y, Z_1, \dots, Z_r)$ have a joint distribution with mean

$$\mu = \begin{bmatrix} \mu_Y \\ \mu_Z \end{bmatrix}_{(r+1) \times 1} \text{ and covariance } \Sigma = \begin{bmatrix} \sigma_{YY} & \sigma'_{YZ} \\ \sigma_{YZ} & \Sigma_{ZZ} \end{bmatrix}_{(r+1) \times (r+1)}$$

Goal: predict Y based on $Z_1, \dots, Z_r$ ← random variables

We predict Y based on $Z_1, \dots, Z_r$ via some function $g(Z_1, \dots, Z_r)$

prediction error-minimized

prediction error = $Y - g(Z_1, \dots, Z_r)$, $g: \mathbb{R}^r \rightarrow \mathbb{R}$

mean prediction squared error = $E[(Y - g(Z_1, \dots, Z_r))^2]$

What function minimizes $g(Z_1, \dots, Z_r)$?

Step 1 $E(Y - a)^2 = h(a)$, $h(a)$ is minimized at $a = E(Y)$

"minimize"

Step 2 $E[(Y - a)^2 | Z] = h(a)$, $h(a)$ is minimized at $E(Y|Z) = \frac{E(Y, Z)}{E(Z)}$

Step 3 $E[(Y - a)^2 | Z_1, \dots, Z_r]$ is minimized at $E(Y | Z_1, \dots, Z_r) = g^*(Z_1, \dots, Z_r)$

$$E[(Y - g(Z_1, \dots, Z_r))^2] = E\{E[(Y - g(Z_1, \dots, Z_r))^2 | Z_1, \dots, Z_r]\} \quad (*)$$

$$\oplus E[(Y - g(Z_1, \dots, Z_r))^2 | Z_1, \dots, Z_r] \geq E[(Y - g^*(Z_1, \dots, Z_r))^2 | Z_1, \dots, Z_r]$$

Taking expectation both sides,

$$E[(Y - g(Z_1, \dots, Z_r))^2] \geq E[(Y - g^*(Z_1, \dots, Z_r))^2] \text{ by } (*)$$

So, the minimizer $E(Y | Z_1, \dots, Z_r)$ is called the regression function of Y on $Z_1, \dots, Z_r$. "Jointly normal $\Rightarrow$ linear regression"

Q: When $\begin{pmatrix} Y \\ Z_1 \\ \vdots \\ Z_r \end{pmatrix} \sim N_{r+1}(\mu, \Sigma)$, then $E(Y | Z_1, \dots, Z_r) = ?$

Then the conditional distribution of Y given $Z_1, \dots, Z_r$ is

$$\star N(\mu_Y + \sigma'_{YZ} \Sigma_{ZZ}^{-1} (Z - \mu_Z), \sigma_{YY} - \sigma'_{YZ} \Sigma_{ZZ}^{-1} \sigma_{YZ})$$

$$\text{so, } E[Y | Z_1, \dots, Z_r] = \mu_Y + \sigma'_{YZ} \Sigma_{ZZ}^{-1} (Z - \mu_Z) \leftarrow \text{constant} \Rightarrow \text{linear}$$

$$= \beta_0 + \beta_Z' Z, \text{ where}$$

$$\beta_0 = (\mu_Y - \sigma'_{YZ} \Sigma_{ZZ}^{-1} \mu_Z), \beta_Z = \Sigma_{ZZ}^{-1} \sigma_{YZ}$$

★ result If the joint distribution is not normal, then

$E(Y | Z_1, \dots, Z_r)$ need not to be linear.

Ex Suppose (Y, Z) has joint pdf, $f_{Y,Z}(y, z) = ze^{-yz-1}$, $0 \leq y, z < \infty$

Find $E(Y | Z) = ?$

$$f_{Y|Z}(y|z) = \frac{f_{Y,Z}(y, z)}{f_Z(z)} = \frac{ze^{-yz-1}}{\int_0^\infty ze^{-yz-1} dy} = ze^{-yz}$$

$$E[Y | Z] = \int_0^\infty y \cdot f_{Y|Z}(y|z) dy = \int_0^\infty y ze^{-yz} dy = \frac{1}{z}$$

so, $E[Y | Z]$ is not linear in Z .

* In general, (without normality) among all linear predictors of y based on $z_1, \dots, z_r$, the one that minimizes the mean squared error is $\beta_0 + \beta'z$, where $\beta_0 = (\mu_y - \sigma_{yz}' \Sigma_{zz}^{-1} \mu_z)$, $\beta = \Sigma_{zz}^{-1} \sigma_{yz}$.

Def A linear predictor of y based on $z_1, \dots, z_r$ has the form $b_0 + b'z$ for some constants $b_0, b = \begin{pmatrix} b_1 \\ \vdots \\ b_r \end{pmatrix}$.

If $f(b_0, b) = E[(y - b_0 - b'z)^2] \geq f(\beta_0, \beta)$ for any b_0, b .

Math B1800

3/20/18, Tue

• Best Linear Predictor

Recall $(y, z_1, \dots, z_r)$ has a joint distribution with mean $\mu = \begin{bmatrix} \mu_y \\ \mu_z \end{bmatrix}$ ($(r+1) \times 1$), covariance $\Sigma = \begin{bmatrix} \sigma_{yy} & \sigma_{yz}' \\ \sigma_{yz} & \Sigma_{zz} \end{bmatrix}$ ($(r+1) \times (r+1)$).

Goal: Find a linear predictor of y based on $z_1, \dots, z_r$, which has smallest possible mean squared error.

A linear predictor has the form $b_0 + b'z$, $z = \begin{bmatrix} z_1 \\ \vdots \\ z_r \end{bmatrix}$.

Its mean squared error: $g(b_0, b) = E[(y - b_0 - b'z)^2]$

Goal: Minimize $g(b_0, b)$ w.r.t. b_0, b .

Result: $g(b_0, b) \geq g(\beta_0, \beta)$, where $\beta_0 = \mu_y - \sigma_{yz}' \Sigma_{zz}^{-1} \mu_z$, $\beta = \Sigma_{zz}^{-1} \sigma_{yz}$.

Proof $g(b_0, b) = E[(y - b_0 - b'z)^2]$, " $(a+b+c)^2 = a^2 + b^2 + c^2 + 2ab + 2bc + 2ac$ "

$$= E[(y - \mu_y) - b'(\bar{z} - \mu_z) - b_0 + \mu_y - b'\mu_z]^2$$

$$= E[(y - \mu_y)^2] + E[(b'(\bar{z} - \mu_z))^2] + (b_0 - \mu_y + b'\mu_z)^2$$

$$+ 0 + 0 - 2E[(y - \mu_y)b'(\bar{z} - \mu_z)] - 2E[(y - \mu_y)(b_0 - \mu_y + b'\mu_z)] - 2E[b'(\bar{z} - \mu_z)(b_0 - \mu_y + b'\mu_z)]$$

$$= \sigma_{yy} + b' \Sigma_{zz} b - 2b' \sigma_{yz} - 2(b_0 - \mu_y + b'\mu_z) \sigma_{yz}'$$

so, $g(b_0, b) \geq g(\mu_y - b'\mu_z, b)$

$$= \sigma_{yy} + b' \Sigma_{zz} b - 2b' \sigma_{yz} \quad \text{--- (*)}$$

$$\frac{\partial g(\mu_y - k' \mu_z, k)}{\partial k} = 0 + 2 \sum_{zz} k - 2 \Sigma_{yz}$$

Set $\frac{\partial g}{\partial k} = 0$, we get $k = \underbrace{\sum_{zz}^{-1} \Sigma_{yz}}_{\beta} \leftarrow \text{"minimizer of } k"$

so, $g(b_0, k) \geq g(\mu_y - k' \mu_z, k)$

$\geq g(\mu_y - \beta' \mu_z, \beta) = g(\beta_0, \beta)$

$\beta_0 = \mu_y - \Sigma_{yz}^{-1} \sum_{zz} \mu_z, \beta = \sum_{zz}^{-1} \Sigma_{yz}$

• Smallest mean squared error among linear predictors:

$$E[(y - \beta_0 - \beta' z)^2] = E[(y - (\mu_y - \Sigma_{yz}^{-1} \sum_{zz} \mu_z) - \Sigma_{yz}^{-1} \sum_{zz} z)^2]$$

$$= E[(y - \mu_y + \Sigma_{yz}^{-1} \sum_{zz} \mu_z - \Sigma_{yz}^{-1} \sum_{zz} z)^2]$$

$$= E[(y - \mu_y - \Sigma_{yz}^{-1} \sum_{zz} (z - \mu_z))^2] \quad \left(\begin{matrix} \beta \\ \Sigma_{yz}^{-1} \Sigma_{zz}^{-1} \Sigma_{yz} \end{matrix} \right)$$

$$= g(\beta_0, \beta) = \sigma_{yy} + \beta' \sum_{zz} \beta - 2 \beta' \Sigma_{yz}$$

by plug-in $\rightarrow = \sigma_{yy} + \Sigma_{yz}^{-1} \sum_{zz} \Sigma_{yz} - 2 \Sigma_{yz}^{-1} \sum_{zz} \Sigma_{yz}$

$$= \sigma_{yy} - \Sigma_{yz}^{-1} \sum_{zz} \Sigma_{yz}$$

$$= \sigma_{yy} \left[1 - \frac{\Sigma_{yz}^{-1} \sum_{zz} \Sigma_{yz}}{\sigma_{yy}} \right]$$

$$= \sigma_{yy} [1 - \rho_{y(z)}^2], \text{ where } \rho_{y(z)} = \pm \sqrt{\frac{\Sigma_{yz}^{-1} \sum_{zz} \Sigma_{yz}}{\sigma_{yy}}}$$

called the "population multiple correlation coefficient between y and z "

Result

$$\max_{b_0, k} \text{corr}(y, b_0 + k' z) = \text{corr}(y, \beta_0 + \beta' z) = \rho_{y(z)}$$

Proof $[\text{corr}(y, b_0 + k' z)]^2 = \frac{[\text{cov}(y, b_0 + k' z)]^2}{\text{var}(y) \text{var}(b_0 + k' z)} = \frac{(k' \Sigma_{yz})^2}{\sigma_{yy} k' \sum_{zz} k}$

Recall using Cauchy-Schwarz Inequality.

$$(k' \Sigma_{yz})^2 = (k' \sum_{zz}^{-1/2} \sum_{zz}^{-1/2} \Sigma_{yz})^2 \leq (k' \sum_{zz} k) (\Sigma_{yz}' \sum_{zz}^{-1} \Sigma_{yz})$$

$$\text{so, } \frac{(k' \Sigma_{yz})^2}{k' \sum_{zz} k} \leq \Sigma_{yz}' \sum_{zz}^{-1} \Sigma_{yz}$$

$$\text{so, } [\text{corr}(y, b_0 + k' z)]^2 \leq \frac{\Sigma_{yz}' \sum_{zz}^{-1} \Sigma_{yz}}{\sigma_{yy}} \quad \left(\begin{matrix} \text{Bad} & \text{"linearity" between } y, z & \text{Good} \end{matrix} \right)$$

"0 ≤ ρ_{y(z)} ≤ 1"

$$= (\text{corr}(y, \beta_0 + \beta' z))^2 = \rho_{y(z)}^2$$

(If y, z has a joint dist., we can predict y based on z , and vice-versa.)

Now, suppose we have $\begin{bmatrix} y \\ z \end{bmatrix}$ is jointly $N_{r+1}(\mu, \Sigma)$, we observe

a sample of size n , $\begin{pmatrix} y_1 \\ z_1 \end{pmatrix}, \dots, \begin{pmatrix} y_n \\ z_n \end{pmatrix}$.

Based on these samples, we find MLE of Best linear predictor, smallest mean squared error, $\rho_{y(z)}$.

We compute the sample mean and sample covariance,

$$\hat{\mu} = \begin{pmatrix} \bar{y} \\ \bar{z} \end{pmatrix}, S = \begin{bmatrix} s_{yy} & s_{yz}' \\ s_{yz} & s_{zz} \end{bmatrix}, \text{ where } s_{yy} = \frac{1}{n-1} \sum_{i=1}^n (y_i - \bar{y})^2$$

$$(s_{yz})_k = \frac{1}{n-1} \sum_{i=1}^n (y_i - \bar{y})(z_{ik} - \bar{z}_k), \quad k=1, \dots, r$$

$$(s_{zz})_{kl} = \frac{1}{n-1} \sum_{i=1}^n (z_{ik} - \bar{z}_k)(z_{il} - \bar{z}_l), \quad k, l=1, 2, \dots, r$$

MLE of μ, Σ are $\hat{\mu}, \frac{n-1}{n} S$

* Recall (Invariance property)

If $\hat{\theta}$ is the MLE of θ , then for any function g ,

$g(\hat{\theta})$ is an MLE of $g(\theta)$.

Recall $\beta_0 = \mu_y - \Sigma_{yz}' \Sigma_{zz}^{-1} \mu_z, \beta = \Sigma_{zz}^{-1} \Sigma_{yz}$

So, $\hat{\beta}_0 = \bar{y} - \bar{z}' s_{zz}^{-1} \bar{z}$ and $\hat{\beta} = s_{zz}^{-1} s_{yz}$

* MLE for Best Linear predictor

$$= \hat{\beta}_0 + \hat{\beta}' \bar{z} = \bar{y} + \bar{z}' s_{zz}^{-1} (\bar{z} - \bar{z})$$

$$\rho_{y(z)}^2 = \frac{s_{yz}' s_{zz}^{-1} s_{yz}}{s_{yy}}$$

$$\hat{\rho}_{y(z)}^2 = \frac{s_{yz}' s_{zz}^{-1} s_{yz}}{s_{yy}}$$

MLE for smallest mean squared error $s_{yy} - s_{yz}' s_{zz}^{-1} s_{yz}$ is

$$\frac{n-1}{n} (s_{yy} - s_{yz}' s_{zz}^{-1} s_{yz})$$

* Multiple response case

$$\begin{bmatrix} y \\ z \end{bmatrix} \begin{matrix} m \times 1 \\ r \times 1 \end{matrix} \text{ have a joint distribution with mean } \mu = \begin{pmatrix} \mu_y \\ \mu_z \end{pmatrix} \text{ and } \text{cov} \begin{pmatrix} y \\ z \end{pmatrix} = \Sigma = \begin{bmatrix} \Sigma_{yy} & \Sigma_{yz}' \\ \Sigma_{yz} & \Sigma_{zz} \end{bmatrix} \begin{matrix} m \times m \\ r \times r \end{matrix}$$

A linear predictor of y based on z has the form

$$b_0 + b'z, \text{ where } b_0 \text{ is a constant vector and } b \text{ is } m \times r \text{ matrix.}$$

We minimize mean square and cross products of errors

$$E[(y - b_0 - b'z)(y - b_0 - b'z)'] = E[\varepsilon \varepsilon']$$

* The best linear predictor of y based on z is

$$\mu_y - \Sigma_{yz}' \Sigma_{zz}^{-1} \mu_z + \Sigma_{yz}' \Sigma_{zz}^{-1} z = \beta_0 + \beta' z = \mu_y + \Sigma_{yz}' \Sigma_{zz}^{-1} (z - \mu_z)$$

"smallest" (in the matrix sense: $A \geq B$ for matrices A, B if $A - B$ is n.n.d.)

* Mean square and products of errors

$$= \Sigma_{yy} - \Sigma_{yz}' \Sigma_{zz}^{-1} \Sigma_{yz} = \Sigma_{yy \cdot z}$$

$$= E[(y - \beta_0 - \beta'z)(y - \beta_0 - \beta'z)']$$

"Partial correlation coefficient" between y_k, y_l eliminating the effect of $\underline{z}$ is the (k, l) corr. coeff. corresponding to the covariance

w/ eliminating the effect of $\underline{z}$. matrix $\Sigma_{yy \cdot z}$

e.g. If $\text{corr}(y_k, y_l)$ is small, y_k, y_l are dependent on $\underline{z}$.

& $\text{corr}(y_k, y_l)$ is large
w/ the effect of $\underline{z}$

If $\text{corr}(y_k, y_l)$ is large, then y_k, y_l are not dependent on $\underline{z}$.
w/ eliminating the effect of $\underline{z}$

& $\text{corr}(y_k, y_l)$ is large they are ~~fairly~~ strongly correlated each other.
w/ the effect of $\underline{z}$

$$\Sigma_{yy \cdot z} = \begin{bmatrix} \sigma_{y_1 y_1 \cdot z} & \sigma_{y_1 y_2 \cdot z} & \dots & \sigma_{y_1 y_m \cdot z} \\ \vdots & \ddots & \ddots & \vdots \\ \sigma_{y_m y_1 \cdot z} & \dots & \dots & \sigma_{y_m y_m \cdot z} \end{bmatrix}$$

so, partial corr. coeff. between y_k, y_l eliminating the effect of $\underline{z}$

$$= \frac{\sigma_{y_k y_l \cdot z}}{\sqrt{\sigma_{y_k y_k \cdot z} \cdot \sigma_{y_l y_l \cdot z}}}$$

Math B7800

3/22/18, Thur

4/10, Tue "Exam 2"

or 4/12, Thur

Chapter 14 of the other text book.

Non-linear Regression

In linear regression model, $y = \beta_0 + \beta_1 x + \varepsilon$, ε is error. y is response, x is predictor.

$E[y] = \beta_0 + \beta_1 x$ (a linear function both in parameters and predictors)

In nonlinear regression, the relation between $E[y]$ and the parameters β_0, β_1 , and the predictor x is no longer linear.

"generalized linear model"

An important and useful subclass is the generalized linear model (GLM).

Here, the relation between $E[y]$ and β_0, β_1 , and x is via

a link function η s.t. $\eta(E[y]) = \beta_0 + \beta_1 x$.

e.g.

Logistic regression, Probit response function,

log-log mean response function.

In many examples, the response variable takes two possible values.

Ex1 While predicting the possibility of heart disease based on ¹diet and ²lifestyle measures.

Ex2 While predicting whether a married women will be in the work force based on ¹family background, ²education, ³household income, etc.

Ex3 While predicting whether a home owner will buy liability insurance.

Ex4 While predicting whether a company will have a legal division or not.

Ex5 While deciding whether to give loan/credit card.

Now, we assume y takes two values: 0 and 1.

Suppose $\pi = P(Y=1)$, so $P(Y=0) = 1-\pi$.

$$E[Y] = \pi, \text{ var}(Y) = \pi(1-\pi)$$

If we try to use linear model, $y_i = \beta_0 + \beta_1 x_i + \varepsilon_i$

Issues: $E[y_i] = \pi_i$ (ith observation)

$$1. \text{var}(\varepsilon_i) = \text{var}(y_i) = \pi_i(1-\pi_i) = E(y_i)(1-E(y_i))$$

(Variance depends on mean)

2. Possible values of ε_i : $1 - E(y_i)$, $0 - E(y_i)$

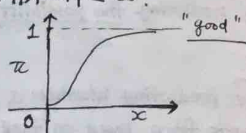
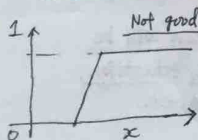
$$\Rightarrow 1 - \pi_i, -\pi_i$$

(Normal approximation for ε_i is not good).

$$* 3. \pi_i = E(y_i) = \beta_0 + \beta_1 x_i \leftarrow E(y_i) = 1\pi_i + 0(1-\pi_i) = \pi_i = P(Y_i=1)$$

Since π_i is a probability, $|\beta_0 + \beta_1 x_i| \leq 1$.

(π_i is dependent on x_i)



Ex y^c = Duration of pregnancy

X = Index for alcohol consumption.

Here, we can use linear model $y^c = \beta_0^c + \beta_1^c x + \varepsilon^c$, where $\varepsilon^c \sim N(0, \sigma^2)$

It is considered to be a premature delivery if $y^c \leq 38$ weeks, otherwise full term delivery.

Define

$$y = \begin{cases} 1 & \text{if } y^c \leq 38 \text{ (pre-term)} \\ 0 & \text{if } y^c > 38 \text{ (full term)} \end{cases}$$

$$\pi = P[Y=1] = P[y^c \leq 38] = P(\beta_0^c + \beta_1^c x + \varepsilon^c \leq 38)$$

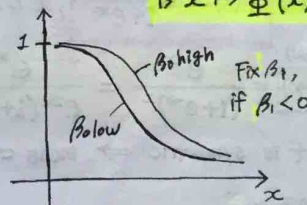
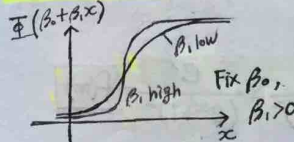
$$= P[\varepsilon^c \leq 38 - \beta_0^c - \beta_1^c x] = P\left[\frac{\varepsilon^c}{\sigma} \leq \frac{38 - \beta_0^c}{\sigma} - \frac{\beta_1^c}{\sigma} x\right]$$

$$= P[Z \leq \beta_0 - \beta_1 x] = \Phi(\beta_0 + \beta_1 x)$$

$$\text{so, } \Phi^{-1}(\pi) = \beta_0 + \beta_1 x$$

$$\Phi^{-1}(E(y)) = \beta_0 + \beta_1 x$$

This is called "probit mean response function" (probit transformation is $x \mapsto \Phi^{-1}(x)$).



Properties of the curve

1. Always lie between 0 and 1.

2. Monotone function.

3. For fixed β_0 , if β_1 increases, then curve becomes more sigmoid. (오차의 정도)

4. If β_1 is fixed and β_0 is changed, then the curve shifts.

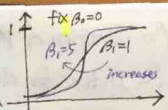
5. If $P(Y=1) = \pi = \Phi(\beta_0 + \beta_1 x)$, then

$$1 - \pi = P(Y=0) = 1 - \Phi(\beta_0 + \beta_1 x) = \Phi(-\beta_0 - \beta_1 x)$$

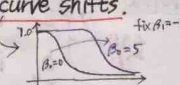
(Using the property $\Phi(z) + \Phi(-z) = 1$)

"(read) of meaning."

* This model is hard to interpret value of β_i (parameter) (probit model)



Sigmoid. (오차의 정도)



Using the property $\Phi(z) + \Phi(-z) = 1$

Logistic Regression

Logistic distribution is a continuous distribution with pdf

pdf $f(x) = \frac{e^x}{(1+e^x)^2}, x \in \mathbb{R}$

cdf $F(x) = \frac{e^x}{1+e^x}, x \in \mathbb{R}$

$$f(-x) = \frac{e^{-x}}{(1+e^{-x})^2} = \frac{e^x}{e^{2x}(1+e^{-x})^2} = \frac{e^x}{(e^x+1)^2} = f(x)$$

So, f is symmetric $\Rightarrow$ mean = 0.

$$e^{2x}(e^{-x} + 2e^{-x} + 1) = (1 + 2e^x + e^{2x}) = (e^x + 1)^2$$

Variance = $\int_{-\infty}^{\infty} x^2 \frac{e^x}{(1+e^x)^2} dx$

f is even iff $f(-x) = f(x)$ for all x in the domain of f .
"even function"

$$= 2 \int_0^{\infty} x^2 \frac{e^x}{(1+e^x)^2} dx$$

"by part"

$$\int u dv = uv - \int v du$$

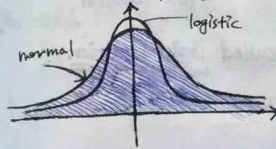
$u = x^2, dv = \frac{e^x}{(1+e^x)^2} dx$
 $v = -\frac{1}{(1+e^x)}$

$$= 2 \left[x^2 \left(-\frac{1}{1+e^x} \right) \Big|_0^{\infty} + \int_0^{\infty} \frac{2x}{1+e^x} dx \right]$$

$\Rightarrow 0$

$$= 4 \int_0^{\infty} \frac{x}{1+e^x} dx = 4 \int_0^{\infty} \log(1+e^{-x}) dx = \frac{\pi^2}{3}$$

check



Now, suppose $Y^c = \beta_0^c + \beta_1^c x + \epsilon^c$, where

ϵ^c has logistic distribution with mean 0, variance σ^2 .

$$Y = \begin{cases} 1 & \text{if } Y^c \leq 38 \\ 0 & \text{if } Y^c > 38 \end{cases}$$

$$\pi = P(Y=1) = P(Y^c \leq 38) = P(\epsilon^c \leq 38 - \beta_0^c - \beta_1^c x)$$

$$= P\left(\frac{\epsilon^c}{\sigma} \frac{\pi}{\sqrt{3}} \leq \frac{\pi}{\sqrt{3}} 38 - \frac{\pi}{\sqrt{3}} \beta_0^c - \frac{\pi}{\sqrt{3}} \beta_1^c x\right)$$

var = 1

$$= \frac{e^{\beta_0 + \beta_1 x}}{1 + e^{\beta_0 + \beta_1 x}} \leftarrow \text{cdf}$$

$$y = \frac{e^x}{1+e^x}, \text{ then } x = \log\left(\frac{y}{1-y}\right)$$

$$y + e^x y = e^x$$

$$\log y + x + \log y = x$$

$$\frac{y}{1-y} = \frac{e^x}{(1+e^x)-e^x} = e^x$$

$$\text{So, } \log\left(\frac{\pi}{1-\pi}\right) = \beta_0 + \beta_1 x$$

The function $y \rightarrow \log_e\left(\frac{y}{1-y}\right)$ is called logit function.

$$\frac{\pi}{1-\pi} = \frac{P(Y=1)}{P(Y=0)} \leftarrow \text{odds ratio.}$$

Math B7800

4/10/18, Tue

"4/26 - Exam II"
(Tue) Thur.

Generalized Linear Model

Logistic Regression / GLM

We have a binary response y and predictor x .

y takes two values: 0 or 1.

" $y = E(y) + \text{error}$ "

For the i th observation, $\pi_i = E(y_i) = P(y_i=1)$.

$y_i = \text{response}$, $x_i = \text{predictor}$

For GLM, $\pi_i = F(\beta_0 + \beta_1 x_i)$, where F is a function (usually F is the cdf of a distribution).

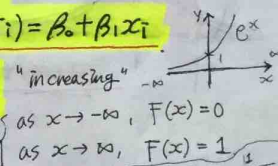
Different choices of F give different GLM. (pdf)

(For a random variable X , the cdf $F_X(x) = P(X \leq x)$)

If $F(x) = \Phi(x)$, where Φ is the cdf of $N(0,1)$. Then we have the probit regression model.

$$\pi_i = \Phi(\beta_0 + \beta_1 x_i) \Leftrightarrow \Phi^{-1}(\pi_i) = \beta_0 + \beta_1 x_i$$

When $F(x) = \frac{e^x}{1+e^x} = 1 - \frac{1}{1+e^x}$
"cdf"
of logistic distribution



Using this F , we get logistic regression model.

$$\pi_i = \frac{e^{\beta_0 + \beta_1 x_i}}{1 + e^{\beta_0 + \beta_1 x_i}} \Leftrightarrow \frac{1}{\pi_i} = 1 + \frac{1}{e^{\beta_0 + \beta_1 x_i}} \Leftrightarrow \frac{\pi_i}{1 - \pi_i} = e^{\beta_0 + \beta_1 x_i}$$

$$\frac{P(y_i=1)}{P(y_i=0)} \text{ "odds" } \Leftrightarrow \log\left(\frac{\pi_i}{1 - \pi_i}\right) = \beta_0 + \beta_1 x_i$$

$$\Leftrightarrow \log \left(\frac{\pi_i}{1-\pi_i} \right) = \beta_0 + \beta_1 x_i \quad \text{"odd ratio"}$$

is called odds = $\frac{P(Y_i=1)}{P(Y_i=0)} = \frac{\pi_i}{1-\pi_i}$

If $F(x) = 1 - e^{-e^x}$ $\rightarrow$ increasing
is the cdf of Gumbel dist.
as $x \rightarrow -\infty$, $F(x) = 0$
as $x \rightarrow \infty$, $F(x) = 1$

(If $y = 1 - e^{-e^x}$, then $x = \log(-\log(1-y))$)

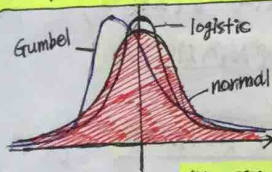
then $F^{-1}(\pi) = \log(-\log(1-\pi))$

• log-log regression function.

$$\text{Here } \pi_i = 1 - e^{-e^{\beta_0 + \beta_1 x_i}} = F(\beta_0 + \beta_1 x_i)$$

$$\Leftrightarrow \log(-\log(1-\pi_i)) = \beta_0 + \beta_1 x_i$$

The pdf of the "three distribution"



Logistic regression is the most popular among the three GLM.

The reason is (a) availability of software for computation;

(b) easy to interpret the parameters.

For logistic regression, $\log \left(\frac{\pi}{1-\pi} \right) = \beta_0 + \beta_1 x$. odd = $\frac{\pi}{1-\pi} = \frac{P(Y_i=1)}{P(Y_i=0)}$

β_1 captures the change in odds for each unit change in x .

Suppose π and $\tilde{\pi}$ correspond to the predictor value x and $x+1$.

$$\text{Then } \log \frac{\pi}{1-\pi} = \beta_0 + \beta_1 x$$

$$\log \frac{\tilde{\pi}}{1-\tilde{\pi}} = \beta_0 + \beta_1 (x+1)$$

$$\text{Subtract } \log \frac{\tilde{\pi}}{1-\tilde{\pi}} - \log \frac{\pi}{1-\pi} = \beta_1$$

$$= \log \frac{\tilde{\pi}/(1-\tilde{\pi})}{\pi/(1-\pi)}$$

$$\star \text{ Odd Ratio (OR)} = \frac{\tilde{\pi}/(1-\tilde{\pi})}{\pi/(1-\pi)} = e^{\beta_1}$$

If $\hat{\beta}_1$ is the MLE of β_1 , then $\hat{OR} = e^{\hat{\beta}_1}$

$\star$ How to find MLE of β_0, β_1 in logistic regression.

Need to find the "likelihood function" * joint dist. of y_i

$L(\beta_0, \beta_1) =$ joint distribution of the responses.

We have n -individuals with known response and predictor values.

y	x
y_1	x_1
$\vdots$	$\vdots$
y_n	x_n

independent

What is the pmf of y_i ?

$$P(Y_i=1) = \pi_i \Leftrightarrow P(Y_i=y) = \pi_i^y (1-\pi_i)^{1-y}$$

$$P(Y_i=0) = 1-\pi_i \quad g_i(y)$$

Then the joint pmf of $y_1, \dots, y_n$

$$g(y_1, \dots, y_n) = \prod_{i=1}^n g_i(y_i) = \prod_{i=1}^n \pi_i^{y_i} (1-\pi_i)^{1-y_i}$$

$$\log[g(y_1, \dots, y_n)] = \sum_{i=1}^n [y_i \log \pi_i + (1-y_i) \log (1-\pi_i)]$$

$$\begin{aligned} L(\beta_0, \beta_1) &= \sum_{i=1}^n [y_i \log \frac{\pi_i}{1-\pi_i} + \log (1-\pi_i)] \\ &= \sum_{i=1}^n [y_i (\beta_0 + \beta_1 x_i) - \log (1 + e^{\beta_0 + \beta_1 x_i})] \end{aligned}$$

$$\pi_i = \frac{e^{\beta_0 + \beta_1 x_i}}{1 + e^{\beta_0 + \beta_1 x_i}}$$

To obtain MLE $\hat{\beta}_0$ and $\hat{\beta}_1$ of

β_0 and β_1 , we need to maximize

$$\log L(\beta_0, \beta_1) = \sum_{i=1}^n y_i (\beta_0 + \beta_1 x_i) - \log (1 + e^{\beta_0 + \beta_1 x_i})$$

"No closed form expression"

Ex (Effect of programming experience in successfully completing projects in time)

$y_i = \begin{cases} 1 & \text{if } i\text{th candidate completed project in time} \\ 0 & \text{otherwise} \end{cases}$

x_i = amount of programming experience (in months), i is index of person

The dataset is for 25 candidates (available in book) ($i=1, \dots, 25$)

Logistic regression was used $\hat{\beta}_1 = 0.615$, $\hat{\beta}_0 = -3.05$.

So the fittest logistic regression model:

$$\log \frac{\pi}{1-\pi} = -3.05 + 0.615x = \beta_0 + \beta_1 x$$

In this case, $\hat{OR} = e^{\hat{\beta}_1} = e^{0.615} = 1.175$.

For two persons, P_1, P_2 , where P_2 has one month more experience.

$$\frac{\text{Odds}_2}{\text{Odds}_1} = 1.175 \Leftrightarrow \frac{\text{Odds}_2 - \text{Odds}_1}{\text{Odds}_1} = 0.175 = 17.5\%$$

So for each additional month of experience, the odds for success increase by 17.5%.
 $\uparrow$ ratio $P(y_i=1)$ per $P(y_i=0)$.

I have two future candidates having 10 months and 15 months of experiences.

$\nearrow (\frac{\circ}{\circ})$

Relative difference of odds

$$\begin{aligned} \hat{OR} &= e^{5\hat{\beta}_1}, \quad 5 = |15-10| \\ &= e^{0.8075} = 2.24 \end{aligned}$$

$$\frac{\text{Odds}_{15}}{\text{Odds}_{10}}$$

$$\frac{\text{Odds}_2 - \text{Odds}_1}{\text{Odds}_1} = 1.24$$

$$e^{3.075} =$$

$$\frac{0.615}{5} \times 2 = \boxed{3.075}$$

$$\frac{0.615}{5} \times 2$$

• Logistic Regression

Last time Example showing the interpretation of the parameter β_1 , connection with odds Ratio.

So far, we consider the case where one y value is present for each x value. But there can be multiple y values corresponding to one x value.

Ex (coupon distribution)

A company gives out 5 different coupons (\$5, \$10, \$15, \$20, \$30) to 200 customers for each category.

count how many customers have redeemed those coupons.

(\$ coupon amount	# customer given	# customers who redeemed
5 (x)	200	55 (y)
10	200	74
15	200	95
20	200	122
30	200	144

sum of
Bernoulli
Binomial

For $i=1, 2, \dots, 5$, $Y_i \sim \text{Bin}(200, \pi_i)$, where $\pi_i = \frac{e^{\beta_0 + \beta_1 x_i}}{1 + e^{\beta_0 + \beta_1 x_i}}$

So, $\log \frac{\pi_i}{1 - \pi_i} = \beta_0 + \beta_1 x_i = \pi_i'$

pmf of Y_i is $g_i(y_i) = P(Y_i = y_i) = \binom{n}{y_i} \pi_i^{y_i} (1 - \pi_i)^{n - y_i}$

$$= \binom{200}{y_i} \pi_i^{y_i} (1 - \pi_i)^{200 - y_i}$$

$$P(Y_i \leq y_i) = \sum_{k=0}^{y_i} \binom{n}{k} \pi_i^k (1 - \pi_i)^{n-k}$$

cdf

$$\log \pi_i = \frac{\log 1 - \log(1 - \pi_i)}{\log \pi_i}$$

$$\beta_0 + \beta_1 x_i$$

$$\log \pi_i = \frac{\log 1}{\log \pi_i}$$

$$\beta_0 + \beta_1 x_i$$

• Joint pmf of $Y_1, \dots, Y_5$

$$\theta(Y_1, \dots, Y_5) = \prod_{i=1}^5 \binom{200}{Y_i} \pi_i^{Y_i} (1-\pi_i)^{200-Y_i}$$

log: $\pi \mapsto \sum$

$$\log g(Y_1, \dots, Y_5) = \sum_{i=1}^5 \left[\log \binom{200}{Y_i} + Y_i \log \pi_i + (200-Y_i) \log (1-\pi_i) \right]$$

$$= \sum_{i=1}^5 \left[\log \binom{200}{Y_i} + Y_i \log \frac{\pi_i}{1-\pi_i} + 200 \log (1-\pi_i) \right]$$

So, the log-likelihood function "maximize this function"

$$\ell(\beta_0, \beta_1) = \log L(\beta_0, \beta_1)$$

$$= C + \sum_{i=1}^5 \left[Y_i (\beta_0 + \beta_1 x_i) - 200 \log (1 + e^{\beta_0 + \beta_1 x_i}) \right]$$

Next, we maximize the above function with respect to (β_0, β_1) using numerical method.

The maximizer will be the MLE. find $\hat{\beta}_0, \hat{\beta}_1$

• The Interpretation of β_1

$$\log \frac{\pi_i}{1-\pi_i} = \beta_0 + \beta_1 x_i$$

So, if $\tilde{x}_i = x_i + 1$ and $\tilde{\pi}_i, \pi_i$ are the corresponding probabilities.

$$\log \frac{\tilde{\pi}_i}{1-\tilde{\pi}_i} - \log \frac{\pi_i}{1-\pi_i} = \beta_1 (\tilde{x}_i - x_i) = \beta_1$$

$$= \log \frac{\tilde{\pi}_i / (1-\tilde{\pi}_i)}{\pi_i / (1-\pi_i)}$$

So, $OR = e^{\beta_1}, \hat{OR} = e^{\hat{\beta}_1} = \log(OR)$

EX Suppose $\hat{\beta}_1 = 0.01$. What is the % increase in odds if \$12 is offered instead of \$4.

The odds ratio estimated $\hat{OR} = e^{\hat{\beta}_1} = e^{0.08} = 1.083$

$$\frac{\text{Odds}_{12}}{\text{Odds}_4} = 1.083 \quad \text{So, } \frac{\text{Odds}_{12} - \text{Odds}_4}{\text{Odds}_4} = 0.083 \leftarrow$$

So, the odds will increase by 8.3% for any x_i .

π_i depends on x_i
Chance

(odds does not depend on x_i)

• Multiple logistic Regression

Here, we have $p-1$ predictors and one response y

Again, y takes two values: 0 or 1

x_i 's are be either categorical or continuous variables.

Ex (Disease occurrence)

$$Y = \begin{cases} 1 & \text{if individual infected} \\ 0 & \text{if not} \end{cases}$$

Four predictors

Two continuous (X_1 , Age, cholesterol level)
Two categorical (X_3, X_4 , Socio economic status, city sector)

Socio economic status

City sector

	middle - X_3	lower - X_4
upper	0	0
Middle	1	0
Lower	0	1

	X_5 - sector
sector 1	0
sector 2	1

so, the predictor $X = (1, X_1, \dots, X_5)'$

so, we have information about n -individuals

we know $(Y_i, X_i), i=1, \dots, n$

The logistic regression model is as follows.

$$Y_i \sim \text{Ber}(\pi_i), \text{ where } \pi_i = \frac{e^{X_i' \beta}}{1 + e^{X_i' \beta}}, \beta = \begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_5 \end{bmatrix}$$

so, the pmf of Y_i is $g_i(Y_i) = \pi_i^{Y_i} (1 - \pi_i)^{1-Y_i}$

so, the joint pmf $g(Y_1, \dots, Y_n) = \prod_{i=1}^n \pi_i^{Y_i} (1 - \pi_i)^{1-Y_i}$

$$\begin{aligned} \log g(Y_1, \dots, Y_n) &= \sum_{i=1}^n \left[Y_i \log \frac{\pi_i}{1 - \pi_i} + \log(1 - \pi_i) \right] \\ &= \sum_{i=1}^n \left[Y_i (X_i' \beta) - \log(1 + e^{X_i' \beta}) \right] \end{aligned}$$

Interpretation of the parameters

e^{β_i} is the odd ratio between two cases where ^①everything except the i th predictor is kept fixed and ^②the i th predictor value is increased by 1.

Suppose $\hat{\beta}_1$ is such that $e^{\hat{\beta}_1} = 1.0309$

this means that the odds for having the disease will go up by 3.09% for each year of age.

Suppose

$e^{\hat{\beta}_5} = 4.82$ (if change "section", odds 5 times. eg. 5 $\rightarrow$ 25)

then the odds of having the disease becomes approx. five time in sector 2 the odds in sector 1.

Inference about the parameters

The inference in this case is for large sample size

The MLE $\hat{\beta}$ approximately follows normal distribution with mean β and covariance $-H(\hat{\beta})$, where

$$H_{ij} = \frac{\partial^2 \log L(\beta)}{\partial \beta_i \partial \beta_j} \bigg|_{\hat{\beta}} \quad (\text{If } \hat{\beta} \text{ given, can you find covariance matrix?})$$

This will allow us to obtain confidence intervals and simultaneous CI for different components of β .

Math B7800

4/17/18, Tue

We are discussing logistic regression.

Ex In a simple logistic regression, $\log \frac{\pi}{1-\pi} = \beta_0 + \beta_1 x$
 if the estimate of β_0 and β_1 are -3 and 1.2 , i.e., $\hat{\beta}_0 = -3$, $\hat{\beta}_1 = 1.2$
 then what is the probability that among 100 future individuals
 having $x=5$, there will be at least 60 individuals with
 positive response.

$P(\text{an individual having } x=5 \text{ gives positive response})$

$$= \frac{e^{\beta_0 + 5\beta_1}}{1 + e^{\beta_0 + 5\beta_1}} = \pi(5)$$

so, # individuals among the future 100 individuals who will
 give positive response. - *

$$\sim \text{Bin}(100, \pi(5))$$

$$P(\# * \geq 60) = \sum_{k=60}^{100} \binom{100}{k} \pi(5)^k (1-\pi(5))^{100-k}$$

$$\widehat{P(\# * \geq 60)} \leftarrow \text{plug in } \hat{\beta}_0 = -3, \hat{\beta}_1 = 1.2$$

Recall: If $\hat{\theta}$ is an MLE of θ , then for any function g ,
 $g(\hat{\theta})$ is MLE of $g(\theta)$.

So, MLE of the probability is

$$\sum_{k=60}^{100} \binom{100}{k} (\hat{\pi}(5))^k (1-\hat{\pi}(5))^{100-k}, \text{ where}$$

$$\hat{\pi}(5) = \frac{e^{-3+1.2 \times 5}}{1+e^{-3+1.2 \times 5}} = \frac{e^3}{1+e^3}$$

Q2: Use normal approximation to approximate the estimate.

Recall: $\text{Bin}(n, p) \approx N(np, npq)$, $q = 1-p$

So the normal approximation is $P(N(100 \cdot \hat{\pi}(5), 100 \cdot \hat{\pi}(5)(1-\hat{\pi}(5)) \geq 60)$

$$= P(N(95.26, 4.5) \geq 60)$$

$$= 1 - \Phi\left(\frac{60-95.26}{\sqrt{4.5}}\right) \text{ standard deviation}$$

Inference for the parameters in logistic regression.

the inference is correct for large sample size $\hat{\beta}$ is approximately normal with mean β and covariance matrix $[-H(\hat{\beta})]^{-1}$ ← this is correct. check the last note.

Where $H(\beta)$ is the matrix of second derivatives

$$H_{ij}(\beta) = \frac{\partial^2 \log L(\beta)}{\partial \beta_i \partial \beta_j}, \quad i, j = 0, 1, \dots, p-1$$

Using this approximation,

We obtain CI for β_k , simultaneous CI for the parameters $\beta_0, \dots, \beta_{k-1}$.

Also, we can test $H_0: \beta_k = 0$ vs $H_1: \beta_k \neq 0$ - (a)

(b) $H_0: \beta_k = 0$ vs $H_1: \beta_k > 0$

(c) $H_0: \beta_k = 1$ vs $H_1: \beta_k \neq 1$

Q: CI for β_k ?

standard normal $N(0, 1)$.

$$100(1-\alpha)\% \text{ CI for } \beta_k = \hat{\beta}_k \pm Z_{\alpha/2} \sqrt{(-H(\hat{\beta}))^{-1}_{k,k}} \quad - (a)$$

How to test $H_0: \beta_k = 0$ vs $H_1: \beta_k > 0$

$$\text{Reject } H_0 \text{ at level } \alpha \text{ if } \hat{\beta}_k > Z_{\alpha} \sqrt{(-H(\hat{\beta}))^{-1}_{k,k}} \quad - (b)$$

How to test $H_0: \beta_k = 1$ vs $H_1: \beta_k \neq 1$

$$\text{Reject } H_0 \text{ at level } \alpha \text{ if } |\hat{\beta}_k - 1| > Z_{\alpha/2} \sqrt{(-H(\hat{\beta}))^{-1}_{k,k}} \quad - (c)$$

$H_0: \beta_k = 1$ vs $H_1: \beta_k < 1$?

$$\text{Reject } H_0 \text{ at level } \alpha \text{ if } \hat{\beta}_k < 1 - Z_{\alpha} \sqrt{(-H(\hat{\beta}))^{-1}_{k,k}}$$

• Likelihood Ratio Test

$q < p$,

$H_0: \beta_q = \beta_{q+1} = \dots = \beta_{p-1} = 0$ vs $H_1: H_0$ is false

$$\Delta = \text{likelihood ratio} = \frac{\sup_{H_0} L(\beta)}{\sup L(\beta)}$$

Note that under H_0 , we also have a logistic regression model

$$\log \frac{\pi}{1-\pi} = \beta_0 + \beta_1 X_1 + \dots + \beta_{q-1} X_{q-1}$$

using similar approach, we can find MLE of $\beta_0, \dots, \beta_{q-1}$ under H_0 .

call the MLE $\hat{\beta}_q$.

$$\text{so, } \Delta = \frac{L\left(\begin{bmatrix} \hat{\beta}_q \\ 0 \end{bmatrix}\right)}{L(\hat{\beta})} \quad \text{small} \rightarrow \text{reject } H_0.$$

$$-2 \log \Delta \approx \chi^2_{p-q}$$

reject H_0 if $-2 \log \Delta > \chi^2_{p-q, \alpha}$
at level α

• Model selection

Define $AIC = -2 \log L(\hat{\beta}) + 2(\# \text{ variables})$ ← minimize

We compare all possible subsets of predictors, when the total # of predictors is ≤ 30 or 40.

If the number of predictors is more than 40,

We should use stepwise methods for finding the optional model.

• Principal Component Analysis (PCA)

population principal component

suppose $\underline{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_p \end{bmatrix}$ is a vector consisting of p features of an individual.

We want to choose linear combinations (any combination of $\underline{x}$)

$$y_1 = a_1' \underline{x} \quad (a_1, \dots, a_p \text{ are } 1 \times 1 \text{ vector of constants})$$

so that $y_1, \dots, y_p$ are uncorrelated and

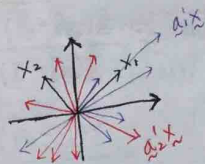
"population" $y_1, \dots, y_p$ have max possible variance.

Suppose Σ is covariance matrix of $\underline{x} = [x_1, \dots, x_p]'$

Then the constant vectors will depend on Σ .

These y_i 's are called (population) principal component of $\underline{x}$

y_1 is the first principal component, y_2 is the second and so on.



Math B7800

4/19/18, Thur

Question 3 from Exam 1

Let $n=50$, $R^2=0.6$.

$Y = X\beta + \varepsilon$, $H_0: \beta_1 = \dots = \beta_{10} = 0$ vs $H_1: H_0$ is false.

$$= \beta_0 + \beta_1 X_1 + \dots + \beta_{10} X_{10} + \varepsilon$$

$$Z = \begin{bmatrix} 1 & Z_{1,1} & \dots & Z_{1,10} \\ \vdots & \vdots & & \vdots \\ 1 & Z_{50,1} & \dots & Z_{50,10} \end{bmatrix} = [Z_1 | Z_2]$$

$$\beta = \begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_{10} \end{bmatrix} = \begin{bmatrix} \beta_{(1)} \\ \beta_{(2)} \end{bmatrix}$$

$$\text{So, } Z_1 = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix}, \quad Z_2 = \begin{bmatrix} Z_{1,1} & \dots & Z_{1,10} \\ \vdots & & \vdots \\ Z_{50,1} & \dots & Z_{50,10} \end{bmatrix}$$

So, $H_0: \beta_{(2)} = 0$ vs $H_1: \beta_{(2)} \neq 0$.

$$\text{LRT: } \frac{Y'(P_Z - P_{Z_1})Y}{Y'(I - P_Z)Y} \quad \begin{matrix} \text{Reject } H_0 \text{ if LRT is large} \\ \text{We would reject } H_0 \text{ at 5\% level} \end{matrix}$$

$$= \frac{Y'(P_Z - P_{Z_1})Y}{Y'(I - P_Z)Y} = R^2 \quad \text{if } R^2 > \frac{q}{40} F_{q,40}(0.05).$$

Principal Component Analysis (PCA)

Suppose $\underline{X}$ is a random vector with mean $\underline{0}$ and covariance matrix Σ .

$\underline{X}$ represents p features of the individuals in a population.
 p is usually large.

Goal: Choose fewer (less than p) linear combinations of features without losing much information.

$y_1, \dots, y_n$

$y_i = \underline{a}_i' \underline{X}$

Steps Find linear combinations $\underline{a}_1, \dots, \underline{a}_p$ s.t.

$\underline{Y} = \begin{bmatrix} \underline{a}_1' \\ \vdots \\ \underline{a}_p' \end{bmatrix} \underline{X}$ must have uncorrelated components and max possible variance.

y_i

$\underline{y}_{i1} = \underline{a}_{i1}' \underline{X}$

Step 1 • First principal component maximize $\underline{a}' \Sigma \underline{a} = \text{var}(\underline{a}' \underline{X})$ w.r.t. $\underline{a}$ s.t. $\underline{a}' \underline{a} = 1$.

We saw $\max_{\underline{a}: \underline{a}' \underline{a} = 1} \underline{a}' \Sigma \underline{a} = \lambda_1 \leftarrow$ largest eigen value of Σ

by spectral decomposition.

$\max_{\underline{a} \neq \underline{0}} \frac{\underline{a}' \Sigma \underline{a}}{\underline{a}' \underline{a}}$

$\underline{a}_1 =$ eigenvector of Σ corresponding to λ_1 achieves the max.

Recall: Spectral decomposition of real symmetric matrix.

$A = U \Lambda U'$ for any symmetric real matrix A , where U is orthogonal and Λ is diagonal.

The diagonal entries of Λ are the eigenvalues of A and $U U' = I$.

If $U = [\underline{u}_1 | \dots | \underline{u}_p]$

$$U \Lambda U' = [\underline{u}_1 | \dots | \underline{u}_p] \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_p \end{bmatrix} \begin{bmatrix} \underline{u}_1 \\ \vdots \\ \underline{u}_p \end{bmatrix} = \sum_{i=1}^p \lambda_i \underline{u}_i \underline{u}_i' = A$$

Let $\Sigma = U \Lambda U'$ be the spectral decomposition where

$\Lambda = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_p \end{bmatrix}$ is s.t. $\lambda_1 \geq \dots \geq \lambda_p \geq 0$ because Σ is always h.n.d.

Note: $\underline{u}_1, \dots, \underline{u}_p$ form a basis of $\mathbb{R}^p$.

so, any $\underline{a} \in \mathbb{R}^p$ can be expressed as $\underline{a} = x_1 \underline{u}_1 + x_2 \underline{u}_2 + \dots + x_p \underline{u}_p$ for some constants $x_1, \dots, x_p$.

$$\underline{a} \cdot \underline{u}_1 = x_1 \cdot 1 + x_2 \cdot 0 + \dots + x_p \cdot 0 = x_1$$

similarly, $x_i = (\underline{a} \cdot \underline{u}_i)$ for $i=1, \dots, p$.

Then any $\underline{a} \in \mathbb{R}^p$ can be written as $\underline{a} = \sum_{i=1}^p (\underline{a} \cdot \underline{u}_i) \underline{u}_i$

$$\text{Then } \underset{\sim u'}{a}' \underset{\sim u}{a} = \left(\sum_{i=1}^P (a \cdot u_i) u_i' \right) \left(\sum_{j=1}^P (a \cdot u_j) u_j \right)$$

$$= \sum_{i=1}^P (a \cdot u_i)^2$$

$$\underset{\sim u}{a}' \underset{\sim u}{a} = \left(\sum_{i=1}^P (a \cdot u_i) u_i' \right) \left(\sum_{j=1}^P \lambda_j u_j u_j' \right) \left(\sum_{k=1}^P (a \cdot u_k) u_k \right)$$

$$(a \cdot u_i) u_i' \lambda_j u_j u_j' (a \cdot u_k) u_k$$

$$= \dots u_i' u_j u_j' u_k = 0 \text{ if } i \neq j \text{ or } j \neq k$$

$$= \sum_{i=1}^P (a \cdot u_i) \lambda_i (a \cdot u_i) = \sum_{i=1}^P \lambda_i (a \cdot u_i)^2 \quad (*)$$

So, what does the problem reduce to?

$$\text{Call } x_i = (a \cdot u_i) \quad (*)$$

$$\max_{x_1, \dots, x_P} \frac{\lambda_1 x_1^2 + \dots + \lambda_P x_P^2}{x_1^2 + \dots + x_P^2} = \lambda_1 \frac{x_1^2}{x_1^2 + \dots + x_P^2} + \dots + \lambda_P \frac{x_P^2}{x_1^2 + \dots + x_P^2}$$

$$= \lambda_1$$

$$= \lambda_1 \alpha_1 + \dots + \lambda_P \alpha_P, \text{ where } \alpha_1 + \dots + \alpha_P = 1$$

This shows that $y_1 = u_1' \tilde{x}$ and $\text{var}(y_1) = \lambda_1$.

of eigen val
of Σ

Second principal component: $y_2 = \tilde{a}_2' \tilde{x}$ s.t. $\text{cov}(y_2, y_1)$

$$= \tilde{a}_2' \sum u_i u_i' u_1 \quad \leftarrow \text{eigen vector of } \Sigma$$

$$= \tilde{a}_2' (\lambda_1 u_1)$$

$$= \lambda_1 \tilde{a}_2' u_1 = 0$$

$$\text{So, } \tilde{a}_2' u_1 = 0$$

$$\underset{\sim u}{a}' \underset{\sim u}{a} = \tilde{a}_2' \tilde{a}_2$$

Now, we need to maximize $\frac{\tilde{a}_2' \tilde{a}_2}{\tilde{a}_2' \tilde{a}_2}$ subject to the constraint $\tilde{a}_2' u_1 = 0$ (normalize)

$$\tilde{a}_2 = \sum_{i=1}^P (a_2 \cdot u_i) u_i$$

$$\text{If } \tilde{a}_2' u_1 = 0, \text{ then } \tilde{a}_2 = \sum_{i=2}^P (a_2 \cdot u_i) u_i$$

$$\text{Also, } \tilde{a}_2' \tilde{a}_2 = \sum_{i=1}^P \lambda_i (a_2 \cdot u_i)^2 = \sum_{i=2}^P \lambda_i (a_2 \cdot u_i)^2$$

$$\frac{\tilde{a}_2' \tilde{a}_2}{\tilde{a}_2' \tilde{a}_2} = \frac{\lambda_2 (a_2 \cdot u_2)^2 + \dots + \lambda_P (a_2 \cdot u_P)^2}{(a_2 \cdot u_2)^2 + \dots + (a_2 \cdot u_P)^2}$$

So, using similar argument,

$$\max_{\tilde{a}_2 \neq 0} \frac{\tilde{a}_2' \tilde{a}_2}{\tilde{a}_2' \tilde{a}_2} = \lambda_2$$

$$\tilde{a}_2' u_1 = 0$$

$$\tilde{a}_2 = u_2 \quad \leftarrow \text{eigen vector satisfies}$$

$$\tilde{a}_2' u_1 = 0 \text{ and } \frac{\tilde{a}_2' \tilde{a}_2}{\tilde{a}_2' \tilde{a}_2} = \lambda_2$$

Thus, $y_2 = u_2' \tilde{x}$, and $\text{var}(y_2) = \lambda_2$.

↑
eigen vector
of Σ

↑
eigen value
of Σ

Following similar argument, $Y_i = u_i' X$, $i=1, \dots, p$

$$\text{Var}(Y_i) = \lambda_i \quad \text{trace}(\Delta)$$

$$\sum_{i=1}^p \text{Var}(Y_i) = \sum_{i=1}^p \lambda_i = \text{trace}(\Sigma) = \sum_{i=1}^p \sigma_{ii}$$

$$\text{tr}(\Sigma) = \text{tr}(P\Delta P') = \text{tr}(\Delta) = \sum_{i=1}^p \frac{\text{Var}(X_i)}{\sigma_{ii}} \quad \left(\begin{array}{l} \text{total} \\ \text{variance} \\ \text{doesn't change} \end{array} \right)$$

• Proportion of total explained by i th principal component (PC)

$$= \frac{\lambda_i}{\lambda_1 + \dots + \lambda_p} \quad \left(\text{this is called portion of variance explained by } i\text{th PC} \right)$$

We will choose first k PC if the corresponding explained variance $\frac{\lambda_1 + \dots + \lambda_k}{\lambda_1 + \dots + \lambda_p}$ is close to 1.

If X has mean μ , covariance Σ , then the (population) PCs are

$$Y_1 = u_1' (X - \mu), \dots, Y_p = u_p' (X - \mu),$$

where $u_1, \dots, u_p$ are the eigenvectors of Σ .

• PC for standardized random vector.

Suppose $X_{p \times 1}$ has mean 0, covariance Σ

$$\text{Define } Z_i = \frac{X_i - \mu_i}{\sqrt{\sigma_{ii}}}, \quad Z = \begin{bmatrix} Z_1 \\ \vdots \\ Z_p \end{bmatrix} \quad \left| \quad \begin{array}{l} \text{Cov}(Z_i, Z_j) \\ = \text{Cov}\left(\frac{X_i - \mu_i}{\sqrt{\sigma_{ii}}}, \frac{X_j - \mu_j}{\sqrt{\sigma_{jj}}}\right) \\ = \frac{\sigma_{ij}}{\sqrt{\sigma_{ii}} \sqrt{\sigma_{jj}}} = \rho_{ij} \end{array} \right.$$

$$E(Z) = 0,$$

$$\text{Cov}(Z) = \rho = \begin{bmatrix} 1 & \rho_{12} & \rho_{13} & \dots & \rho_{1p} \\ \rho_{21} & 1 & & & \\ \vdots & & \ddots & & \\ \rho_{p1} & \rho_{p2} & \dots & \dots & 1 \end{bmatrix} \quad \left| \quad \begin{array}{l} \text{Let } V = \text{Diag}(\Sigma) \\ = \begin{bmatrix} \sigma_{11} & & 0 \\ & \ddots & \\ 0 & & \sigma_{pp} \end{bmatrix} \end{array} \right.$$

$$= (V^{1/2})^{-1} \Sigma (V^{1/2})^{-1}$$

The PCs of Z are $\tilde{Y}_1, \dots, \tilde{Y}_p$, where

$$\tilde{Y}_i = \tilde{u}_i' Z, \quad \text{where } \tilde{u}_1, \dots, \tilde{u}_p \text{ are eigen vectors of } \rho$$

$$= \tilde{u}_i' (V^{1/2})^{-1} (X - \mu)$$

Principal Components

Last time population PC for standardized covariance.

• special covariance matrices

case 1: when Σ is diagonal, $\Sigma = \begin{bmatrix} \sigma_{11} & 0 \\ 0 & \ddots \\ 0 & \ddots & \sigma_{pp} \end{bmatrix}$

Here, the eigenvalues of Σ are $\sigma_{11}, \dots, \sigma_{pp}$, we order them in decreasing order $\sigma_{(1)} \geq \dots \geq \sigma_{(p)}$.

$\mathbf{X} = \begin{pmatrix} X_1 \\ \vdots \\ X_p \end{pmatrix}$ • First principal component = $(X_{(1)} - \mu_{(1)})$
 • The i th PC is $(X_{(i)} - \mu_{(i)})$

Ex If $\Sigma = \begin{bmatrix} 2 & 0 \\ 0 & 3 \end{bmatrix}$ and $\mathbf{X}$ has mean μ and cov Σ .

then $PC1 = X_3 - \mu_3$

$PC2 = X_1 - \mu_1$

$PC3 = X_2 - \mu_2$

Corresponding correlation matrix = $\mathbf{I}$. ← all eigenvalues are equal important.

case 2: All pairwise correlations of the components are equal to ρ .

Suppose $\text{Var}(X_i) = \sigma_{ii}$.

then $\text{cov}(X_i, X_j) = \begin{cases} \rho \sqrt{\sigma_{ii}} \sqrt{\sigma_{jj}} & \text{if } i \neq j \\ \sigma_{ii} & \text{if } i = j \end{cases}$

so,

$$\Sigma = \begin{bmatrix} \sigma_{11} & \rho \sqrt{\sigma_{11}} \sqrt{\sigma_{22}} & \dots & \rho \sqrt{\sigma_{11}} \sqrt{\sigma_{pp}} \\ \vdots & \vdots & \ddots & \vdots \\ \rho \sqrt{\sigma_{pp}} \sqrt{\sigma_{11}} & \dots & \dots & \sigma_{pp} \end{bmatrix}$$

corresponding correlation matrix $= \begin{bmatrix} 1 & \rho & \dots & \rho \\ \rho & 1 & \dots & \rho \\ \vdots & \vdots & \ddots & \vdots \\ \rho & \rho & \dots & 1 \end{bmatrix} = \rho$

Step 1: $\begin{bmatrix} 1 & \dots & 1 \\ \vdots & \ddots & \vdots \\ 1 & \dots & 1 \end{bmatrix}_{n \times n} = \mathbf{1}_n \mathbf{1}_n'$, $\mathbf{1}_n$ is $n \times 1$ vector of all 1

$\mathbf{1}_n' \mathbf{1}_n = n$
the eigenvalues of $\mathbf{1}_n \mathbf{1}_n'$ are $(n, 0, 0, \dots, 0)$

$(\mathbf{1}_n \mathbf{1}_n') \cdot \mathbf{1}_n = n \cdot \mathbf{1}_n$. So, $\mathbf{1}_n$ is an eigenvector of $\mathbf{1}_n \mathbf{1}_n'$ corresponding to eigenvalue n .

If $\mathbf{u}$ is an eigenvector of $\mathbf{1}_n \mathbf{1}_n'$ corresponding to 0, then

$(\mathbf{1}_n \mathbf{1}_n') \mathbf{u} = \mathbf{0} = \mathbf{1}_n (\mathbf{1}_n' \mathbf{u}) = (\mathbf{1}_n' \mathbf{u}) \mathbf{1}_n$. So, $\mathbf{1}_n' \mathbf{u} = 0$

Def the eigenspace of a matrix A corresponding to eigenvalue λ is $\{x: Ax = \lambda x\}$.

In our case, the eigenspace of $\mathbf{1}_n \mathbf{1}_n'$ corresponding to eigenvalue is $\text{Null}(\mathbf{1}_n')$.

$\dim(N(A)) = \dim \text{col}(A) - \text{rank}$.

In this case, $\dim(\text{eigenspace}) = n - 1$.

Any basis of $N(\mathbf{1}_n')$, or in particular, any set of mutually orthogonal vectors $u_1, \dots, u_{n-1}$ satisfying $\mathbf{1}_n' u_i = 0, i = 1, \dots, n-1$.

and be used as $PC_2, \dots, PC_n$ for $\mathbf{1}_n \mathbf{1}_n'$.

$\left[\frac{1}{\sqrt{2}}, \frac{-1}{\sqrt{2}}, 0, \dots, 0 \right]$ "normalized" "mutually orthogonal"

$\left[\frac{1}{\sqrt{2 \times 3}}, \frac{-1}{\sqrt{2 \times 3}}, \frac{-2}{\sqrt{2 \times 3}}, 0, \dots, 0 \right]$

$\left[\frac{1}{\sqrt{(i-1)i}}, \frac{1}{\sqrt{(i-1)i}}, \dots, \frac{1}{\sqrt{(i-1)i}}, \frac{-(i-1)}{\sqrt{(i-1)i}}, 0, 0, \dots, 0 \right] = e_i$

$\frac{(i-1) + (i-1)^2}{(i-1)i} = 1 = \text{norm}$

for $i = 2, 3, \dots, n$.

then $e_1 = \mathbf{1}_n, e_2, \dots, e_n$ are independent eigenvectors of $\mathbf{1}_n \mathbf{1}_n'$.

In the example, the correlation matrix was

$$\begin{bmatrix} 1 & \rho & \dots & \rho \\ \rho & 1 & & \\ \vdots & & \ddots & \\ \rho & \dots & & 1 \end{bmatrix}_{p \times p} = \rho \mathbf{1}_p \mathbf{1}_p' + (1-\rho) \mathbf{I}_p$$

If $\lambda_1, \dots, \lambda_p$ are the eigenvalues of $A_{p \times p}$ with corresponding eigenvectors $u_1, \dots, u_p$, then what are the eigenvalues and eigenvectors of $a\mathbf{I}_p + bA$?

$a + b\lambda_1, a + b\lambda_2, \dots, a + b\lambda_p$ are the eigenvalues with eigenvectors $u_1, \dots, u_p$.

$$A u_i = \lambda_i u_i$$

$$bA u_i = b\lambda_i u_i$$

$$a u_i + bA u_i = a u_i + b\lambda_i u_i = (a + b\lambda_i) u_i$$

$$(a\mathbf{I} + bA) u_i$$

So, the eigen values of $\rho \mathbf{1}_p \mathbf{1}_p' + (1-\rho) \mathbf{I}_p$ are

$$\rho p + (1-\rho), 1-\rho, 1-\rho, \dots, 1-\rho.$$

$$\parallel$$

$$1 + (p-1)\rho$$

$$\text{So, } PC_1 = \left(\frac{1}{\sqrt{p}}, \frac{1}{\sqrt{p}}, \dots, \frac{1}{\sqrt{p}} \right)' (X - \bar{X})$$

$$PC_i = \left(\frac{1}{\sqrt{(i-1)i}}, \frac{1}{\sqrt{(i-1)i}}, \dots, \frac{1}{\sqrt{(i-1)i}}, \frac{-(i-1)}{\sqrt{(i-1)i}}, 0, \dots, 0 \right)' (X - \bar{X})$$

for $i=2, \dots, p$.

proportion of variance explained by the first component.

total variance: sum of eigenvalues.

$$\text{total variance} = p.$$

$$\text{So, the proportion} = \frac{1 + (p-1)\rho}{p} = \rho + \frac{1-\rho}{p}$$

If ρ is large and/or p is large, then this proportion is quite large.

If $\rho = 0.8, p = 5$, then the proportion is 0.84.

So, it is enough to consider the first component.

Sample PC (Principal Components)

Suppose $X_1, \dots, X_n$ are n iid sample from a population with mean $\mu_{p \times 1}$ and covariance $\Sigma_{p \times p}$.

Obtain sample mean $\bar{X}$ and sample cov $S = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})(X_i - \bar{X})'$

then obtain eigenvalues $\hat{\lambda}_1 \geq \dots \geq \hat{\lambda}_p$ of S with corresponding eigenvectors $\hat{e}_1, \dots, \hat{e}_p$.

then the i th sample PC is $\hat{e}_i' (X - \bar{X})$

If S is the sample covariance matrix, then for any constant vector a , what is the sample variance of $a'X_1, \dots, a'X_n$?

$$\text{Ans } a'Sa$$

For two constant vectors a, b , sample cov between $a'X, b'X$ is $a'Sb$.

keeping these in mind, the sample PC maximizes sample variance subject to uncorrelated components.

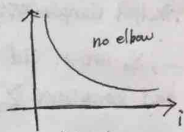
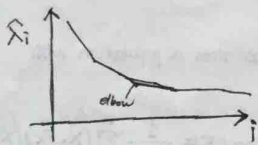
so, we want to find linear combinations $\hat{e}_1, \dots, \hat{e}_p$ s.t.

if $y_i = \hat{e}_i'(X - \bar{X})$, then the components $y = \begin{pmatrix} y_1 \\ \vdots \\ y_p \end{pmatrix}$

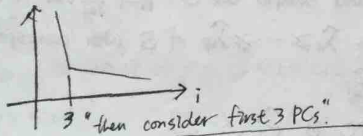
have zero sample covariance and max sample variance.

Number of PC

look at the scree plot ($i, \hat{\lambda}_i$) and look for 'elbow'.



"hard to conclude and find it"



"then consider first 3 PCs."

Math B7800

5/1/18, Tue

Final exam is similar with Exam 1. with 4 page sheets. (5/22) calculator.

Sample PC

$X_1, X_2, \dots, X_n$ are iid $p \times 1$ random vectors from a population with mean μ and cov Σ .

Obtain sample mean $\bar{X}$ and sample covariance matrix S .

Obtain eigenvalues, eigenvectors ($\hat{\lambda}_i, \hat{e}_i$) $i=1, \dots, p$ for the matrix S . then the sample PC are given by $\hat{\lambda}_1 \geq \dots \geq \hat{\lambda}_p$.

$\hat{e}_i'(X - \bar{X})$, $i=1, \dots, p$.

$$\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})'$$

Interpretation

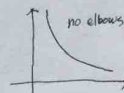
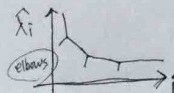
Total sample variance = $\text{trace}(S)$

$$= \sum_{i=1}^p \hat{\lambda}_i$$

$$= \sum_{i=1}^p \underbrace{\text{sample variance of } \hat{e}_i'(X - \bar{X})}_{= \hat{\lambda}_i}$$

The PC corresponding to small $\hat{\lambda}_i$ are dropped.

Either use Scree plot or a thumb rule applied to standardized sample PC.



"hard to conclude and find it"



"then consider first 3 PCs."

Standardized sample PC

Initially we have $X_1, \dots, X_n$.

sample mean $\bar{X}$, sample cov = S , $D = \text{Diag}(S)$.

$$\tilde{X}_i = D^{-1/2}(X_i - \bar{X}) \quad \text{standardized } X_i$$

$\tilde{X}_1, \dots, \tilde{X}_n$ are standardized version of X_i 's.

$$\bar{\tilde{X}} = 0 \quad \text{and} \quad \frac{1}{n-1} \sum_{i=1}^n \tilde{X}_i \tilde{X}_i' = \text{sample cov matrix of } \tilde{X}_1, \dots, \tilde{X}_n.$$

$$= R \quad (\text{sample correlation matrix for } X_1, \dots, X_n)$$

$\text{trace}(R) = P$, If $\hat{\lambda}_i, i=1, \dots, P$ are the eigenvalues of R .

$$\text{then } \sum_{i=1}^P \hat{\lambda}_i = P$$

The PC for which $\hat{\lambda}_i \leq 1$ are dropped.

Standardized sample PC are

$\hat{E}_i' Z = \hat{E}_i' D^{-1/2}(X - \bar{X})$, where $\hat{E}_i$ is an eigenvector corresponding to $\hat{\lambda}_i$ and $\hat{E}_1, \dots, \hat{E}_P$ are orthogonal.

(3) the distribution of $\hat{\lambda}_i$ and $\hat{E}_i, i=1, \dots, P$ are asymptotically independent.

X Suppose the spectral decomposition of $S_{3 \times 3}$

$$\begin{aligned} \hat{\lambda}_1 &= 10 & \hat{E}_1 &= \begin{bmatrix} 1/\sqrt{3} \\ 1/\sqrt{3} \\ 1/\sqrt{3} \end{bmatrix}, & \hat{E}_2 &= \begin{bmatrix} 1/\sqrt{2} \\ -1/\sqrt{2} \\ 0 \end{bmatrix}, & \hat{E}_3 &= \begin{bmatrix} 1/\sqrt{6} \\ 1/\sqrt{6} \\ -2/\sqrt{6} \end{bmatrix} \\ \hat{\lambda}_2 &= 5 \\ \hat{\lambda}_3 &= 1 \end{aligned}$$

Q: Find 95% CI for $\lambda_1, \lambda_2, \lambda_3$. (n is large), $n=80$.

$$\begin{aligned} \text{For } \lambda_1, \quad \frac{\sqrt{80}(10 - \lambda_1)}{-1.96} &\leq \frac{\sqrt{80}(10 - \lambda_1)}{\sqrt{2 \cdot 10^2}} \leq \frac{1.96}{\sqrt{2 \cdot 10^2}} \quad \alpha = 0.05 \\ &\downarrow \\ ? &\leq \lambda_1 \leq ? \end{aligned}$$

Sample PC estimate population PC.

Recall $Y_i = e_i' (X - \mu)$, $i=1, \dots, p$ where the population PC.

Recall $\text{cov}(X) = \Sigma$, Σ has eigenvalues $\lambda_1 \geq \dots \geq \lambda_p$, with corresponding eigenvectors $e_1, \dots, e_p$.

$$\text{cov}(Y) = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_p \end{bmatrix} = \Lambda, \quad \text{cov}(Y_i, Y_j) = e_i' \Sigma e_j = \begin{cases} 0 & \text{if } i \neq j \\ \lambda_i & \text{if } i = j \end{cases}$$

large sample inference

$$(1) \sqrt{n} \left[\begin{pmatrix} \hat{\lambda}_1 \\ \vdots \\ \hat{\lambda}_p \end{pmatrix} - \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_p \end{pmatrix} \right] \approx N_p(0, 2\Lambda^2)$$

find confidence interval

$$(2) \sqrt{n}(\hat{e}_k - e_k) \approx N_p(0, E_k), \text{ where } \hat{e}_k \leftarrow \text{rank } 1.$$

$X \sim N^p(0, \Sigma)$ rank is 1.

$$E_k = \sum_{i=1, i \neq k}^p \frac{\lambda_i \lambda_k}{(\lambda_i - \lambda_k)^2} e_i e_i' \leftarrow \text{rank } p-1 \Rightarrow \text{singular matrix}$$

Note that rank (E_k) is $p-1$.

This is fine because $\hat{e}_k$ satisfies $\sum_{i=1}^p (\hat{e}_{ki})^2 = 1$, and so the vector $\hat{e}_k$ lies on $p-1$ dimensional sphere.

Q: Find 99% CI for e_{21} .

$$-Z_{\alpha/2} \leq \frac{\sqrt{80} (1/\sqrt{2} - e_k)}{\sqrt{(\hat{E}_2)_{1,1}}} \leq Z_{\alpha/2}$$

$\alpha = 0.01$

$$E_2 = \sum_{i=1, i \neq 2}^3 \frac{\lambda_i \lambda_k}{(\lambda_i - \lambda_k)^2} e_i e_i'$$

$$= \frac{\lambda_1 \lambda_2}{(\lambda_1 - \lambda_2)^2} e_1 e_1'$$

$$+ \frac{\lambda_2 \lambda_3}{(\lambda_2 - \lambda_3)^2} e_3 e_3'$$

$$(\hat{E}_2)_{1,1} = \frac{10 \cdot 5}{5^2} \frac{1}{3} + \frac{5}{16} \frac{1}{6}$$

Q: Find 95% CR for $\begin{pmatrix} e_{21} \\ e_{22} \end{pmatrix}$.

$$\sqrt{80} \begin{pmatrix} \hat{e}_{21} - e_{21} \\ \hat{e}_{22} - e_{22} \end{pmatrix} \sim N_2\left(0, \begin{bmatrix} a & b \\ b & c \end{bmatrix}\right)$$

$$a = (\hat{E}_2)_{1,1}$$

$$\hat{E}_2 = \frac{10 \times 5}{5^2} \hat{e}_1 \hat{e}_1' + \frac{5 \times 1}{16} \hat{e}_3 \hat{e}_3'$$

$$\sqrt{80} A^{-1/2} \begin{pmatrix} \hat{e}_{21} - e_{21} \\ \hat{e}_{22} - e_{22} \end{pmatrix} \sim N_2(0, I_2)$$

$$80 \begin{pmatrix} \hat{e}_{21} - e_{21} \\ \hat{e}_{22} - e_{22} \end{pmatrix}' A^{-1} \begin{pmatrix} \hat{e}_{21} - e_{21} \\ \hat{e}_{22} - e_{22} \end{pmatrix} \approx \chi^2_2$$

$$\text{CR} = \left\{ \begin{pmatrix} e_{21} \\ e_{22} \end{pmatrix} : T \leq \chi^2_{2, 0.05} \right\}$$

Test $x_1, \dots, x_n$ iid with mean μ and cov Σ and correlation matrix ρ .

H_0 : All pairwise correlations are same $\rho = \rho_0 = \begin{bmatrix} 1 & \rho & \dots & \rho \\ \rho & 1 & & \\ & & \ddots & \\ \rho & \dots & \rho & 1 \end{bmatrix}$

H_1 : H_0 is false.

• Likelihood Ratio test at level α :

First obtain sample correlation matrix R , $R_{ij} = \frac{S_{ij}}{\sqrt{S_{ii}S_{jj}}}$

$\bar{r}_k$ = average of the off-diagonal entries in the k^{th} row of R
 $k = 1, \dots, p$.

$$\bar{r}_k = \frac{1}{p-1} \sum_{i \neq k} R_{ik}$$

$\bar{r}$ = average of all off diagonal = $\frac{1}{p(p-1)} \sum_{i \neq k} R_{i,k}$

$$\Gamma = \frac{(p-1)(1-(1-\bar{r})^2)}{p-(p-2)(1-\bar{r})^2}$$

$$T = \frac{p-1}{(1-\bar{r})^2} \left[\sum_{i \neq k} (\bar{r}_{ik} - \bar{r})^2 - \Gamma \sum_k (\bar{r}_k - \bar{r})^2 \right]$$

$$\stackrel{H_0}{\sim} \chi^2_{(p+1)(p-2)/2}$$

Math B7800

5/3/18/ Thur

ch.11 classification

In this set up, we suppose that the data (observation vector, feature vector) comes from one of the classes (or groups).

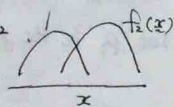
The class labels of the data vectors are unknown.

The set of possible class is known.

We would like to predict the class labels.

Two classes Suppose π_1 and π_2 denote the two possible classes that the data vectors x can come from.

The pdf of x is $f_1(x)$ if it is from π_1
 $f_2(x)$ if it is from π_2

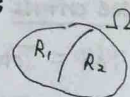


Ω = set of all possible values of x .

We want to divide Ω into two disjoint regions

R_1 and R_2

so that $x \in R_i$ will be classified to π_i .



Misclassification Probability

$$\begin{aligned} P(2|1) &= P(\text{observation from } \pi_1 \text{ is classified to } \pi_2) \\ &= P(\text{observation classified to } \pi_2 \mid \text{it came from } \pi_1) \\ &= \int_{R_2} f_1(x) dx \end{aligned}$$

Similarly,

$$P(1|2) = \int_{R_1} f_2(x) dx$$

In many cases, we have prior knowledge about the probability that x comes from π_1 or π_2 .

Let p_i be the prior prob. of π_i , i.e., $P(X \in \pi_i) = p_i$
 $p_1 + p_2 = 1$

★ One criteria to minimize

TMP = Total misclassification prob.

$$= P(x \text{ is misclassified})$$

$$= P(2|1)p_1 + P(1|2)p_2$$

$$= p_1 \int_{R_2} f_1(x) dx + p_2 \int_{R_1} f_2(x) dx$$

~~Cost~~

Cost involved in Misclassification

	π_1	π_2	"Actual"
your classification	π_1	$\circ$	$C(1 2)$
	π_2	$\circ$	$C(2 1)$

Average of Expected cost of Misclassification
 $ECM = E(\text{cost})$
 $= p_1 P(2|1) C(2|1) + p_2 P(1|2) C(1|2)$

$$= p_1 \int_{R_2} f_1(x) dx C(2|1) + p_2 \int_{R_1} f_2(x) dx C(1|2)$$

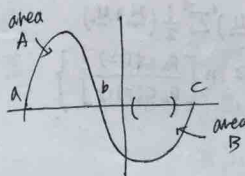
Assume that we know $p_1, p_2, C(1|2), C(2|1), f_1(x), f_2(x)$.

Goal: find R_1, R_2 so that ECM is minimized.

$$\int_{R_1} f_1(x) dx + \int_{R_2} f_1(x) dx = 1$$

$$\text{So, } ECM = p_1 \left(1 - \int_{R_1} f_1(x) dx \right) C(2|1) + p_2 \int_{R_1} f_2(x) dx C(1|2)$$

$$= \int_{R_1} [p_2 C(1|2) f_2(x) - p_1 C(2|1) f_1(x)] dx + p_1 C(2|1)$$



$\int_R h(x) dx$ is minimum or max
 when $R = [b, c]$ when $R = [a, b]$

$$-B \leq \int_R f(x) dx \leq A$$

So, using similar argument, ECM will be minimized

$$\text{if } R_1 = \left\{ x: \underbrace{\frac{f_1(x)}{f_2(x)}}_{\substack{\text{ratio} \\ \text{of pdf}}} \geq \underbrace{\frac{C(1|2)P_2}{C(2|1)P_1}}_{\substack{\text{ratio} \\ \text{of cost}}} \right\}$$

Suppose π_1 is $N_p(\mu_1, \Sigma)$ where $\mu_1 \neq \mu_2$
and π_2 is $N_p(\mu_2, \Sigma)$ but same covariance.

Assume μ_1, μ_2, Σ are known cov.

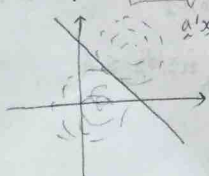
$$\text{Then } f_1(x) = \frac{1}{(2\pi)^{p/2} |\Sigma|^{1/2}} \exp\left[-\frac{1}{2}(x-\mu_1)'\Sigma^{-1}(x-\mu_1)\right]$$

$$f_2(x) = \frac{1}{(2\pi)^{p/2} |\Sigma|^{1/2}} \exp\left[-\frac{1}{2}(x-\mu_2)'\Sigma^{-1}(x-\mu_2)\right]$$

$$\frac{f_1(x)}{f_2(x)} = \exp\left[-\frac{1}{2}\{(x-\mu_1)'\Sigma^{-1}(x-\mu_1) - (x-\mu_2)'\Sigma^{-1}(x-\mu_2)\}\right]$$

$$= \exp\left[(\mu_1 - \mu_2)'\Sigma^{-1}\left(x - \frac{\mu_1 + \mu_2}{2}\right)\right]$$

$$\text{So, } R_1 = \left\{ x: \underbrace{(\mu_1 - \mu_2)'\Sigma^{-1}x - (\mu_1 - \mu_2)'\Sigma^{-1}\frac{\mu_1 + \mu_2}{2}}_{a'x - b \geq c} \geq \ln\left[\frac{P_2 C(1|2)}{P_1 C(2|1)}\right] \right\}$$



If μ_1, μ_2, Σ are not known, but we have a "training sample"

(n_1 observations $x_{1i}, i=1, \dots, n_1$ from π_1 and
 n_2 observations $x_{2j}, j=1, \dots, n_2$ from π_2).

Then we estimate μ_1 by $\bar{x}_1$
 μ_2 by $\bar{x}_2$.

Σ by S_{pooled} , where

$$S_{pooled} = \frac{n_1-1}{n_1-1+n_2-1} S_1 + \frac{n_2-1}{n_1-1+n_2-1} S_2 \text{ and}$$

$$S_j = \frac{1}{n_j-1} \sum_{i=1}^{n_j} (x_{ji} - \bar{x}_j)(x_{ji} - \bar{x}_j)'$$

We replace μ_1, μ_2, Σ by $\bar{x}_1, \bar{x}_2, S_{pooled}$ in R_1 .

Note: $n_1-1+n_2-1 \geq p$.

Suppose π_1, π_2 are $N_p(\mu_1, \Sigma_1), N_p(\mu_2, \Sigma_2)$, then

$$\frac{f_1(x)}{f_2(x)} = \frac{|\Sigma_2|^{1/2}}{|\Sigma_1|^{1/2}} \exp\left[-\frac{1}{2}(x-\mu_1)'\Sigma_1^{-1}(x-\mu_1) + \frac{1}{2}(x-\mu_2)'\Sigma_2^{-1}(x-\mu_2)\right] = \frac{f_1(x)}{f_2(x)}$$

so, (R_1)

$$R_1 = \left\{ x: \frac{1}{2}x'(\Sigma_2^{-1} - \Sigma_1^{-1})x + (\mu_1'\Sigma_1^{-1} - \mu_2'\Sigma_2^{-1})x - \frac{1}{2}(\mu_1'\Sigma_1^{-1}\mu_1 - \mu_2'\Sigma_2^{-1}\mu_2) + \frac{1}{2}\ln\frac{|\Sigma_2|}{|\Sigma_1|} \geq \ln\left[\frac{P_2 C(1|2)}{P_1 C(2|1)}\right] \right\}$$

This is a quadratic separator.

When $\mu_1, \mu_2, \Sigma_1, \Sigma_2$ are not known, but we have ~~a~~
a training sample, then estimate μ_i by $\bar{x}_i$ and
 Σ_i by S_i .

Math 87800

5/8/18, Tue

Classification

Recall that in this set up the observations are coming from
one of g possible groups: $\pi_1, \pi_2, \dots, \pi_g$.

case 1: The pdfs for different groups are known.

case 2: The pdfs' are unknown.

case 2.1: The pdfs' are parametric with unknown parameters.

case 2.2: The pdfs' are nonparametric.

In case 1, we minimize either TMP or ECM (Expected Cost of misclassification probability)

For $g=2$

$$\text{TMP} = P_1 \int_{R_2} f_1(x) dx + P_2 \int_{R_1} f_2(x) dx$$

P_i = prior probability of π_i .

$$\text{ECM} = P_1 c(2|1) \int_{R_2} f_1(x) dx + P_2 c(1|2) \int_{R_1} f_2(x) dx$$

Where

		predict	
		π_1	π_2
Actual	π_1	0	$c(2 1)$
	π_2	$c(1 2)$	0

The choice of R_1 and R_2 that minimizes ECM

$$R_1 = \left\{ x: \frac{f_1(x)}{f_2(x)} > \frac{C(1|2)P_2}{C(2|1)P_1} \right\}$$

Note: if $C(2|1) = C(1|2)$, then

minimizing ECM and TMP are equivalent.

$$R_2 = \left\{ x: \frac{f_1(x)}{f_2(x)} < \frac{C(1|2)P_2}{C(2|1)P_1} \right\}$$

Optimum error rate (OER) = smallest possible error of a classification rule.

Ex Suppose π_1 and π_2 denote $N_p(\mu_1, \Sigma), N_p(\mu_2, \Sigma)$ with known parameters. Suppose $C(2|1) = C(1|2), P_1 = P_2 = \frac{1}{2}$.

Find OER?

The classification rule that minimizes ECM or TMP is given by

$$R_1 = \left\{ x: \frac{f_1(x)}{f_2(x)} \geq 1 \right\}, R_2 = \left\{ x: \frac{f_1(x)}{f_2(x)} < 1 \right\}$$

$$\frac{f_1(x)}{f_2(x)} = \exp \left[-\frac{1}{2} (x - \mu_1)' \Sigma^{-1} (x - \mu_1) + (x - \mu_2)' \Sigma^{-1} (x - \mu_2) \right]$$

$$\frac{f_1(x)}{f_2(x)} \geq 1 \iff \log \frac{f_1(x)}{f_2(x)} \geq 0 \iff (\mu_1 - \mu_2)' \Sigma^{-1} x - (\mu_1 - \mu_2)' \Sigma^{-1} \frac{1}{2} (\mu_1 + \mu_2) \geq 0$$

$$\text{so, } R_1 = \left\{ x: (\mu_1 - \mu_2)' \Sigma^{-1} x \geq (\mu_1 - \mu_2)' \Sigma^{-1} \frac{1}{2} (\mu_1 + \mu_2) \right\}$$

$$R_2 = \left\{ x: (\mu_1 - \mu_2)' \Sigma^{-1} x < (\mu_1 - \mu_2)' \Sigma^{-1} \frac{1}{2} (\mu_1 + \mu_2) \right\}$$

$$\text{TMP} = P_1 P(2|1) + P_2 P(1|2)$$

$$P(2|1) = P(x \in R_2 | x \sim N(\mu_1, \Sigma))$$

$$P(x \in R_2), \text{ where } x \sim N(\mu_1, \Sigma)$$

$$\text{so, } (\mu_1 - \mu_2)' \Sigma^{-1} x \sim N((\mu_1 - \mu_2)' \Sigma^{-1} \mu_1, (\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2))$$

$$P(x \in R_2) = P((\mu_1 - \mu_2)' \Sigma^{-1} x < (\mu_1 - \mu_2)' \Sigma^{-1} \frac{1}{2} (\mu_1 + \mu_2))$$

$$= P \left[\frac{(\mu_1 - \mu_2)' \Sigma^{-1} (x - \mu_1)}{\sqrt{(\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2)}} < \frac{(\mu_1 - \mu_2)' \Sigma^{-1} \left[\frac{1}{2} (\mu_1 + \mu_2) - \mu_1 \right]}{\sqrt{(\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2)}} \right]$$

$$= P \left[Z < -\frac{1}{2} \sqrt{(\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2)} \right]$$

$$= \Phi \left(-\frac{1}{2} \Delta \right)$$

$$P(1|2) = P(x \in R_1) \text{ if } x \sim N(\mu_2, \Sigma)$$

$$= P \left[(\mu_1 - \mu_2)' \Sigma^{-1} x \geq (\mu_1 - \mu_2)' \Sigma^{-1} \frac{1}{2} (\mu_1 + \mu_2) \right], \text{ where}$$

$$(\mu_1 - \mu_2)' \Sigma^{-1} x \sim N((\mu_1 - \mu_2)' \Sigma^{-1} \mu_2, (\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2))$$

$$= P \left[\frac{(\mu_1 - \mu_2)' \Sigma^{-1} (x - \mu_2)}{\sqrt{(\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2)}} \geq \frac{(\mu_1 - \mu_2)' \Sigma^{-1} \left[\frac{1}{2} (\mu_1 + \mu_2) - \mu_2 \right]}{\sqrt{(\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2)}} \right]$$

$$= P \left[Z \geq \frac{1}{2} \Delta \right] = 1 - \Phi \left(\frac{1}{2} \Delta \right) = \Phi \left(-\frac{1}{2} \Delta \right) \leftarrow \text{"OER"}$$

$$= P[Z > \frac{1}{2} \Delta] = 1 - \Phi(\frac{1}{2} \Delta) = \Phi(-\frac{1}{2} \Delta)$$

$$\text{So, OER} = \Phi(-\frac{\Delta}{2})$$

$\Delta = \sqrt{(\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 - \mu_2)}$ is a "distance" between μ_1, μ_2 .

This is called Mahalanobis (PC Mahalanobis) distance between two means.

Case 2.1 pdfs are parametric with unknown parameters.

We have a training sample for which the group label is known. Then we estimate the unknown parameters using the training sample.

Then use the classification rule used in case 1 with the parameters replaced with their estimates.

Here, we use $\hat{R}_1, \hat{R}_2$ instead of R_1 and R_2 .

The corresponding rates are called Actual Error Rate (AER).

$$\text{AER} = P_1 \int_{\hat{R}_2} f_1(x) dx + P_2 \int_{\hat{R}_1} f_2(x) dx$$

AER is not known, as it involves the unknown pdfs' $f_1(x), f_2(x)$.

We estimate AER from the training data.

AER (APER, apparent error rate)

$$= \frac{n_{1m} + n_{2m}}{n_1 + n_2} = \text{proportion of misclassification within training sample.}$$

Training Data	Predicted		
	π_1	π_2	
Actual π_1	n_{1c}	n_{1m}	n_1
Actual π_2	n_{2m}	n_{2c}	n_2

Logistic Regression

Starting with training data for which group labels are known. We can treat the group labels as the response variable.

For the i th training sample,

$$y_i = \begin{cases} 1 & \text{if it is from } \pi_1 \\ 2 & \text{if it is from } \pi_2 \end{cases}$$

The data: $x_1, \dots, x_n$ ($P \times 1$)

logistic model

$$\text{logit}(P(Y=1)) = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p$$

$$= \log \left(\frac{P(Y=1)}{P(Y=0)} \right)$$

First estimate: $\beta_0, \beta_1, \dots, \beta_p$.

Then allocate x to π_1 if

$$\hat{\beta}_0 + \hat{\beta}_1 x_1 + \dots + \hat{\beta}_p x_p > 0.$$

Then allocate x to π_i if

$$\hat{\beta}_0 + \hat{\beta}_1 x_1 + \dots + \hat{\beta}_p x_p > 0.$$

When there are $g \geq 2$ populations $\pi_1, \dots, \pi_g$ with

pdf $f_1(x), \dots, f_g(x)$.

$C(k|i)$ = cost of misclassification of an observation coming from π_i to π_k .

$P_1, P_2, \dots, P_g$ are prior probabilities.

$$ECM(1) = C(2|1)f_2(x) + C(3|1)f_3(x) + \dots + C(g|1)f_g(x)$$

$$ECM(i) = \sum_{\substack{k=1 \\ k \neq i}}^g C(k|i)f_k(x)$$

$$ECM = P_1 ECM(1) + P_2 ECM(2) + \dots + P_g ECM(g).$$

Allocate x to π_k if

$$\sum_{\substack{i=1 \\ i \neq k}}^g P_i C(k|i) f_i(x) \text{ is the smallest.}$$

Math B7800

5/10/18, Thur

No date and time for variables on the project.

Last time

Classification with g groups/classes.

↓
supervised learning

$\pi_1, \pi_2, \dots, \pi_g$ are g groups with prior probabilities $P_1, P_2, \dots, P_g$.

The pdfs are $f_1(x), \dots, f_g(x)$

Case 1: pdfs are known.

$$ECM = P_1 ECM(1) + \dots + P_g ECM(g).$$

cost of misclassification

$C(k|i)$ = cost for misclassifying an individual from π_i to π_k .

$$ECM(1) = C(2|1) \cdot P(2|1) + C(3|1) \cdot P(3|1) + \dots + C(g|1) P(g|1)$$

$$ECM(i) = \sum_{\substack{k=1 \\ k \neq i}}^g C(k|i) P(k|i), i = 1, 2, \dots, g.$$

$$ECM = \sum_{i=1}^g P_i \sum_{\substack{k=1 \\ k \neq i}}^g C(k|i) P(k|i)$$

$$= P_1 ECM(1) + \dots + P_g ECM(g)$$

$P(k|i)$ = prob. of
classifying
an observation
from π_i to π_k

The classification rule that minimizes ECM will allocate x

to π_k if $\sum_{i=1}^g f_i(x) P_i C(k|i)$ is minimum.

$$x_k = \arg \min_{i \neq k} \sum_{i=1}^g f_i(x) P_i C(k|i)$$

When $g=2$

we allocate x to π_1 if $f_1(x)P_1C(2|1) \geq f_2(x)P_2C(1|2)$ minimum.

$$\Leftrightarrow \frac{f_1(x)}{f_2(x)} \geq \frac{C(1|2)P_2}{C(2|1)P_1}$$

When the costs $C(k|i)$ are all equal, then it is enough to compare $\beta_k = \sum_{i=1}^g P_i f_i(x)$.

Choose k such that $\beta_k = \min_i \beta_i$

$$\beta_k = \sum_{i=1}^g P_i f_i(x) - P_k f_k(x) = C - P_k f_k(x)$$

so, minimizing $\beta_k \Leftrightarrow$ maximizing $P_k f_k(x)$

Posterior prob. of a group

posterior prob. for group i

$$= P(\pi_i | x) = P(X \text{ comes from } \pi_i | X=x)$$

$$= \frac{P_i f_i(x)}{\sum_{k=1}^g P_k f_k(x)}$$

so, the allocation rule that maximizes the posterior prob.

will allocate x to π_i if $P_i f_i(x) = \max_k P_k f_k(x)$

When the costs are equal, then minimizing ECM

$\Leftrightarrow$ Maximizing posterior.

Ex If π_i is $N_p(\mu_i, \Sigma_i)$, $i=1, \dots, g$.

Suppose equal costs. then Allocate x to π_i if $P_i f_i(x)$ is the max.

Equivalently, $\log(P_i f_i(x))$ is the max.

$$\Leftrightarrow \log(P_i) - \frac{1}{2} \log(|\Sigma_i|) - \frac{1}{2} (x - \mu_i)' \Sigma_i^{-1} (x - \mu_i) \quad \text{"quadratic"}$$

is the max.

Ex If π_i is $N_p(\mu_i, \Sigma)$, $i=1, \dots, g$, ^{same Σ}

then allocate x to π_i if

$$\log(P_i) + \mu_i' \Sigma^{-1} x - \frac{1}{2} \mu_i' \Sigma^{-1} \mu_i \quad \text{"linear"}$$

is the max.

Chapter 12 Clustering (unsupervised learning).

No training data. $\rightarrow$ No fitted model.

example M12

Setup: We have a bunch of "objects", which we need to classify into different groups / clusters

Ex1 Community Detection.

Ex2 Segregating genes.

Ex3 Segregating patients / healthy

First, assume that our observations are points in $\mathbb{R}^p$.

$$\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^p \text{ e.g. } p=2$$



Two kinds of algorithms

1. Hierarchical

2. Non-Hierarchical

1. Let $D_{n \times n}$ be the distance matrix, $D_{ij} = \|\mathbf{x}_i - \mathbf{x}_j\|$

Step 0 We have singleton clusters $\{\mathbf{x}_1\}, \{\mathbf{x}_2\}, \dots, \{\mathbf{x}_n\}$

D = Distance matrix.

Step 1 find clusters U and V s.t. $\text{dist}(U, V)$ is the minimum among all distances of pairs of clusters.

Step 2 Merge clusters U and V . so delete U, V and add $U \cup V$.

[e.g. after step 0, if $\|\mathbf{x}_1 - \mathbf{x}_2\| = \min_{i \neq j} \|\mathbf{x}_i - \mathbf{x}_j\|$

then new cluster structure is $\{\mathbf{x}_1, \mathbf{x}_2\}, \{\mathbf{x}_3\}, \dots, \{\mathbf{x}_n\}$.

Step 3 Adjust the distance matrix.

Delete the rows and columns corresponding to clusters U and V .

Add a new row and column for $U \cup V$ in $D_{(n-1) \times (n-1)}$

Step 4 Repeat step 2 and 3 until you get one cluster $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$

Distance between clusters

(a) single link $\text{dist}(U, x) = \min \{d(x, y) : y \in U\}$

(b) average link $\text{dist}(U, x) = \text{dist}(x, \text{centroid of } U)$

2. Non-

• Nonheirarchical Clustering

(k means clustering algorithm)

Assume that there are k clusters.

$y_1, \dots, y_k$

Goal: Find out the location of k centers and

allocation of the points $x_1, \dots, x_n$ to

$A(x_1), \dots, A(x_n) \in \{y_1, \dots, y_k\}$ so that

$\sum_{i=1}^n \|x_i - A(x_i)\|$ is minimum among all possible centers

and all possible allocation rules. "Np hard"

so, "approximation"

Step 1 Start with any k groups.

Step 2 Obtain the k centroids these groups.
(means)

Step 3 Allocate each vector to its closest centroid.
This defines new k groups.

Repeat Step 2 and 3 until the sum of the distances between the vectors and centroids stabilize.