

# 437 Midterm 1

## Projection Matrix.

In the context of Linear Algebra,

Projection Matrix is the orthogonal  $P = P^T$

projection onto the column space of  $X$   $\downarrow$

s.t.  $P^2 = P$  (idempotent) and  $P$  is symmetric

1. [8] Let  $C = I_n - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T$  be the centering operator on  $\mathbb{R}^n$ .

(a) [2] Show that  $C$  is positive semi-definite (PSD).

(a).

先证  $C$  是

Projection

Matrix

$\downarrow$

Projection Matrix

是 PSD 的

WIS:  $X^T C X \geq 0, \forall X \in \mathbb{R}^n$

We claim  $C$  satisfies that

$C = C^T$  and  $C^2 = C$ .

For  $C^T = (I - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T)^T = I^T - \frac{1}{n} (\mathbf{1}_n \mathbf{1}_n^T)^T = I - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T = C$

And  $C^2 = (I - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T)(I - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T) = I^2 - \frac{2}{n} \mathbf{1}_n \mathbf{1}_n^T + \frac{1}{n^2} \mathbf{1}_n \mathbf{1}_n^T \mathbf{1}_n \mathbf{1}_n^T = I^2 - \frac{2}{n} \mathbf{1}_n \mathbf{1}_n^T + \frac{1}{n^2} \mathbf{1}_n (\mathbf{1}_n^T \mathbf{1}_n) \mathbf{1}_n^T = I^2 - \frac{2}{n} \mathbf{1}_n \mathbf{1}_n^T + \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T = I - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T = C$

(b) [2] Explain why the rank of  $C$  is  $n-1$ .

① Orthogonal Complement:

For  $S \subseteq \mathbb{R}^n$ ,  $S^\perp$  is the orthogonal

complement s.t.  $S^\perp = \{s \in S : \langle s, s \rangle = 0, \forall s \in S\}$

$\dim(\mathbb{R}^n) = \dim(S) + \dim(S^\perp)$

② Projection Complement.

For projection Matrix  $P$ , let  $Q = I - P$ ,  $Q$  is also a projection Matrix but projects to  $S^\perp$  for  $P$  projects to  $S$ .

$C = I_n - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T$  ( $I_n$  spans  $\mathbb{R}^n$ )

$= I_n - J$ , let  $J = \mathbf{1}_n \mathbf{1}_n^T = (\frac{1}{\sqrt{n}} \mathbf{1}_n)(\frac{1}{\sqrt{n}} \mathbf{1}_n^T) = \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T$ , For  $\|u\|=1$ , in the form

(c) [2] Determine the eigenvalues of  $C$  and their multiplicities. (Recall: the multiplicity of an eigenvalue  $\lambda$  is the number of linearly independent eigenvectors corresponding to  $\lambda$ .)

By (a),  $C^2 = C \Rightarrow \lambda^2 = \lambda = 1$

As  $\text{Rank}(C) = n-1$ , and  $\lambda = 1$  or  $0$ ,

We have  $1$ 's multiplicity is  $n-1$ .

$0$ 's multiplicity is  $1$ .

We conclude  $\frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T$  is

projection matrix that projects anything onto  $\text{span}(\mathbf{1})$ .

$\dim(\text{Image}(C)) = 1$

$\Rightarrow \dim(\text{Image}(C)) = n - \dim(\text{Image}(J))$

(d) [2] Show that  $C \mathbf{1}_n = 0$ , i.e.,  $\mathbf{1}_n$  is an eigenvector corresponding to the zero eigenvalue.

$C \mathbf{1}_n = (I_n - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T) \mathbf{1}_n$

$= \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix} - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^T \mathbf{1}_n = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix} - \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix} = 0 \Rightarrow C \mathbf{1}_n = 0 \cdot \mathbf{1}_n$

$\Rightarrow \mathbf{1}_n$  is the eigenvector of  $C$  with  $0$  to be its eigenvalue

其特征向量  
 $AV = \lambda V$   
特征向量对应的标量  
特征值

手动  
计算  
后验证

## SAMPLE MEAN

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

~~SAMPLE COV~~

$$\frac{1}{n} \sum_{i=1}^n \begin{pmatrix} x_i \\ y_i \end{pmatrix} = \begin{pmatrix} \bar{x} \\ \bar{y} \end{pmatrix}$$

## SAMPLE COV

$$S = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})^T = \frac{X^T C X}{n} \text{ For}$$

$C$  to be the centering operator. 对已center的, 则是  $\frac{X^T X}{n}$ .

2. [8] Let  $X = \begin{pmatrix} f_1 & f_2 & f_3 \end{pmatrix} = \begin{pmatrix} x_1 & x_2 & x_3 \\ \vdots & \vdots & \vdots \\ x_n & x_n & x_n \end{pmatrix} \in \mathbb{R}^{10 \times 3}$ , be a data matrix where rows correspond to observations and columns correspond to variables. Assume that  $X$  is column-centered (each column has mean 0) and column-orthogonal ( $X^T X = I_3$ ).

(a) [2] Find the sample mean vector  $\bar{x}$  and the sample covariance matrix  $S$  of  $X$ .

Mean:  $X = \begin{bmatrix} x_{11} & x_{12} & x_{13} \\ x_{21} & x_{22} & x_{23} \\ \vdots & \vdots & \vdots \\ x_{n1} & x_{n2} & x_{n3} \end{bmatrix} \Rightarrow \begin{bmatrix} \sum_{i=1}^n x_{i1} & \sum_{i=1}^n x_{i2} & \sum_{i=1}^n x_{i3} \\ \vdots & \vdots & \vdots \end{bmatrix}^T$

样本均值是每一列平均值相加之后的向量, 和行无关. 结果为 #column x 1 的 Vector.

$\bar{x} = \begin{bmatrix} \bar{x}_1 \\ \bar{x}_2 \\ \bar{x}_3 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} = \vec{0} \in \mathbb{R}^{3 \times 1}$

VARIANCE:  $S = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})^T = \frac{1}{n} I_3 = \begin{bmatrix} \frac{1}{10} & 0 & 0 \\ 0 & \frac{1}{10} & 0 \\ 0 & 0 & \frac{1}{10} \end{bmatrix}$

(b) [2] Suppose  $f_1$  represents measurements in meters,  $f_2$  in centimeters, and  $f_3$  in millimeters. Convert all measurements to centimeters by defining  $Y = \begin{pmatrix} 100f_1 & f_2 & 0.1f_3 \end{pmatrix} \in \mathbb{R}^{10 \times 3}$ . What is the sample covariance matrix of  $Y$ ?

①. Conversion:  $X \rightarrow Y \Rightarrow \begin{bmatrix} 100f_{1,1} & f_{2,1} & 0.1f_{3,1} \\ \vdots & \vdots & \vdots \\ 100f_{1,n} & f_{2,n} & 0.1f_{3,n} \end{bmatrix} = Y$

$\begin{bmatrix} f_{1,1} & f_{1,2} & f_{1,3} \\ \vdots & \vdots & \vdots \end{bmatrix} \begin{bmatrix} 100 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0.1 \end{bmatrix} = \begin{bmatrix} f_{1,1} \cdot 100 + f_{2,1} \cdot 0 + f_{3,1} \cdot 0 & f_{1,1} \cdot 0 + f_{2,1} \cdot 1 + f_{3,1} \cdot 0 & f_{1,1} \cdot 0 + f_{2,1} \cdot 0 + f_{3,1} \cdot 0.1 \\ \vdots & \vdots & \vdots \end{bmatrix}$

(c) [2] Let  $v \in \mathbb{R}^3$  be a vector with unit norm, i.e.,  $\|v\| = 1$ . Define  $z = Xv$ . Find the sample mean of  $z$ .

$$z = Xv, v \in \mathbb{R}^3, \|v\| = 1$$

Consider  $\bar{z} = \frac{1}{n} \sum_{i=1}^n z_i = \frac{1}{n} \sum_{i=1}^n Xv = \frac{1}{n} X \sum_{i=1}^n v = \frac{1}{n} X \cdot 0 = \bar{z} - 0 = \bar{z} = \frac{1}{n} Xv$

(d) [2] What is the sample variance of  $z$ ?

known centered

$$S_z = \frac{z^T z}{n} = \frac{v^T X^T X v}{n} = \frac{v^T I_3 v}{n} = \frac{\|v\|^2}{n} = \frac{1}{10}$$

$\Rightarrow z = z$   
 $\Rightarrow z$  is column centered,  $z$ 's mean is 0

$\Rightarrow \bar{z} = \bar{z}$   
 $\Rightarrow \bar{z} = 0$   
 $\Rightarrow \bar{z} = 0$

②. SAMPLE COV OF  $Y$ .

$$\frac{Y^T Y}{n} = \frac{D^T X^T X D}{n} = \frac{1}{10} D^T D = \frac{1}{10} \begin{bmatrix} 100 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0.1 \end{bmatrix} = \begin{bmatrix} 10 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0.01 \end{bmatrix}$$

$\Rightarrow D^2$  Ans

# Linear Transformation of MVN.

For  $y \sim N_n(\mu, \Sigma)$ , let  $z = Ay + b$   
then  $z \sim N_m(A\mu + b, A\Sigma A^T)$

TRANSPOSE  
矩阵转作.  $(AB)^T = B^T A^T$   
 $(A^T)^T = A$

3. [16] Consider the linear regression model

$$y = X\beta + \varepsilon, \quad \varepsilon \sim N_n(0, \sigma^2 I_n),$$

where  $X \in \mathbb{R}^{n \times p}$  is a fixed column-orthogonal data matrix ( $X^T X = I_p$ ),  $\beta \in \mathbb{R}^p$  is a vector of regression coefficients, and  $\varepsilon$  is Gaussian noise. Assume that  $\sigma^2$  is known.

(a) [2] What is the distribution of  $y$ ? Write it down in terms of  $X$ ,  $\beta$ , and  $\sigma^2$ .

$$E(y) = E(X\beta + \varepsilon) = X\beta + E(\varepsilon) = X\beta =$$

$\uparrow$  Constant in context  $\Rightarrow E(\text{Constant}) = \text{Constant}$

$$E(\varepsilon) = 0 \text{ by Assumption} \quad \text{Constant Variance} = 0$$

$$\text{Var}(y) = \text{Var}(X\beta + \varepsilon) = 0 + \text{Var}(\varepsilon) = \sigma^2 I_n$$

$y \sim N_n(X\beta, \sigma^2 I_n)$

(b) [2] Note that for column-orthogonal  $X$ , the regression coefficient vector is  $\hat{\beta} = X^T y$ . Show that the fitted values  $\hat{y} = X\hat{\beta}$  can be written as  $\hat{y} = Py$ , where  $P$  is a projection matrix. Specify the space onto which  $P$  projects.

Fitted value:  $\hat{y} = X\hat{\beta}$

$$\hat{y} = X\hat{\beta} = X(X^T y) = (XX^T)y = Py$$

For  $\hat{\beta}$ , Use Least Squares Estimates.

$$\hat{\beta} = (X^T X)^{-1} (X^T y) = (I_p)^{-1} (X^T y) = X^T y$$

$\uparrow$  Column Orthogonal

For  $(XX^T)(XX^T) = XX^T = P \Rightarrow \text{Idempotent}$

And  $P^T = (XX^T)^T = (X^T)^T X^T = XX^T = P$   
 $\Rightarrow$  Symmetric  
both satisfies  $\Rightarrow$  Projection matrix

MVN  
Linear Transformation

(c) [2] What is the distribution of  $\hat{y}$ ?

$$\hat{y} \sim N_n(X\beta, \sigma^2 I_n)$$

For  $E(\hat{y}) = P E(y)$  from (b).

$$E(\hat{y}) = P(X\beta) = (XX^T)(X\beta) = X\beta = E(y)$$

$Ay + b, b=0, A=P$

(d) [2] Show that the residuals  $r = y - \hat{y}$  can be expressed as  $r = P_{\perp} y$  for some projection matrix  $P_{\perp}$ . Specify the space onto which  $P_{\perp}$  projects.

For  $\text{Var}(\hat{y}) = P(\text{Var}(y))P^T$   $\downarrow$  MVN variance  $A\Sigma A^T, A=P \text{ here.}$

$$= P(\sigma^2 I_n)P^T = \sigma^2 (P I_n P^T) = \sigma^2 P P^T = \sigma^2 P^2 = \sigma^2 XX^T$$

$\uparrow P^2 = P$   
 $\uparrow P^T = P$

So  $\hat{y} \sim N_n(X\beta, \sigma^2 XX^T)$

$C(X)$ : Column Space:  $X$  所有列向量可组成的空间

(d): RESIDUAL = OBSERVED - FITTED =  $y - \hat{y} = r$   
Response Response

$r = y - \hat{y} = I_n y - P y = (I_n - P)y$ , For  $P = XX^T, I = (I_n - XX^T)y \Rightarrow P_{\perp} = (I_n - XX^T)$

$P_{\perp}$  projects onto  $C(X)^{\perp}$

$\hookrightarrow$  91 右说. 正交补空间的投影

①  $XX^T = P \Rightarrow I - XX^T = P_{\perp}$

(互为正交补)  
ORTHOGONAL COMPLEMENT

②

(e) [2] What is the distribution of the residuals  $r$ ?

$$E(r) = E(P_{\perp}y) = E[(I_n - P)y] = (I_n - P)E(y) = (I_n - P)(X\beta) = X\beta - XX^T X\beta = 0$$

$$Var(r) = (I_n - P)Var(y)(I_n - P)^T = (I_n - P)(\sigma^2 I_n)(I_n - P)^T = \sigma^2 (I_n - P)^2 \xrightarrow{\text{idempotent}} \sigma^2 (I_n - P) = \sigma^2 (I_n - XX^T)$$

(Matrix Form)  
 $Var(AY) = AVar(Y)A^T$

So  $r \sim N_n(0, \sigma^2(I_n - XX^T))$

(f) [2] Explain why the residuals  $r$  and the fitted values  $\hat{y}$  are independent. (MVN.)

$r$  and  $\hat{y}$  independent

$\Downarrow$

WTS:  $Cov(r, \hat{y}) = 0$

$Cov(Ay, By) = AVar(y)B^T$

UNI  
VARI  
ATE

$Cov(X, Y) = E[(X - E(X))(Y - E(Y))^T]$

$Cov(\hat{y}, r) = Cov(Py, P_{\perp}y) = PVar(y)P_{\perp}^T = P(\sigma^2 I_n)P_{\perp}^T$

$= \sigma^2 P(I_n - P) = 0 \Rightarrow \hat{y} \text{ and } r \text{ are uncorrelated, thus independent.}$

(PP<sup>T</sup>=0)

(g) [2] Write down the joint distribution of  $\begin{pmatrix} \hat{y} \\ r \end{pmatrix}$ .

$\begin{pmatrix} \hat{y} \\ r \end{pmatrix} \sim N_{2n} \left( \begin{pmatrix} X\beta \\ 0 \end{pmatrix}, \sigma^2 \begin{pmatrix} XX^T & 0 \\ 0 & I_n - XX^T \end{pmatrix} \right)$

(h) [2] What is the conditional distribution of  $r|\hat{y}$ ?

$r \text{ and } \hat{y} \text{ are independent}$   
 $\Rightarrow r|\hat{y} \sim N_n(0, \sigma^2(I_n - XX^T))$

[END OF EXAMINATION; total points = 32]

# MIDTERM 2

①. Multivariate Normal Distribution

②. MVN Linear Transformation

③. Matrix Calculus.

④. For vector  $\beta$ ,  $\beta^T A \beta$ 's derivative,  $\partial_{\beta} [\beta^T A \beta]$  is  $2A\beta$ .

①.  $f(y) = (2\pi)^{-\frac{n}{2}} |\Sigma|^{-\frac{1}{2}} e^{-\frac{1}{2}(y-\mu)^T \Sigma^{-1}(y-\mu)}$

②. For constant  $A \in \mathbb{R}^{p \times n}$  and  $y \sim N_n(\mu, \Sigma)$ ,  $z = Ay \sim N(A\mu, A\Sigma A^T)$

is  $2A\beta$ .

1. [8] Consider the linear regression model

$$y = X\beta + \epsilon,$$

where  $X = \begin{pmatrix} -x_1^T \\ \vdots \\ -x_n^T \end{pmatrix} \in \mathbb{R}^{n \times p}$  is a fixed column-orthogonal design matrix ( $X^T X = I_p$ ),  $\beta \in \mathbb{R}^p$  is a vector of unknown regression coefficients, and  $\epsilon \sim N_n(0, \sigma^2 I_n)$  is a Gaussian noise vector. Assume that  $\sigma^2$  is fixed and known. From Midterm 1, we know that the response vector follows a multivariate normal distribution:

$$y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \sim N_n(X\beta, \sigma^2 I_n).$$

(a) [2] Write down the log-likelihood function  $\ell(y_1, \dots, y_n; \beta)$ , and explain why finding the maximum likelihood estimate (MLE) for  $\beta$  is equivalent to solving the least-squares problem, i.e., minimizing  $\|y - X\beta\|^2$  with respect to  $\beta$ .

MLE  $\Leftarrow$  Least Squares.

Goal of optimization is the same

For model assumption,  $y \sim N_n(X\beta, \sigma^2 I_n)$ ,  $\mu = X\beta$  and  $\Sigma = \sigma^2 I_n$ ,  $\det \Sigma = |\sigma^2 I_n| = (\sigma^2)^n$

$$f(y; \beta) = (2\pi)^{-\frac{n}{2}} (\sigma^2)^{-\frac{n}{2}} e^{-\frac{1}{2}(y-X\beta)^T (\frac{1}{\sigma^2} I_n) (y-X\beta)}$$

with  $\Sigma^{-1} = (\sigma^2 I_n)^{-1} = \frac{1}{\sigma^2} I_n$

$$\ell(y; \beta) = -\frac{n}{2} \ln(2\pi) - \frac{n}{2} \ln(\sigma^2) - \frac{1}{2\sigma^2} \|y - X\beta\|^2$$

↑  
Constant

↑  
Assumption Constant

$\Rightarrow$  Maximize  $(\|y - X\beta\|^2)$  to maximize  $\ell$

↓  
Minimize  $\|y - X\beta\|^2$  in OLS.

□

TRANSPOSE

$$(A-B)^T = A^T - B^T$$

$$(AB)^T = B^T A^T$$

(b) [2] From the previous part, we conclude that the MLE is  $\hat{\beta} = X^T y$ . What is the distribution of  $\hat{\beta}$ ?

$$\begin{aligned} \partial_{\beta} [(y-X\beta)^T (y-X\beta)] &= \partial_{\beta} [y^T y - 2y^T X\beta + \beta^T (X^T X) \beta] \Rightarrow 2I_p \beta = 2y^T X \\ &\Rightarrow \beta = y^T X \quad (\text{Final Result}) \end{aligned}$$

$$= y^T y - y^T X \beta - \beta^T X^T y + \beta^T X^T X \beta$$

$$\begin{aligned} y &: n \times 1 \\ X &: n \times p \\ \beta &: p \times 1 \end{aligned}$$

$\Rightarrow 1 \times 1$  For  $y^T X \beta \Rightarrow$  result is scalar

Transpose of scalar is still scalar  $\Rightarrow$  Merge into 1.

VARIANCE  
MTRX FORM

$$\text{Var}(Ay) = A \text{Var}(y) A^T$$

$$\begin{aligned} \hat{\beta} &= X^T y \\ E(\hat{\beta}) &= E(X^T y) = X^T E(y) = X^T X \beta = \beta \\ &\Rightarrow \hat{\beta} \text{ is unbiased estimator.} \end{aligned}$$

$$\begin{aligned} \text{Var}(\hat{\beta}) &= \text{Var}(X^T y) = X^T \text{Var}(y) X \\ &= X^T \sigma^2 I_n X = \sigma^2 X^T X = \sigma^2 I_p \\ &\Rightarrow \hat{\beta} \sim N_p(\beta, \sigma^2 I_p) \end{aligned}$$

# HYPOTHESIS TESTING

1. test statistic in context. <sup>IR</sup> MONO

MULTI <sup>IR</sup>  $R^n$

$$z_j = \frac{\hat{\beta}_j - 0}{\sqrt{\hat{\sigma}^2}} = \frac{\hat{\beta}_j}{\hat{\sigma}} \sim N(0,1)$$

$$\hat{\sigma}^2 = \left( \frac{\hat{\beta}_j}{\hat{\sigma}} \right)^2 = \hat{\beta}_j^T (G^2)^{-1} \hat{\beta}_j$$

$$D^2 = (X - \mu)^T Z^{-1} (X - \mu)$$

$$\text{test statistic: } (\hat{\beta} - 0)^T (G^2 I_p)^{-1} (\hat{\beta} - 0)$$

in context

(c) [2] Suppose you would like to test the hypothesis

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

$$\Rightarrow \frac{1}{\hat{\sigma}^2} \hat{\beta}^T I_p \hat{\beta} = \frac{\|\hat{\beta}\|^2}{\hat{\sigma}^2} = \sum_{i=1}^p \left( \frac{\hat{\beta}_i}{\hat{\sigma}} \right)^2$$

Compute the test statistic corresponding to this hypothesis and indicate what distribution this statistic follows under the null hypothesis.

这部分在 Multivariate Testing 里, MVN note

$$f^2 = (\hat{\beta} - 0)^T (G^2 I_p)^{-1} (\hat{\beta} - 0) = \frac{\|\hat{\beta}\|^2}{\hat{\sigma}^2} \sim f^2(p)$$

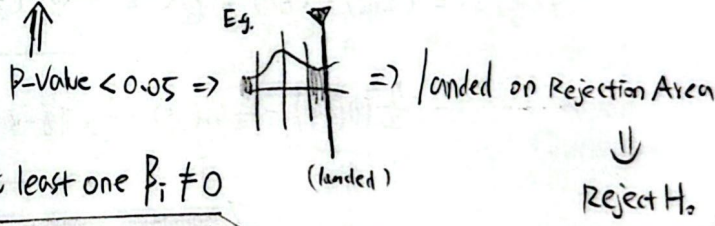
cl.f. = p

(Chi-square Distribution)

$$\text{test statistic} = \frac{\text{observed Bias}}{\text{Environmental Noise}}$$

(d) [2] Suppose the p-value you obtained is less than 0.05. What can we conclude about the components

$$\text{of } \beta = \begin{pmatrix} \beta_1 \\ \vdots \\ \beta_p \end{pmatrix}?$$



Reject  $H_0 \Leftrightarrow$  At least one  $\beta_i \neq 0$

Elas-Mixed signals (kurtosis will change):

$$y_1 \sim (m_1, \sigma_1^2), y_2 \sim (m_2, \sigma_2^2), \text{ w.r. } K(y_1, y_2) = \frac{\sigma_1^4 K(y_1) + \sigma_2^4 K(y_2)}{(\sigma_1^2 + \sigma_2^2)^2}$$

$$K(\alpha_1 z_1 + \alpha_2 z_2) = \alpha_1^4 K(z_1) + \alpha_2^4 K(z_2), \text{ Assume } y_1 \text{ and } y_2 \text{ are independent.}$$

①. Centering, kurtosis is translation invariant, i.e.  $K(y) = K(y - m)$ , we can assume  $m_1 = 0, m_2 = 0$  WLOG.

Let  $\tilde{y}_i = y_i - m_i$ , we have  $\tilde{y}_i$ 's mean = 0 and variance is  $\sigma_{\text{sum}}^2 = \sigma_1^2 + \sigma_2^2$  by independence. For  $S = y_1 + y_2$ ,

By definition,  $K(y) = E\left[\frac{(y-m)^4}{\sigma^4}\right] - 3$  and for zero-mean variables,  $\tilde{y} = K(\tilde{y}) = \frac{E(\tilde{y}^4)}{(E(\tilde{y}^2))^2} - 3 = \frac{E(S^4)}{E^2(S)}$

Expanding  $(y_1 + y_2)^4$ , gives  $y_1^4 + 4y_1^3 y_2 + 6y_1^2 y_2^2 + 4y_1 y_2^3 + y_2^4 \Rightarrow E(\text{polynomial})$

分子与分母本质是在对 Varian 进行操作.

For  $\tilde{y}$ , we have  $E(y^4) = (K(y) + 3) \sigma^4$

$$E(y_1^4) = (K_1 \sigma_1^4 + 3 \sigma_1^4)$$

$$E(y_2^4) = (K_2 \sigma_2^4 + 3 \sigma_2^4)$$

$$E(4y_1^3 y_2) = 4 E(y_1^3) E(y_2) = 0$$

$$E(6y_1^2 y_2^2) = 6 E(y_1^2) E(y_2^2) = 6 \sigma_1^2 \sigma_2^2$$

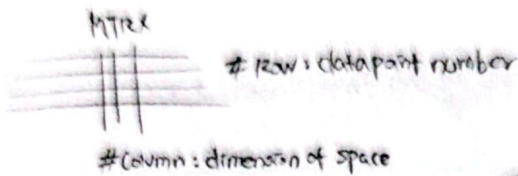
$$E(4y_1 y_2^3) = 4 E(y_1) E(y_2^3) = 0$$

$$E(y_2^4) = (K_2 \sigma_2^4 + 3 \sigma_2^4)$$

$$E(S^4) = K_1 \sigma_1^4 + K_2 \sigma_2^4 + 3(\sigma_1^4 + \sigma_2^4) + 6 \sigma_1^2 \sigma_2^2$$

$\leftarrow$  使用完全平方公式化简

$$K(S) = \frac{K_1 \sigma_1^4 + K_2 \sigma_2^4 + 3(\sigma_1^4 + \sigma_2^4)}{(\sigma_1^2 + \sigma_2^2)^2} - 3 = \frac{K_1 \sigma_1^4 + K_2 \sigma_2^4}{(\sigma_1^2 + \sigma_2^2)^2} + 3 - 3 = \frac{\sigma_1^4 K(y_1) + \sigma_2^4 K(y_2)}{(\sigma_1^2 + \sigma_2^2)^2}$$



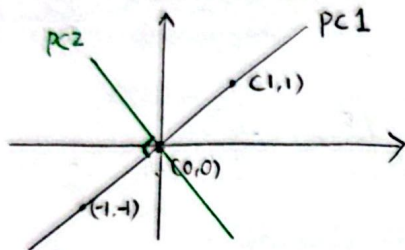
Row: data point  $\begin{cases} (0,0) \\ (1,1) \\ (-1,-1) \end{cases} \Rightarrow \text{Flattened into } \rightarrow 2D \text{ plane.}$

2. [8] Consider the following dataset with 3 observations and 2 features:

$$X = \begin{pmatrix} 0 & 0 \\ 1 & 1 \\ -1 & -1 \end{pmatrix} \in \mathbb{R}^{3 \times 2}$$

原始空间为  $\mathbb{R}^2$   
数据点为 3

- (a) [2] Find the first principal component direction  $v_1$ . (Hint: drawing a scatterplot may help.)



$$\vec{v}_1 = \begin{pmatrix} 1/\sqrt{2} \\ 1/\sqrt{2} \end{pmatrix} \text{ for unit norm. } \vec{v}_1$$

长度

只有在线上, 三点的分布不会 (从 projection 的角度去思考)

有畸变. 若不在 PC1, 则:  $\frac{1}{\sqrt{2}}$  点与点间距为 1,  $\frac{1}{\sqrt{2}}$  点与点间距  $\sqrt{2}$   
差别: 1

- (b) [2] How much variance is explained by the first principal component?

Variance Explained by PC: projection (dot product) 将原始数据投影至新坐标轴之上.

$$z_{(1)} = x_i^T v_1 = \text{Variance explained by PC1}$$

$$\begin{cases} 0 \cdot \frac{1}{\sqrt{2}} + 0 \cdot \frac{1}{\sqrt{2}} = 0 \\ 1 \cdot \frac{1}{\sqrt{2}} + 1 \cdot \frac{1}{\sqrt{2}} = \sqrt{2} \\ -1 + \frac{1}{\sqrt{2}} + (-1) \cdot \frac{1}{\sqrt{2}} = -\sqrt{2} \end{cases} \Rightarrow \bar{z}_{(1)} = 0 + \sqrt{2} + (-\sqrt{2}) = 0 \Rightarrow \lambda_1 = \text{Var}(z_{(1)}) = \frac{1}{3} \sum (z_i - \bar{z})^2$$

$$z_{(1)} = \begin{bmatrix} 0 \\ \sqrt{2} \\ -\sqrt{2} \end{bmatrix} \in \mathbb{R}^3$$

$$= \frac{1}{3} \sum (z_i)^2 = \frac{1}{3} (0 + 2 + 2) = \frac{4}{3}$$

- (c) [2] Find the second principal component direction  $v_2$ . How much variance is explained by the second component?

$$v_2 \perp v_1, \langle v_2, v_1 \rangle = 0 \Rightarrow v_2 = \begin{bmatrix} -1/\sqrt{2} \\ 1/\sqrt{2} \end{bmatrix} \leftarrow \text{交换后取反}$$

$$\bar{z}_{(2)} = 0 + 0 + 0 = 0 \Rightarrow \lambda_2 = \text{Var}(z_{(2)}) = \frac{1}{3} \sum (z_i - \bar{z})^2$$

$$z_{(2)} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}$$

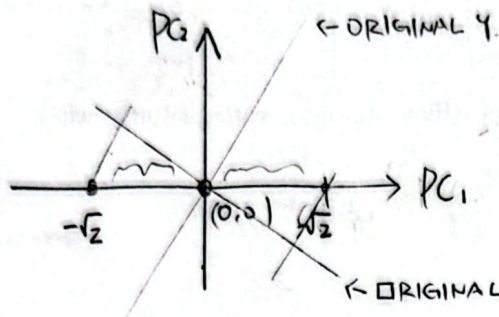
$$= 0 \Rightarrow \text{Variance Explained is 0.}$$

(b) 更直接的说明

$$K(a, b) = \langle a, b \rangle \Rightarrow \Phi(a) = a, \Phi(b) = b \Rightarrow \text{No Transformation Applied.}$$

$$= \langle \Phi(a), \Phi(b) \rangle \text{ (in context)}$$

(d) [2] Draw a biplot using the first and second principal components, and indicate the proportion of variance explained by each component.



3 points has no projection on PC2.

$$PVE_1 = \frac{PC1 \text{ Variance}}{\text{Total Variance}} = \frac{\frac{4}{3}}{\frac{4}{3} + 0} = 1$$

$$PVE_2 = \frac{0}{\frac{4}{3} + 0} = 0$$

3. [4] Suppose we apply kernel PCA to a dataset  $X = \begin{pmatrix} -x_1^T \\ \vdots \\ -x_n^T \end{pmatrix} \in \mathbb{R}^{n \times p}$  using a kernel function  $K(a, b)$ . We compute the kernel matrix  $K \in \mathbb{R}^{n \times n}$  with entries  $K_{ij} = K(x_i, x_j)$ .

(a) [2] Is the kernel matrix  $K$  always positive semidefinite? Explain why or why not.

kernel PCA:  $\Phi$ : New Feature Space Coordinate  
 $K$ : Gram Matrix.

Matrix A is PSD if for  $\forall v \neq 0, v \in \mathbb{R}^n$  we have  $v^T A v \geq 0$ .

For  $\Phi$ , let  $\Phi: \mathbb{R}^p \rightarrow \mathbb{R}^d$ , then  $K_{ij} = \Phi(x_i)^T \Phi(x_j)$ ,  $K = \Phi \Phi^T$ . So  $v^T K v = v^T \Phi \Phi^T v = \|\Phi^T v\|^2 \geq 0$ .

(b) [2] Explain why using the linear kernel  $K(a, b) = \langle a, b \rangle$  in kernel PCA is equivalent to performing standard PCA on the original dataset.

(SVD)

Linear kernel PCA = standard PCA; proof.

For  $K(a, b) = \langle a, b \rangle = a^T b \Rightarrow \Phi(x) = x$ .

(分解)  $\begin{pmatrix} v^T \\ v^T v = 1 \end{pmatrix}$

$K_{ij} = x_i^T x_j \Rightarrow K = X X^T \in \mathbb{R}^{n \times n}$ , Goal:  $(X X^T) \alpha = \lambda \alpha$ , For  $X X^T$ , SVD

SVD 适用于

任何矩阵

$S \propto X^T X$ , Goal:  $(X^T X) w = \lambda w$ .

D: 对角线上是  $\sigma_i$ , 其它为 0 的 diagonal matrix

Sample Covariance Matrix

U: 正交矩阵  $\in \mathbb{R}^{n \times n}$  For  $X = U D V^T$  by SVD,  $X^T X = (V D U^T)(U D V^T) = V D^2 V^T$

$U U^T = I_n$

$\Rightarrow$  Goal is Same

As  $X^T X = X X^T$  in 1

$$\begin{pmatrix} 1 \\ 1 \end{pmatrix} \begin{pmatrix} -1 \\ 1 \end{pmatrix} = \begin{bmatrix} f_1 & f_1^T + f_2 & f_2^T \end{bmatrix}$$

4. [6] Consider two centered feature columns  $f_1 \in \mathbb{R}^n$  and  $f_2 \in \mathbb{R}^n$  that are uncorrelated and have unit sample variance:

$$\text{指的方差} \rightarrow \begin{pmatrix} s_{f_1}^2 \\ s_{f_2}^2 \end{pmatrix} = s_{f_1}^2 = s_{f_2}^2 = 1, \quad s_{f_1 f_2} = 0.$$

We observe two datasets:

$$\text{Var}(f_1) = 1, \quad \text{Var}(f_2) = 1$$

$$X = \begin{pmatrix} | & | \\ f_1 & f_2 \\ | & | \end{pmatrix} \in \mathbb{R}^{n \times 2}, \quad Y = \begin{pmatrix} | \\ f_1 \\ | \end{pmatrix} \in \mathbb{R}^{n \times 1}.$$

$$\text{Cov}(f_1, f_2) = 0$$

(a) [2] Write down the sample covariance matrices  $S_X$ ,  $S_Y$ , and  $S_{XY}$  required for canonical correlation analysis.

NOTATION:  $S_X$ : X's sample covariance matrix

$S_{XY}$ : X and Y's cross covariance matrix

$$S_X = \begin{pmatrix} \text{Var}(f_1) & \text{Cov}(f_1, f_2) \\ \text{Cov}(f_2, f_1) & \text{Var}(f_2) \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \in \mathbb{R}^{2 \times 2}$$

$$S_Y = 1 \in \mathbb{R}$$

$$S_{XY} = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \in \mathbb{R}^{2 \times 1}$$

有关Block表达与scalar/input表达

相互间的转译以及更多, 看note Q4解析 M2

(b) [2] What is the first canonical correlation between  $X$  and  $Y$ ? (Hint: you can solve this question without computing an SVD, just use the structure of the data.)

有关 First/second PCs  
解释与见 M2 Q4 解析

Ans: CCA finds linear combination of columns  $X$  (i.e.  $f_1$  and  $f_2$ ) that are max correlated with  $y$  (i.e.  $f_1$ )

Since  $y$  coincides with the 1st column of  $X$

We can make this correlation = 1

$$\text{i.e. } U = 1 \cdot f_1 + 0 \cdot f_2$$

$$V = 1 \cdot f_1$$

Max Corr between  $X, Y$

1

如何让  $X$  与  $Y$  最像?  $\Rightarrow$  因为  $X$  有  $Y$

那丢掉  $Y$  的就可

(c) [2] Determine the corresponding canonical directions  $u_1$  for  $X$  and  $v_1$  for  $Y$ .

For  $X$ , we only ~~seek for~~ <sup>keep</sup> first column and omit second.

$$\text{So } U_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \text{ meaning: } \in \mathbb{R}^{2 \times 1}$$

For  $Y$ , we keep everything  $\Rightarrow V_1 = 1 \in \mathbb{R}$

To verify Subject to constraints of  $U_1^T S_X U_1 = V_1^T S_Y V_1 = 1$

$$\text{Var}(X u_1) = 1 \text{ As } U_1^T S_X U_1 = \begin{pmatrix} 1 & 0 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} = 1$$

$$\text{Var}(X v_1) = 1 \text{ As } V_1^T S_Y V_1 = 1 \cdot 1 = 1$$

$$U_1^T S_{XY} V_1 = \begin{pmatrix} 1 & 0 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} = 1$$

# FACTOR ANALYSIS

MODEL DEFINITION

$$X = \mu + LZ + \epsilon$$

$$X = LZ + \epsilon$$

$L \rightarrow \mathbb{R}^{p \times k}$

$\mathbb{R}^{p \times 1}$   $\mathbb{R}^{k \times 1}$   $\mathbb{R}^{p \times 1}$

$Z \sim N(0, I_k)$

$\epsilon \sim N_p(0, \Psi)$

$\Sigma$  and  $Z$  are independent

CONDITIONAL & JOINT NORMAL DISTRIBUTION

$$\text{For } \begin{pmatrix} x \\ z \end{pmatrix} \sim N_{p+k} \left( \begin{pmatrix} \mu_x \\ \mu_z \end{pmatrix}, \begin{pmatrix} \Sigma_{xx} & \Sigma_{xz} \\ \Sigma_{zx} & \Sigma_{zz} \end{pmatrix} \right)$$

Joint. We have  $X|Z \sim N(E(X|Z), \text{Var}(X|Z))$

With  $E(X|Z) = \mu_x + \Sigma_{xz} \Sigma_{zz}^{-1} (Z - \mu_z)$   $\text{Var}(X|Z) = \Sigma_{xx} - \Sigma_{xz} \Sigma_{zz}^{-1} \Sigma_{zx}$

(固定公式)

Q1. LOGICFLOW: Marginal  $\rightarrow$  Joint  $\rightarrow$  Conditional

$$E(X) = E(LZ + \epsilon) = LE(Z) + E(\epsilon) = 0 + 0 = 0$$

## Note 9

$$\text{Var}(X) = \text{Var}(LZ + \epsilon) = L \text{Var}(Z) L^T + \text{Var}(\epsilon) = LL^T + \Psi$$

Joint: For  $\text{Cov}(X, Z)$

$$E[(X - \mu_x)(Z - \mu_z)^T] = E(XZ^T) = E((LZ + \epsilon)Z^T) = E(LZZ^T + \epsilon Z^T)$$

$$= E(L(ZZ^T) + E(\epsilon\epsilon^T)) = L(I) + E(\epsilon\epsilon^T) = L(I) + \Psi = L$$

2. Show that PVE by  $j$ -th factor in FA via PCA is  $\frac{\lambda_j}{\text{tr}(\Sigma)}$  where  $\lambda_j$  is  $j$ -th eval of  $\Sigma$ .

3. Is  $PVE = \frac{L L_j^T}{\text{tr}(\Sigma)}$  rotation invariant? i.e. will you get the same value of PVE if use  $\tilde{L}$  instead? Are  $PVE_j = \frac{\sum_{i=1}^p \epsilon_{ij}^2}{\text{tr}(\Sigma)}$  rotation invariant?

4. FA Solution is not guaranteed to exist. Consider:

$$\Sigma = \begin{pmatrix} 1 & 0.9 & 0.7 \\ 0.9 & 1 & 0.4 \\ 0.7 & 0.4 & 1 \end{pmatrix} \text{ for } x = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}$$

We want to use  $r = 1$ , i.e. find  $L = \begin{pmatrix} \ell_{11} \\ \ell_{21} \\ \ell_{31} \end{pmatrix} \in \mathbb{R}^{3 \times 1}$  and  $\Psi = \text{diag}(\psi_1, \psi_2, \psi_3)$  Such

that  $\Sigma = LL^T + \Psi$ .

- Show that  $\ell_{11}\ell_{21} = 0.9, \ell_{21}\ell_{31} = 0.4, \ell_{11}\ell_{31} = 0.7$
- Find value of  $\ell_{11}$
- Show that  $\ell_{11} = \text{Cor}(x_1, x_1)$  and explain why there is no solution to  $\Sigma = LL^T + \Psi$
- Find value of  $\psi_1$
- Show that it's not possible for  $\psi_1 = \text{var}(\epsilon_1)$

5. Finding scores via conditional distribution. Consider  $y = \begin{pmatrix} x \\ z \end{pmatrix}$

- Find joint distribution for  $y$
- Find conditional distribution  $z|x$
- Argue how to use formula for  $E(z|x)$  to find scores  $z_1 \dots z_n$
- Why  $E(z|x)$  is a good estimate?

## Note 10 #

1. Assume that  $P \in \mathbb{R}^{p \times p}$  is a permutation matrix that permutes  $i$ -th and  $j$ -th elements in  $\mathbb{R}^p$ , i.e.

$$E(ZZ^T) = \begin{bmatrix} E(z_1^2) & E(z_1 z_2) \\ E(z_2 z_1) & E(z_2^2) \end{bmatrix}$$

For  $Z = \begin{bmatrix} z_1 \\ z_2 \end{bmatrix}$ , as an example.

Here, in context,  $Z \sim N(0, I)$

So  $E(z_1) = 0 = E(z_2)$

$$\text{Var}(z_1) = E(z_1^2) = 1$$

$$= E(\epsilon\epsilon^T) + E(\epsilon\epsilon^T)$$

$$= L(I) + E(\epsilon\epsilon^T)$$

$$= L \cdot I = L$$

by independence Assumption

$$E(X|Z) = \mu_x + \Sigma_{xz} \Sigma_{zz}^{-1} (Z - \mu_z)$$

$$= 0 + L I^{-1} (Z - 0) = LZ$$

$$\text{Var}(X|Z) = \Sigma_{xx} - \Sigma_{xz} \Sigma_{zz}^{-1} \Sigma_{zx} = (LL^T + \Psi) - L I^{-1} L^T$$

$$X|Z \sim N(LZ, \Psi) = LL^T + \Psi - LL^T = \Psi$$

So we conclude  $Z \sim N(0, I) \Rightarrow X|Z \sim N(LZ, \Psi)$

## Q2. 3. FA via PCA.

By spectral decomposition, covariance matrix could be rewritten as:  $\Sigma = \sum_{i=1}^p \lambda_i v_i v_i^T$  For  $(\lambda_i, v_i)$  to be the eigenvalue and eigenvector.

By selecting  $r$  number of lead eigenvectors & its corresponding eigenvalues, we have  $\ell_i = \sqrt{\lambda_i} v_i$

## Q3. 4. VARIMAX CRITERION.

Introducing orthogonal matrix  $Q$  s.t.  $QQ^T = I$

$$\text{So } Z = LQ$$

Purpose: Make as many values close to 0 as possible. Resulting few large loadings.

Row Vector: 行向量  $\in \mathbb{R}^{1 \times n}$  Column Vector: 列向量  $\in \mathbb{R}^{n \times 1}$

Q1. Assume  $P$  is a permutation matrix that permute  $i$ th and  $j$ th elements in  $\mathbb{R}^P$ .

i.e. if  $x = \begin{pmatrix} x_1 \\ \vdots \\ x_i \\ \vdots \\ x_j \\ \vdots \\ x_p \end{pmatrix}$  then  $Px = \begin{pmatrix} x_1 \\ \vdots \\ x_j \\ \vdots \\ x_i \\ \vdots \\ x_p \end{pmatrix}$

Permutation Matrix:  $n \times n$  square matrix.

used switching row when  $Px$ , switching column when  $x^T P$

All permutation matrix are orthonormal matrix, as  $PP^T = I \Rightarrow P^T = P^{-1}$ , meaning  $P$  is involutory matrix.

$P_{kk} = 1$  for  $k \neq i$  or  $j$

$P_{ij} = P_{ji} = 1$   
 $P_{ii} = P_{jj} = 0$

a. Write down  $P$  explicitly

b. Write down  $P^{-1}$  explicitly

(Hint: note that  $P^{-1}(Px) = P^{-1} \begin{pmatrix} x_1 \\ \vdots \\ x_j \\ \vdots \\ x_i \\ \vdots \\ x_p \end{pmatrix} = x = \begin{pmatrix} x_1 \\ \vdots \\ x_i \\ \vdots \\ x_j \\ \vdots \\ x_p \end{pmatrix}$ )

c. If  $z = \begin{pmatrix} z_1 \\ \vdots \\ z_p \end{pmatrix}$  and  $L = (\ell_1 \dots \ell_p)$  write  $\tilde{z} = Pz$  and  $\tilde{L} = LP^{-1}$  in terms of  $L$  and  $z$

d. Using equation  $Lz = \ell_1 z_1 + \dots + \ell_p z_p$  show that  $\tilde{L}\tilde{z} = Lz$

2. Uncorrelated  $\Rightarrow$  independent. Consider random vector  $x = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$  with joint distribution

$x = \begin{cases} (0, 1) \\ (0, -1) \\ (1, 0) \\ (-1, 0) \end{cases}$  with probability  $1/4$

o Show that  $\text{cor}(x_1, x_2) = 0$

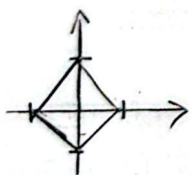
o Find marginal distributions  $f_{x_1}(x_1), f_{x_2}(x_2)$

o Show that  $x_1$  and  $x_2$  are not independent i.e.  $f_{x_1, x_2}(x_1, x_2) \neq f_{x_1}(x_1) \cdot f_{x_2}(x_2)$

3. If  $y_1 \sim (\mu_1, \sigma_1^2)$  and  $y_2 \sim (\mu_2, \sigma_2^2)$  derive that

$X(y_1 + y_2) = \frac{\sigma_1^2 X(y_1) + \sigma_2^2 X(y_2)}{(\sigma_1^2 + \sigma_2^2)^2}$

Q2. Uncorrelated  $\neq$  independent. For  $E(X) = \sum_{\text{All possible } x} X_i \cdot \text{Probability}$



By the question set-up, following the definition of  $X$ ,

We have:  $E(X_1) = 0 \cdot \frac{1}{4} + 0 \cdot \frac{1}{4} + (-1) \cdot \frac{1}{4} + 1 \cdot \frac{1}{4} = 0$

$E(X_2) = 1 \cdot \frac{1}{4} + (-1) \cdot \frac{1}{4} + \frac{1}{4} \cdot 0 + \frac{1}{4} \cdot 0 = 0$

And  $E(X_1 X_2) = 0$  For  $x_1, x_2 \in \{0, 1, -1, 0\}$   
 $= \{0, 0, 0, 0\}$

So  $\text{COV}(X_1, X_2) = E(X_1 X_2) - E(X_1)E(X_2) = 0 - 0 = 0$

By definition of independence, if  $x_1$  and  $x_2$  are independent, then  $f_{x_1, x_2}(x_1, x_2) = f_{x_1}(x_1) f_{x_2}(x_2)$ .

Here,  $P(X_1 = 0) = P(0, 1) + P(0, -1) = \frac{1}{2}$

Similarly,  $P(X_2 = 0) = \frac{1}{2}$   $P(X_1 = 1) = \frac{1}{4}$   $P(X_2 = 1) = \frac{1}{4}$   $P(X_1 = -1) = \frac{1}{4}$   $P(X_2 = -1) = \frac{1}{4}$

For  $x = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$ ,  $P(X_1 = 0, X_2 = 1) = \frac{1}{4}$ ,  $f_{x_1}(x_1 = 0) f_{x_2}(x_2 = 1) = \frac{1}{2} \cdot \frac{1}{4} = \frac{1}{8} \neq \frac{1}{4} \Rightarrow$  Not Independent

Question 1. (a).  $P = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$   $i$ th  $j$ th.

define  $P \in \mathbb{R}^{n \times n}$  for  $P_{kk} = 1$  for  $k \neq i$  or  $j$   
and  $P_{ij} = P_{ji} = 1$ , with  $P_{ii} = P_{jj} = 0$

(b). Consider  $P^{-1}$ , since  $PPx = x$

$\Rightarrow P^2 = I$

(c). For  $\tilde{z} = Pz$ ,  $L \in \mathbb{R}^{1 \times n}$

$= P \cdot \begin{bmatrix} \ell_1 \\ \vdots \\ \ell_n \end{bmatrix}$  For  $\tilde{L} = LP^{-1}$

$= \begin{bmatrix} \ell_1 \\ \vdots \\ \ell_n \end{bmatrix} (\ell_1, \ell_2, \dots, \ell_n) P^{-1} = (\ell_1, \ell_2, \dots, \ell_n)$   
 $\uparrow \quad \uparrow$   
 $i$ th  $j$ th (swapped)  $j$ th  $i$ th

(d). For  $\tilde{L}\tilde{z}$ ,

We have  $[\ell_1 \dots \ell_j \dots \ell_i \dots \ell_p] \begin{bmatrix} z_1 \\ \vdots \\ z_j \\ \vdots \\ z_i \\ \vdots \\ z_p \end{bmatrix} = \sum_{k=1}^p \ell_k z_k$

Since  $\sum_{k \neq i, j} (\ell_k + \tilde{\ell}_k) + (\tilde{\ell}_i \tilde{z}_i) + (\tilde{\ell}_j \tilde{z}_j) = Lz$

$= \sum_{k \neq i, j} (\ell_k + \tilde{\ell}_k) + (\tilde{\ell}_i \tilde{z}_i) + (\tilde{\ell}_j \tilde{z}_j)$

$= \sum_{k=1}^p \ell_k z_k = Lz$

So PCA based on correlation is sufficient for signal separation.

Q3: see page

63

bottom

## Q2 Solution.

$$PVE_j = \frac{\|l_j\|^2}{\text{tr}(\Sigma)} \leftarrow \begin{matrix} j^{\text{th}} \text{ Factor} \\ \text{Explained Variance.} \end{matrix}$$

For  $L = (l_1, \dots, l_r)$   $\text{tr}(\Sigma) = \sum_{j=1}^r \lambda_j = \text{total variance}$

Note that: in FA via PCA,  $L = (\sqrt{\lambda_1} v_1, \dots, \sqrt{\lambda_r} v_r)$ , so the  $j^{\text{th}}$  load vector is defined by  $\sqrt{\lambda_j} v_j$ , for  $\lambda_j$  to be  $j^{\text{th}}$  biggest eigenvalue.  $v_j$  is the normalized Eigenvector, i.e.  $v_j^T v_j = 1$  (unit)

$$\text{So } \|l_j\|^2 = l_j^T l_j = (\sqrt{\lambda_j} v_j)^T (\sqrt{\lambda_j} v_j) = (\sqrt{\lambda_j})^2 (v_j^T v_j) = (\sqrt{\lambda_j})^2 = \lambda_j$$

$$\Rightarrow PVE_j = \frac{\lambda_j}{\text{tr}(\Sigma)}$$

Which means  $PVE_j$  is  $\lambda_j / \text{tr}(\Sigma)$  for  $j^{\text{th}}$  eigenvalue of  $\Sigma$ .

## Q3 Solution

i.e. Never change

Q: Is PVE overall and  $PVE_j$  vs. all invariant after rotation?

A: Yes. For overall, No for  $PVE_j$ .

Proof: WTS:  $\|\tilde{L}\|_F^2 = \|L\|_F^2$ , by definition,  $\|\tilde{L}\|_F^2 = \text{tr}(\tilde{L}^T \tilde{L})$

$$\text{trace}(\tilde{L}^T \tilde{L}) = \text{trace}(L^T (Q Q^T) L) = \text{tr}(L^T L) = \text{tr}(\Sigma) = \|L\|_F^2$$

(Note: orthogonal  $Q$ ,  $Q Q^T = Q^T Q = I$ )  $\Rightarrow I$  For orthogonal  $Q$

②. WTS  $\|\tilde{l}_j\|^2 \neq \|l_j\|^2$ .

For  $\tilde{L} = (\tilde{l}_1, \dots, \tilde{l}_r) = (L q_1, \dots, L q_r)$   $\uparrow$   $r^{\text{th}}$  column of  $Q$

For  $\|\tilde{l}_j\|^2$ ,

$$= \tilde{l}_j^T \tilde{l}_j = (L q_j)^T (L q_j) = q_j^T L^T L q_j, \text{ this expression can take different value depending on } q_j$$

## Q4. Identifiability:

In factoranalysis model, if given  $\Sigma$  and find solution of  $L$  and  $\Psi$ , then we are solving a system of equations.

OFF-DIAGONAL Elements are decided by

$LL^T$ , while diagonal elements are decided

by both  $\Psi$  and  $LL^T$  for  $\Psi$  is a diagonal

$$\text{matrix. } \text{Var}(X_i) = \text{trace}(\Sigma) = \text{tr}(LL^T) + \text{tr}(\Psi) = \sum l_{ik}^2 + \Psi_i$$

## Q4

Consider when  $r=1$ ,  $r$  is dimensionality of latent vector  $Z$ . So  $r=1, p=3 \Rightarrow L = \begin{pmatrix} l_{11} \\ l_{21} \\ l_{31} \end{pmatrix} \in \mathbb{R}^{3 \times 1}$

In this case,  $LL^T = \begin{pmatrix} l_{11} & l_{21} & l_{31} \end{pmatrix} \begin{pmatrix} l_{11} \\ l_{21} \\ l_{31} \end{pmatrix} = \begin{pmatrix} l_{11}^2 & l_{11} l_{21} & l_{11} l_{31} \\ l_{21} l_{11} & l_{21}^2 & l_{21} l_{31} \\ l_{31} l_{11} & l_{31} l_{21} & l_{31}^2 \end{pmatrix} \in \mathbb{R}^{3 \times 3}$

For  $r=1$ , the correlation between, for  $\Sigma_{12}$ , which is  $l_{11} l_{21}$ . example,  $\Sigma_{11}$  and  $\Sigma_{22}$ , will be explained solely by

$$\Sigma = LL^T + \Psi \Leftrightarrow \begin{pmatrix} \Sigma_{11} & \Sigma_{12} & \Sigma_{13} \\ \Sigma_{21} & \Sigma_{22} & \Sigma_{23} \\ \Sigma_{31} & \Sigma_{32} & \Sigma_{33} \end{pmatrix} \Leftrightarrow \text{Solve non-diagonal elements First}$$

$$\text{For } \begin{pmatrix} 1 & 0.9 & 0.7 \\ 0.9 & 1 & 0.4 \\ 0.7 & 0.4 & 1 \end{pmatrix} \Rightarrow l_{11}^2 = \frac{(l_{11} l_{21})(l_{11} l_{31})}{l_{21} l_{31}} = 1.575 \Rightarrow l_{11} \approx 1.255$$

$$\text{However, } \text{Cov}(X_1, Z_1) = E(X_1 Z_1) = E[(LZ + \epsilon) Z] = l_{11} E[Z^2] = l_{11}$$

$$\text{And } \text{Cov}(X_1, Z_1) = \frac{\text{Cov}(X_1, Z_1)}{\sqrt{\text{Var}(X_1)} \sqrt{\text{Var}(Z_1)}} = \frac{l_{11}}{1 \cdot 1} = l_{11}, \text{ but } l_{11} = 1.255$$

## Q5 SCORING.

The joint distribution is:  $\begin{pmatrix} X \\ Z \end{pmatrix} \sim \mathcal{N}_{\text{mult}} \left( \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} LL^T + \Psi & L \\ L^T & I \end{pmatrix} \right)$

For  $Z|X$ , we have

$$\begin{cases} \mu_{Z|X} = \mu_Z + \Sigma_{ZX} \Sigma_{XX}^{-1} (X - \mu_X) \\ \Sigma_{Z|X} = \Sigma_{ZZ} - \Sigma_{ZX} \Sigma_{XX}^{-1} \Sigma_{XZ} \end{cases}$$

The score:

$$E(Z|X) = 0 + L^T (LL^T + \Psi)^{-1} \cdot (X - 0) = L^T (LL^T + \Psi)^{-1} X$$

$$V(Z|X) = I - L^T (LL^T + \Psi)^{-1} L$$

We seek for  $\hat{z} = E(Z|X)$  s.t.  $E\|Z - \hat{z}\|^2$  is minimized (具体数学证明待补充)

(Regression Scoring via Bayesian Stats)

$l_{11} > 1 \approx 1.255$ , and  $P$   
We conclude a contradiction and thus  $l_{11} > 1$  is not possible.

$$\Sigma_{11} = l_{11}^2 + \Psi_{11}$$

$$1 = 1.575 + \Psi_{11}$$

$$\Rightarrow \Psi_{11} < 0$$

but  $\text{Var}(\epsilon) \geq 0$

As variance

non-negative

Conclude a