

Mahalanobis transformation

Univariate $Y = \frac{X - E[X]}{\sqrt{\text{Var}[X]}} = \frac{X}{\sqrt{\text{Var}[X]}} - \frac{E[X]}{\sqrt{\text{Var}[X]}} \Rightarrow E[Y] = 0, \text{Var}[Y] = 1$

X multivariate. $E[\underline{X}] = \underline{\mu}$, $\text{Cov}[\underline{X}] = \underline{\Sigma}$

Can we find $\underline{Y} = \underline{A}\underline{X} + \underline{b}$ such that $E[\underline{Y}] = 0$
and $\text{Cov}[\underline{Y}] = \underline{I}$

Must have: $\underline{A}E[\underline{X}] + \underline{b} = \underline{A}\underline{\mu} + \underline{b} = 0$

2. $\text{Cov}[\underline{Y}] = \underline{A}\underline{\Sigma}\underline{A}^T = \underline{I}$. (assume positive definite)

Start with 2. Define $\underline{\Sigma}^{\frac{1}{2}} = \underline{P}\underline{\Lambda}^{\frac{1}{2}}\underline{P}^T$ and $\underline{\Sigma}^{-\frac{1}{2}} = \underline{P}\underline{\Lambda}^{-\frac{1}{2}}\underline{P}^T$

where $\underline{\Lambda}^{\frac{1}{2}} = \begin{bmatrix} \sqrt{\lambda_1} & & 0 \\ & \ddots & \\ 0 & & \sqrt{\lambda_p} \end{bmatrix}$ and $\underline{\Lambda}^{-\frac{1}{2}} = \begin{bmatrix} 1/\sqrt{\lambda_1} & & 0 \\ & \ddots & \\ 0 & & 1/\sqrt{\lambda_p} \end{bmatrix}$

Then $\underline{\Sigma}^{\frac{1}{2}}\underline{\Sigma}^{-\frac{1}{2}} = \underline{P}\underline{\Lambda}^{\frac{1}{2}}\underbrace{\underline{P}^T\underline{P}}_{\underline{I}}\underline{\Lambda}^{-\frac{1}{2}}\underline{P}^T = \underline{P}\underline{P}^T = \underline{I}$

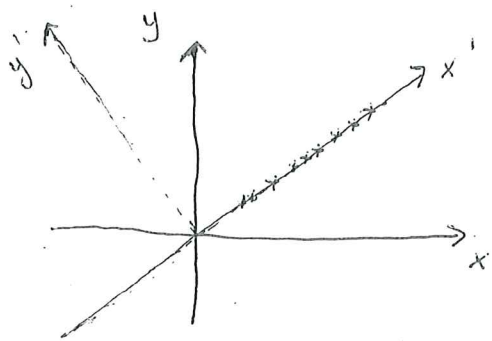
With $\underline{A} = \underline{\Sigma}^{-\frac{1}{2}}$ we get ~~for~~ $\text{Cov}[\underline{Y}] = \underline{A}\underline{\Sigma}\underline{A}^T$

$$= \underline{P}\underline{\Lambda}^{-\frac{1}{2}}\underbrace{\underline{P}^T\underline{P}}_{\underline{I}}\underbrace{\underline{\Lambda}^{\frac{1}{2}}\underline{\Lambda}^{-\frac{1}{2}}}_{\underline{I}}\underline{P}^T = \underline{P}\underline{P}^T = \underline{I}$$

and $E[\underline{Y}] = \underline{\Sigma}^{-\frac{1}{2}}\underline{\mu} + \underline{b} = 0 \Rightarrow \underline{b} = -\underline{\Sigma}^{-\frac{1}{2}}\underline{\mu}$

Hence the transformation is: $\underline{Y} = \underline{\Sigma}^{-\frac{1}{2}}(\underline{X} - \underline{\mu})$

Principal components



Principal component analysis is a method for re-expressing multivariate data. It allows the researcher to reorient the data such that the first few dimensions account for as much of the variation in the original variables as possible. Most often it is used for dimension reduction and data interpretation.

Let $\underline{x}^T = (x_1, x_2, \dots, x_p)$ and let $\underline{\Sigma} = \text{cov}(\underline{x})$

A linear combination of $\underline{x}$ is given by $y = \underline{a}^T \underline{x}$

$$\text{and } \text{Var}[y] = E[\underline{a}^T (\underline{x} - \underline{\mu})(\underline{x} - \underline{\mu})^T \underline{a}] = \underline{a}^T \underline{\Sigma} \underline{a}$$

Principal components are uncorrelated linear combinations of $\underline{x}$ constructed such that

1. The first one maximizes $\text{Var}[\underline{a}_1^T \underline{x}]$
 $\underline{a}_1^T \underline{a}_1 = 1$

2. The second one maximizes $\text{Var}[\underline{a}_2^T \underline{x}]$ under the
 $\underline{a}_2^T \underline{a}_2 = 1$

$$\text{constraint } \text{Cov}(\underline{a}_1^T \underline{x}, \underline{a}_2^T \underline{x}) = \underline{a}_1^T \underline{\Sigma} \underline{a}_2 = 0$$

3. The third one maximizes $\text{Var}[\underline{a}_3^T \underline{x}]$ under the
 $\underline{a}_3^T \underline{a}_3 = 1$

$$\text{constraint } \underline{a}_i^T \underline{\Sigma} \underline{a}_j = 0, \quad i < j$$

and so on.

Rp1. Let $\underline{X}$ have ^(positive definite) covariance matrix $\underline{\Sigma}$ with eigenvalues

$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p \geq 0$ and eigenvectors $\underline{e}_1, \underline{e}_2, \dots, \underline{e}_p$. The

i -th principal component is given by $Y_i = \underline{e}_i^T \underline{X}$, $i = 1, 2, \dots, p$.

$\text{Var}[Y_i] = \lambda_i$ and $\text{Cov}[Y_i, Y_j] = 0$, $i \neq j$.

Proof $\underline{Y} = \underline{P}^T \underline{X} \Rightarrow \text{Cov}[\underline{Y}] = \underline{P}^T \underline{\Sigma} \underline{P} = \underbrace{\underline{P}^T \underline{P}}_{\underline{I}} \underbrace{\underline{P}^T \underline{P}}_{\underline{I}} = \underline{I}$

$$= \begin{bmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & \dots & \lambda_p \end{bmatrix}$$

Further: $\max_{\underline{a}^T \underline{a} = 1} \text{Var}[\underline{a}^T \underline{X}] = \max_{\underline{a}^T \underline{a} = 1} \frac{\text{Var}[\underline{a}^T \underline{X}]}{\underline{a}^T \underline{a} = 1} = \max_{\underline{a}^T \underline{a} = 1} \frac{\underline{a}^T \underline{\Sigma} \underline{a}}{\underline{a}^T \underline{a}}$

$$= \max_{\underline{a}^T \underline{a} = 1} \frac{\underline{a}^T \underline{P} \underline{P}^T \underline{a}}{\underline{a}^T \underline{P} \underline{P}^T \underline{a}} \stackrel{\text{we want only } \underline{a} \neq 0}{=} \max_{(\underline{a}^T \underline{a} = 1)} \frac{\underline{y}^T \underline{I} \underline{y}}{\underline{y}^T \underline{y}} = \max_{(\underline{a}^T \underline{a} = 1)} \frac{\sum_{i=1}^p \lambda_i y_i^2}{\sum_{i=1}^p y_i^2} \leq \lambda_1$$

But $\underline{a} = \underline{e}_1$ gives $\frac{\underline{e}_1^T \underline{\Sigma} \underline{e}_1}{\underline{e}_1^T \underline{e}_1} = \frac{\underline{e}_1^T \underline{\Sigma} \underline{e}_1}{\underline{e}_1^T \underline{e}_1} = \lambda_1 \Rightarrow \max_{\underline{a}^T \underline{a} = 1} \text{Var}[\underline{a}^T \underline{X}] = \lambda_1$

$\Rightarrow Y_1 = \underline{e}_1^T \underline{X}$

For $Y_2 = \underline{a}_2^T \underline{X}$ we shall have $\left\{ \begin{array}{l} \underline{e}_1^T \underline{\Sigma} \underline{a}_2 = 0 \\ \underline{a}_2^T \underline{\Sigma} \underline{e}_1 = 0 \end{array} \right. \Rightarrow \underline{a}_2^T \underline{e}_1 = 0$

and $\underline{y}^T = [0, y_2, \dots, y_p]$

This implies $\max_{\underline{a}_2^T \underline{a}_2 = 1} \frac{\sum_{i=2}^p \lambda_i y_i^2}{\sum_{i=2}^p y_i^2} \leq \lambda_2$

But $\underline{a}_2 = \underline{e}_2$ gives $\frac{\underline{a}_2^T \underline{\Sigma} \underline{a}_2}{\underline{a}_2^T \underline{a}_2} = \frac{\underline{e}_2^T \underline{\Sigma} \underline{e}_2}{\underline{e}_2^T \underline{e}_2} = \lambda_2 \Rightarrow \max_{\underline{a}_2^T \underline{a}_2} \text{Var}[\underline{a}_2^T \underline{X}] = \lambda_2$
 $\underline{a}_2^T \underline{\Sigma} \underline{e}_1 = 0$

and $Y_2 = \underline{e}_2^T \underline{X}$

and so on.

RP2 Let $\underline{X}$ have covariance matrix $\text{Cov}[\underline{X}] = \underline{\Sigma}$ (Σ pos definite)

and let $Y_i = \underline{e}_i^T \underline{X}$, $i = 1, 2, \dots, p$ be the principal components.

$$\text{Then } \text{tr}(\underline{\Sigma}) = \sum_{i=1}^p \sigma_{ii} = \sum_{i=1}^p \lambda_i = \sum_{i=1}^p \text{Var}(Y_i)$$

Proof. $\text{tr}(\underline{\Sigma}) = \text{tr}(\underline{P} \underline{\Lambda} \underline{P}^T) = \text{tr}(\underline{P}^T \underline{P} \underline{\Lambda}) = \text{tr}(\underline{\Lambda})$