

# MÉTODOS ESTATÍSTICOS PARA METEOROLOGIA

**Ralph dos Santos Silva**

<https://www.ralphsilva.org/meteorologia.html>

Instituto de Matemática  
Centro de Ciências Matemáticas e da Natureza  
Universidade Federal do Rio de Janeiro

## Referências

- Wilks, D. S. (2011). *Statistical Methods in the Atmospheric Sciences*, 3<sup>th</sup> Edition. Academic Press.
- Johnson e Wichern (2007). *Applied Multivariate Statistical Analysis*, 6<sup>th</sup> Edition. Prentice-Hall.

## Organização de dados

- A matéria prima a ser trabalhada aqui é um conjunto de dados multivariados, isto é, várias variáveis são observadas sobre diversos indivíduos ou objetos. Nosso objetivo será apresentar uma forma conveniente de organizar estes dados e de representá-los graficamente.
- Suponha que estejamos diante de um problema em que  $p$  variáveis foram observadas para uma amostra de  $n$  elementos. Assim, a observação para o  $i$ -ésimo elemento da amostra será um vetor  $p$ -variado denotador por  $\mathbf{x}_{.i}$  ou  $\underline{x}_{.i}$  tal que

$$\mathbf{x}_{.i}^{\top} = (x_{1i}, x_{2i}, \dots, x_{pi}), \quad i = 1, 2, \dots, n,$$

em que  $x_{ji}$  representa a observação da  $j$ -ésima variável do  $i$ -ésimo elemento da amostra,  $j = 1, 2, \dots, p$ .

- A coleção de dados observados pode ser representada por meio de uma matriz  $\mathbf{X}$  de dimensão  $p \times n$  como segue

$$\mathbf{X} = \begin{bmatrix} x_{11} & x_{12} & \dots & x_{1n} \\ x_{21} & x_{22} & \dots & x_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ x_{p1} & x_{p2} & \dots & x_{pn} \end{bmatrix}.$$

- Assim, as linhas da matriz  $\mathbf{X}$  representam as  $p$  variáveis medidas e, as colunas, as  $n$  unidades amostrais.
- Podemos representar a matriz  $\mathbf{X}$  através de suas linhas

$$\mathbf{X} = \begin{bmatrix} \mathbf{x}_{1.}^\top \\ \mathbf{x}_{2.}^\top \\ \vdots \\ \mathbf{x}_{p.}^\top \end{bmatrix},$$

sendo  $\mathbf{x}_{j.}^\top = (x_{j1}, x_{j2}, \dots, x_{jn})$  as  $n$  observações da  $j$ -ésima variável,  $j = 1, 2, \dots, p$ .

- Também podemos representar a matriz  $\mathbf{X}$  através de suas colunas. Adotando aqui a notação  $\mathbf{x}_{.i}$ ,  $i = 1, 2, \dots, n$ , para designar a  $i$ -ésima coluna de  $\mathbf{X}$ , temos

$$\mathbf{X} = [\mathbf{x}_{.1} \quad \mathbf{x}_{.2} \quad \dots \quad \mathbf{x}_{.n}],$$

em que cada  $\mathbf{x}_{.i}$  é um vetor  $p \times 1$ ,  $i = 1, 2, \dots, n$ .

## Exemplo 1:

- Uma seleção de 4 notas fiscais de uma livraria universitária foi obtida de modo a investigar a natureza das vendas. Cada nota forneceu o número de livros vendidos e o valor total da venda (em dólares). Obteve-se a seguinte matriz de dados, na qual a primeira linha indica o número de livros vendidos e, a segunda, o valor da venda.

$$\mathbf{X} = \begin{bmatrix} 4 & 5 & 4 & 3 \\ 42 & 52 & 48 & 58 \end{bmatrix}.$$

- A representação dos dados desta forma permite o cálculo de quantidades numéricas de interesse de forma eficiente e fácil.

## Estatísticas descritivas

1. A **média amostral** para a  $j$ -ésima variável observada pode ser representada como

$$\bar{x}_{j.} = \frac{1}{n} \sum_{i=1}^n x_{ji}, \quad j = 1, 2, \dots, p.$$

Assim, podemos definir o vetor de médias amostral  $\bar{\mathbf{x}}$  como

$$\bar{\mathbf{x}} = \begin{bmatrix} \bar{x}_{1.} \\ \bar{x}_{2.} \\ \vdots \\ \bar{x}_{p.} \end{bmatrix}.$$

Algebricamente,

$$\bar{\mathbf{x}} = \frac{1}{n} \mathbf{X} \mathbf{1},$$

sendo  $\mathbf{1}$  um vetor  $n \times 1$  com todos os elementos iguais a 1.

2. Uma medida de dispersão para a  $j$ -ésima variável é dada pela **variância amostral**

$$s_{jj} = \frac{1}{n-1} \sum_{i=1}^n (x_{ji} - \bar{x}_{j.})^2, \quad j = 1, 2, \dots, p.$$

3. A **covariância amostral** entre a  $j$ -ésima e a  $k$ -ésima variáveis é dada por

$$s_{jk} = \frac{1}{n-1} \sum_{i=1}^n (x_{ji} - \bar{x}_{j.})(x_{ki} - \bar{x}_{k.}), \quad j, k = 1, 2, \dots, p \text{ e } j \neq k.$$

- Podemos então representar de forma organizada as informações sobre variabilidade através da **matriz de (variâncias e) covariâncias amostral** dada por

$$\mathbf{S} = \begin{bmatrix} s_{11} & s_{12} & \dots & s_{1p} \\ s_{21} & s_{22} & \dots & s_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ s_{p1} & s_{p2} & \dots & s_{pp} \end{bmatrix}.$$

Observe que a matriz de covariâncias é uma matriz simétrica:  $\mathbf{S} = \mathbf{S}^T$ .  
(A matriz de covariância  $\mathbf{S}$  também é positiva definida.)

- Algebricamente, podemos escrever a matriz de dados menos o vetor de médias na forma

$$\mathbf{X} - \bar{\mathbf{x}}\mathbf{1}^\top,$$

e a matriz de covariâncias na forma

$$\mathbf{S} = \frac{1}{n-1}(\mathbf{X} - \bar{\mathbf{x}}\mathbf{1}^\top)(\mathbf{X} - \bar{\mathbf{x}}\mathbf{1}^\top)^\top.$$

4. Podemos também definir a **matriz de correlações amostral**  $\mathbf{R}$ , com elementos  $r_{jk} = \frac{S_{jk}}{\sqrt{S_{jj}S_{kk}}}$ ,

$$\mathbf{R} = \begin{bmatrix} r_{11} & r_{12} & \dots & r_{1p} \\ r_{21} & r_{22} & \dots & r_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ r_{p1} & r_{p2} & \dots & r_{pp} \end{bmatrix}.$$

Observe que a matriz de correlações é uma matriz simétrica:  $\mathbf{R} = \mathbf{R}^\top$ .  
(A matriz de correlação  $\mathbf{R}$  também é positiva definida.)

- Algebricamente, a matriz de correlações pode ser escrita na forma

$$\mathbf{S} = \frac{1}{n-1} \left[ \Delta^{-1/2} (\mathbf{X} - \bar{\mathbf{x}} \mathbf{1}^\top) \right] \left[ \Delta^{-1/2} (\mathbf{X} - \bar{\mathbf{x}} \mathbf{1}^\top) \right]^\top$$

com

$$\Delta^{1/2} = \text{diag}(s_{11}^{1/2}, s_{22}^{1/2}, \dots, s_{pp}^{1/2})$$

uma matriz diagonal  $p \times p$ , e  $\Delta^{-1/2} = \left[ \Delta^{1/2} \right]^{-1}$ .

- Observe que podemos relacionar algebricamente as matrizes  $\mathbf{S}$  e  $\mathbf{R}$  tal que

$$\mathbf{S} = \Delta^{1/2} \mathbf{R} \Delta^{1/2} \quad \text{e} \quad \mathbf{R} = \Delta^{-1/2} \mathbf{S} \Delta^{-1/2}.$$

- Em muitas aplicações as somas dos desvios quadrados da média e dos produtos cruzados de tais desvios são utilizadas. Adotaremos aqui a notação:

$$W_{jj} = \sum_{i=1}^n (x_{ji} - \bar{x}_{j.})^2, \quad j = 1, 2, \dots, p, \quad \text{e}$$

$$W_{jk} = \sum_{i=1}^n (x_{ji} - \bar{x}_{j.})(x_{ki} - \bar{x}_{k.}), \quad j, k = 1, 2, \dots, p \quad \text{e} \quad j \neq k.$$

Assim, definimos a matriz ***W*** de somas dos desvios quadrados da média e produtos cruzados dos desvios da média por

$$\mathbf{W} = \begin{bmatrix} W_{11} & W_{12} & \dots & W_{1p} \\ W_{21} & W_{22} & \dots & W_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ W_{p1} & W_{p2} & \dots & W_{pp} \end{bmatrix} = (\mathbf{X} - \bar{\mathbf{x}}\mathbf{1}^\top)(\mathbf{X} - \bar{\mathbf{x}}\mathbf{1}^\top)^\top.$$

## Exemplo 1: (continuação)

Para os dados do Exemplo 1 e utilizando o programa R, faça os cálculos:

- do vetor de médias; e das matrizes **S** e **R**. (Ver exemplo\_01.r)

```
p      <- 2                # numero de variaveis
n      <- 4                # tamanho da amostra
X      <- matrix(0,p,n)    # definindo X
X[1, ] <- c( 4,  5,  4,  3) # valores de X_{1i}
X[2, ] <- c(42, 52, 48, 58) # valores de X_{2i}
ones   <- matrix(1,n,1)   # vetor de uns
xbar   <- (X%*%ones) / n   # vetor de medias
sdX    <- apply(X,1,"sd")  # d. padrao das variaveis
D.half <- diag(sdX,p,p)    # matriz Delta^{1/2}
X.xbar <- X-xbar%*%t(ones) # matriz (X - xbarra*ones')
W      <- X.xbar%*%t(X.xbar) # somas dos desvios ao
                                     # quadrados e cruzados
S      <- W/(n-1)          # matriz de covariancias
ID.half <- solve(D.half)   # matriz Delta^{-1/2}
R      <- ID.half%*%S%*%ID.half # matriz de correlacoes
```

```

> xbar # vetor de medias
      [,1]
[1,]     4
[2,]    50
> S     # matriz de covariancias
           [,1]      [,2]
[1,]  0.6666667 -2.00000
[2,] -2.0000000 45.33333
> R     # matriz de correlacoes
           [,1]      [,2]
[1,]  1.0000000 -0.3638034
[2,] -0.3638034  1.0000000

```

**Usando as funções do R, podemos calcular:**

```

apply(X,1,"mean") # vetor (linha) de medias
var(t(X))         # matriz de covariancias
cor(t(X))         # matriz de correlacoes
# IMPORTANTE: note a transposicao da matriz X.

```

## Representação gráfica

- Podemos utilizar gráficos de dispersão para variáveis duas a duas.  
(Ver exemplo\_02.r)

```
# Note que o R entende variaveis por colunas e
# as amostras por linha. Entao e' necessario
# transpor a matriz de dados
p      <- 3                                # numero de variaveis
n      <- 10000                            # tamanho da amostra
X      <- matrix(rnorm(p*n),p,n) # definindo X
#pdf(file="dispersao_pairs.pdf")
#par(mfrow=c(1,1),lwd=2.0,cex.lab=1.5,cex.axis=1.5,
#    lab=c(10,5,5),mar=c(0,1,0,2.5),xpd=T,cex.main=2.0)
pairs(t(X),pch=15)
#dev.off()

apply(X,1,"mean") # vetor (linhas) de medias
var(t(X))         # matriz de covariancias
cor(t(X))         # matriz de correlacoes
```

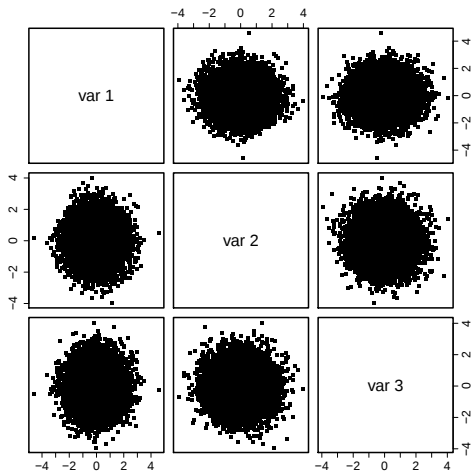


Figura 1: Exemplo de gráfico de dispersão de variáveis duas a duas.

## Distâncias

- A maior parte das técnicas multivariadas baseia-se no simples conceito de distância.
- Estamos habituados a distância usual chamada **distância euclidiana**, tal que se  $P(x_1, x_2, \dots, x_p)$  e  $Q(\mu_1, \mu_2, \dots, \mu_p)$  são dois pontos em  $\mathbb{R}^p$ , a distância entre  $P$  e  $Q$  é dada por

$$d_e(P, Q) = \sqrt{(x_1 - \mu_1)^2 + (x_2 - \mu_2)^2 + \dots + (x_p - \mu_p)^2}.$$

- Porém, a distância euclidiana pode não ser adequada em muitos problemas, dependendo da natureza das variáveis envolvidas.
- Isto ocorre devido ao fato de que na distância euclidiana cada coordenada contribui igualmente para o cálculo da mesma.
- Quando as coordenadas representam medições que são sujeitas à flutuações aleatórias de magnitudes diferentes, é frequentemente desejável ponderar coordenadas sujeitas à maior variabilidade com um peso menor do que àquelas sujeitas a uma menor variabilidade.

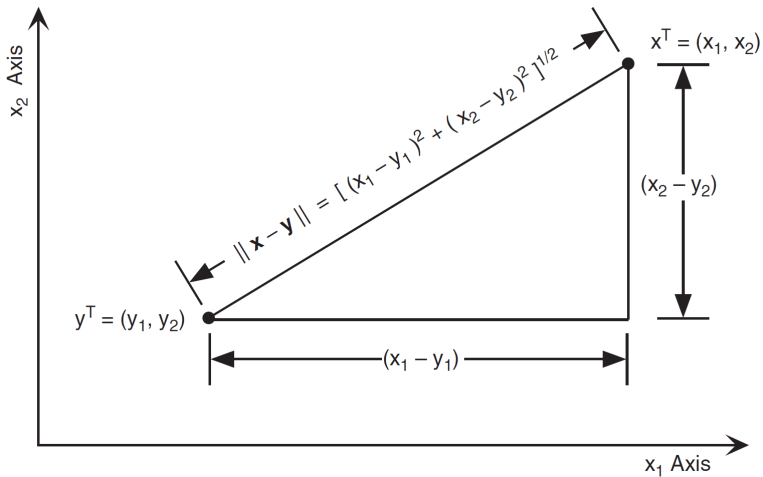


Figura 2: Exemplo da distância euclidiana entre os pontos  $\mathbf{x}$  e  $\mathbf{y}$ .

- Deseja-se uma nova medida de distância que leve em conta as diferenças em variabilidade entre as diversas variáveis incluídas na análise e a presença de correlação entre os pares de variáveis.
- Suponha primeiro um conjunto de  $p$  variáveis não correlacionadas, com variâncias distintas. Assim, de forma a equilibrar a contribuição das diversas variáveis ao cálculo da distância, podemos ponderá-las de forma inversamente proporcional aos seus desvios padrão ( $\sqrt{s_{jj}}$ ) e calculando a distância euclideana

$$d_e(P, Q) = \sqrt{(\mathbf{x} - \boldsymbol{\mu})^\top \mathbf{D}^{-1} (\mathbf{x} - \boldsymbol{\mu})},$$

sendo  $\mathbf{D}$  uma matriz diagonal com elementos  $s_{11}, s_{22}, \dots, s_{pp}$ ,  $\mathbf{x} = (x_1, x_2, \dots, x_p)^\top$  e  $\boldsymbol{\mu} = (\mu_1, \mu_2, \dots, \mu_p)^\top$ .

- Uma medida de distância que leva em conta as covariâncias entre as variáveis é dada por

$$d_e(P, Q) = \sqrt{(\mathbf{x} - \boldsymbol{\mu})^\top \mathbf{S}^{-1} (\mathbf{x} - \boldsymbol{\mu})}, \quad \text{(distância de Mahalanobis),}$$

sendo  $\mathbf{S}$  a matriz de covariâncias.

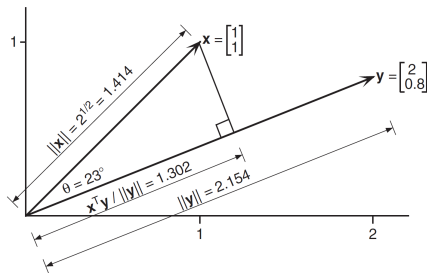


Figura 3: Exemplo do ângulo entre os vetores  $\mathbf{x}$  e  $\mathbf{y}$ .

- Tamanho de um vetor  $\mathbf{x}$ :  $\|\mathbf{x}\| = \left[ \sum_{i=1}^p x_i^2 \right]^{1/2}$  (norma euclidiana).
- Projeção de  $\mathbf{x}$  em  $\mathbf{y}$ :  $L_{\mathbf{x},\mathbf{y}} = \frac{\mathbf{x}^\top \mathbf{y}}{\|\mathbf{y}\|}$ .
- O ângulo  $\theta$  entre dois vetores:  $\theta = \cos^{-1} \left[ \frac{\mathbf{x}^\top \mathbf{y}}{\|\mathbf{x}\| \|\mathbf{y}\|} \right]$ .
- Se  $\mathbf{x}^\top \mathbf{y} = 0$ , então  $\theta = 90^\circ$ , i.é., vetores perpendiculares (ortogonal).

## Distribuição normal (univariada)

Seja  $X$  uma variável aleatória contínua com distribuição normal com média  $\mu$  e variância  $\sigma^2$ . Então, a função de densidade de probabilidade é dada por

$$f(x|\mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp \left\{ -\frac{1}{2\sigma^2} (x - \mu)^2 \right\},$$

para  $x \in \mathbb{R}$ ,  $\mu \in \mathbb{R}$  e  $\sigma^2 \in \mathbb{R}^+$ .

Temos que

$$\Pr(|X - \mu| \leq \sigma) \simeq 0,683;$$

$$\Pr(|X - \mu| \leq 2\sigma) \simeq 0,954; \quad \text{e}$$

$$\Pr(|X - \mu| \leq 3\sigma) \simeq 0,997.$$

Sabemos também que  $E(X) = \mu$  e  $\text{Var}(X) = \sigma^2$ .

## Distâncias

Sabemos que a distância euclidiana entre dois pontos,  $P(\mathbf{x})$  e  $Q(\boldsymbol{\mu})$ , no  $\mathbb{R}^p$  é dado por

$$\begin{aligned}d_e(P, Q) &= \sqrt{(x_1 - \mu_1)^2 + (x_2 - \mu_2)^2 + \cdots + (x_p - \mu_p)^2} \\&= \sqrt{(\mathbf{x} - \boldsymbol{\mu})^\top (\mathbf{x} - \boldsymbol{\mu})}.\end{aligned}$$

Se  $d = 1$ , então  $d_e(x, \mu) = \sqrt{(x - \mu)^2}$ .

Uma generalização da distância euclidiana é dada por

$$d_e(P, Q) = \sqrt{(\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Psi} (\mathbf{x} - \boldsymbol{\mu})}.$$

Neste caso, temos pesos diferentes para cada valor do vetor e também levamos em conta as possíveis interações (correlações).

No núcleo da função de densidade da normal univariada, temos

$$\left(\frac{x - \mu}{\sigma}\right)^2 = \frac{(x - \mu)^2}{\sigma^2} = (x - \mu)(\sigma^{-2})(x - \mu).$$

Isto lembra o quadrado da distância euclidiana modificada para  $d = 1$ .

Assim, o **núcleo de uma normal multivariada** é dado pelo quadrado da distância euclidiana modificada entre os pontos  $\mathbf{x}$  e  $\boldsymbol{\mu}$  no  $\mathbb{R}^p$ , isto é, por

$$(\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu}),$$

sendo a matriz de pesos dada por  $\boldsymbol{\Sigma}^{-1}$ , e  $\boldsymbol{\Sigma}$  corresponde a matriz de covariâncias (simétrica e positiva definida).

Note que quanto maior a variância de determinada variável (componente do vetor aleatório  $\mathbf{X}$ ), menor será seu peso no cálculo da distância em questão.

Para determinar a função de densidade de probabilidade da normal multivariada precisamos somente da constante de normalização. Prova-se que a constante é dada por

$$(2\pi)^{-p/2} |\boldsymbol{\Sigma}|^{-1/2},$$

com  $|\mathbf{A}| = \det(\mathbf{A})$  por definição.

## Função de densidade de probabilidade da normal multivariada

**Definição:** A função de densidade de probabilidade da normal multivariada - com vetor de médias  $\boldsymbol{\mu}$  e matriz de covariâncias  $\boldsymbol{\Sigma}$  - é dada por

$$\begin{aligned}f(\mathbf{x}|\boldsymbol{\mu}, \boldsymbol{\Sigma}) &= \frac{1}{(2\pi)^{p/2}|\boldsymbol{\Sigma}|^{1/2}} \exp\left\{-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu})\right\} \\&= (2\pi)^{-p/2}|\boldsymbol{\Sigma}|^{-1/2} \exp\left\{-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu})\right\},\end{aligned}$$

com  $\mathbf{x} \in \mathbb{R}^p$ ,  $\boldsymbol{\mu} \in \mathbb{R}^p$  e  $\boldsymbol{\Sigma}$  uma matriz  $p \times p$  simétrica e positiva definida.

Prova-se que  $E(\mathbf{X}) = \boldsymbol{\mu}$  e  $\text{Var}(\mathbf{X}) = \boldsymbol{\Sigma}$  (matriz simétrica e positiva definida).

**Notação:**  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ .

## Exemplo: Normal bivariada

$$\boldsymbol{\mu} = \begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix} = \begin{bmatrix} E(X_1) \\ E(X_2) \end{bmatrix}$$

e

$$\boldsymbol{\Sigma} = \begin{bmatrix} \sigma_{11} & \sigma_{12} \\ \sigma_{21} & \sigma_{22} \end{bmatrix} = \begin{bmatrix} \text{Var}(X_1) & \text{Cov}(X_1, X_2) \\ \text{Cov}(X_2, X_1) & \text{Var}(X_2) \end{bmatrix},$$

sendo  $\text{Cov}(X_1, X_2) = \text{Cov}(X_2, X_1) \Rightarrow \sigma_{12} = \sigma_{21}$ . Logo,

$$\begin{aligned} \boldsymbol{\Sigma}^{-1} &= \frac{1}{\sigma_{11}\sigma_{22} - \sigma_{12}^2} \begin{bmatrix} \sigma_{22} & -\sigma_{12} \\ -\sigma_{12} & \sigma_{11} \end{bmatrix} \\ &= \frac{1}{\sigma_{11}\sigma_{22}(1 - \rho_{12}^2)} \begin{bmatrix} \sigma_{22} & -\rho_{12}\sqrt{\sigma_{11}\sigma_{22}} \\ -\rho_{12}\sqrt{\sigma_{11}\sigma_{22}} & \sigma_{11} \end{bmatrix}, \end{aligned}$$

pois  $\rho_{12} = \frac{\sigma_{12}}{\sqrt{\sigma_{11}\sigma_{22}}}$  (correlação entre  $X_1$  e  $X_2$ ).

Segue-se que

$$\begin{aligned}
 & (\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu}) \\
 = & \begin{bmatrix} x_1 - \mu_1 & x_2 - \mu_2 \end{bmatrix} \frac{1}{\sigma_{11}\sigma_{22}(1 - \rho_{12}^2)} \begin{bmatrix} \sigma_{22} & -\rho_{12}\sqrt{\sigma_{11}\sigma_{22}} \\ -\rho_{12}\sqrt{\sigma_{11}\sigma_{22}} & \sigma_{11} \end{bmatrix} \begin{bmatrix} x_1 - \mu_1 \\ x_2 - \mu_2 \end{bmatrix} \\
 = & \frac{\sigma_{22}(x_1 - \mu_1)^2 + \sigma_{11}(x_2 - \mu_2)^2 - 2\rho_{12}\sqrt{\sigma_{11}\sigma_{22}}(x_1 - \mu_1)(x_2 - \mu_2)}{\sigma_{11}\sigma_{22}(1 - \rho_{12}^2)} \\
 = & \frac{1}{1 - \rho_{12}^2} \left[ \left( \frac{x_1 - \mu_1}{\sqrt{\sigma_{11}}} \right)^2 + \left( \frac{x_2 - \mu_2}{\sqrt{\sigma_{22}}} \right)^2 - 2\rho_{12} \left( \frac{x_1 - \mu_1}{\sqrt{\sigma_{11}}} \right) \left( \frac{x_2 - \mu_2}{\sqrt{\sigma_{22}}} \right) \right].
 \end{aligned}$$

Note que se as variáveis  $X_1$  e  $X_2$  são não correlacionadas, tal que  $\rho_{12} = 0$ , então a densidade conjunta pode ser reescrita como o produto de duas densidades normais. Neste caso,  $f(x_1, x_2 | \boldsymbol{\mu}, \boldsymbol{\Sigma}) = f(x_1 | \mu_1, \sigma_{11})f(x_2 | \mu_2, \sigma_{22})$ , e, portanto,  $X_1$  e  $X_2$  são variáveis aleatórias independentes.

```
library(mvtnorm)
nx <- 100; ny <- 100
x <- seq(-4,4,length=nx); y <- seq(-3,3,length=ny);
z <- matrix(NA,nx,ny)
mu <- c(0,0); S <- matrix(c(1, .9, .9, 1),2,2)
for(i in 1:nx)
  for(j in 1:ny)
    z[i,j] <- dmvtorm(c(x[i],y[j]), mean=mu,sigma=S)
# Grafico 1
x11(); persp(x,y,z, phi = 45, theta = 30,col=3)
# Grafico 2
x11(); contour(x,y,z,col=2,lwd=2)
#-----
w <- rmvtorm(100000,mean=mu,sigma=S)
#Grafico 3
x11(); par(mfrow=c(2,2))
plot(w,pch=".",main="",xlab="x",ylab="y")
hist(w[,2],prob=T,nclass=50,main="",xlab="y")
lines(y,dnorm(y),lwd=2,col=2)
hist(w[,1],prob=T,nclass=50,main="",xlab="x")
lines(x,dnorm(x),lwd=2,col=2)
```

## Algumas propriedades da normal multivariada

Seja  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ . Então,

$$\mathbb{E}(\mathbf{X}) = \mathbb{E}[(X_1, \dots, X_p)^\top] = (\mathbb{E}(X_1), \dots, \mathbb{E}(X_p))^\top = (\mu_1, \dots, \mu_p)^\top = \boldsymbol{\mu}.$$

$$\text{Var}(\mathbf{X}) = \mathbb{E}[(\mathbf{X} - \boldsymbol{\mu})(\mathbf{X} - \boldsymbol{\mu})^\top] = \mathbb{E}(\mathbf{X}\mathbf{X}^\top) - \boldsymbol{\mu}\boldsymbol{\mu}^\top = \boldsymbol{\Sigma},$$

em que  $\mathbb{E}(\mathbf{X}) = \boldsymbol{\mu}$ .

Note que  $\text{Var}(\mathbf{X})_{j,k} = \text{Cov}(X_j, X_k)$  para  $j, k = 1, \dots, p$ .

A matriz de covariâncias,  $\boldsymbol{\Sigma}$ , é simétrica e positiva definida.

Seja  $\mathbf{A}$  uma matriz  $m \times p$ ,  $\mathbf{b}$  um vetor  $m \times 1$  e  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ . Então,

$$\mathbf{Y} = \mathbf{A}\mathbf{X} + \mathbf{b} \sim \mathcal{N}_m(\boldsymbol{\xi}, \boldsymbol{\Psi}),$$

em que  $\boldsymbol{\xi} = \mathbf{A}\boldsymbol{\mu} + \mathbf{b}$  e  $\boldsymbol{\Psi} = \mathbf{A}\boldsymbol{\Sigma}\mathbf{A}^\top$ .

**Combinação linear de variáveis aleatórias normais (correlacionadas ou não) tem distribuição normal.**

## Decomposição da matriz $\Sigma$

- Podemos sempre representar a distribuição normal multivariada como uma transformação linear de variáveis aleatórias normais univariadas?
- Podemos sempre representar a distribuição normal multivariada como uma transformação linear de variáveis aleatórias normais **padrão** univariadas?
- Podemos sempre ver a distribuição normal multivariada como uma elipse  $p$ -dimensional?

**Sim** é a resposta para todas as perguntas acima.

Podemos, por exemplo, empregar a decomposição em valores singulares (SVD) a matriz de covariâncias  $\Sigma$  para obter

$$\Sigma = \mathbf{V}\mathbf{D}\mathbf{V}^T,$$

sendo  $\mathbf{V}$  uma matriz  $p \times p$  ortonormal (as colunas da matriz são ortogonais e cada uma tem norma 1) e  $\mathbf{D}$  uma matriz  $p \times p$  diagonal com elementos  $d_i$ ,  $i = 1, \dots, p$ , positivos.

Isto significa que qualquer vetor normalmente distribuído  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ , com  $\boldsymbol{\Sigma} = \mathbf{V}\mathbf{D}\mathbf{V}^\top$ , pode ser escrito como  $\mathbf{X} = \mathbf{V}(\mathbf{W} - \boldsymbol{\mu}) + \boldsymbol{\mu}$ , em que  $\mathbf{W} \sim \mathcal{N}_p(\boldsymbol{\mu}, \mathbf{D})$  e as coordenadas de  $\mathbf{W}$ ,  $W_j$ , são independentes e normalmente distribuídas com média  $\mu_j$  e variância  $d_j$ ,  $j = 1, \dots, p$ .

Suponha que  $\mathbf{Z} \sim \mathcal{N}_p(\mathbf{0}, \mathbf{I}_p)$  e que  $\boldsymbol{\Sigma} = \mathbf{V}\mathbf{D}\mathbf{V}^\top = \mathbf{V}\mathbf{D}^{1/2}\mathbf{D}^{1/2}\mathbf{V}^\top = \mathbf{H}\mathbf{H}^\top$ , em que  $\mathbf{D}^{1/2}$  é a matriz formada pelas raízes quadradas dos elementos da diagonal de  $\mathbf{D}$ , e  $\mathbf{H} = \mathbf{V}\mathbf{D}^{1/2}$ . Então, para  $\mathbf{X} = \mathbf{H}\mathbf{Z} + \boldsymbol{\mu}$ , temos

$$\begin{aligned} \mathbb{E}(\mathbf{X}) &= \mathbb{E}(\mathbf{H}\mathbf{Z} + \boldsymbol{\mu}) = \mathbf{H}\mathbb{E}(\mathbf{Z}) + \boldsymbol{\mu} = \mathbf{H}\mathbf{0} + \boldsymbol{\mu} = \boldsymbol{\mu} \quad \text{e} \\ \text{Var}(\mathbf{X}) &= \text{Var}(\mathbf{H}\mathbf{Z}) = \mathbf{H}\text{Var}(\mathbf{Z})\mathbf{H}^\top = \mathbf{H}\mathbf{I}\mathbf{H}^\top \\ &= \mathbf{V}\mathbf{D}^{1/2}\mathbf{D}^{1/2}\mathbf{V}^\top = \mathbf{V}\mathbf{D}\mathbf{V}^\top = \boldsymbol{\Sigma}. \end{aligned}$$

Portanto,  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ .

Também podemos utilizar a **decomposição de Cholesky** para obter  $\boldsymbol{\Sigma} = \mathbf{U}^\top \mathbf{U} = \mathbf{L}\mathbf{L}^\top$ , em que  $\mathbf{U}$  e  $\mathbf{L}$  são matrizes triangulares superior e inferior, respectivamente. Neste caso, temos que  $\mathbf{L} = \mathbf{U}^\top$ .

**Gere valores da normal multivariada utilizando o SVD no programa R.**

## Gerando valores da normal multivariada

(Ver exemplo\_04.r)

```
p      <- 3          # dimensao
n      <- 100000     # tamanho da amostra
# estrutura de covariancia (correlacao)
SIGMA <- matrix(c(1,0.9,0.8,0.9,1,0.7,0.8,0.7,1),p,p,T)
# normais independentes
x      <- matrix(rnorm(p*n),n,p) # note a transposicao
#-----
# Usando o SVD
S      <- svd(SIGMA)          # decomposicao de SIGMA
T      <- S$u%*%diag(sqrt(S$d)) # construcao da matriz T
y      <- x%*%t(T)            # gerando valores
cor(y); SIGMA
#-----
# Usando a decomposicao de Cholesky
T      <- chol(SIGMA)
y      <- x%*%T
cor(y); SIGMA
```

## Distribuição normal multivariada: marginais e condicionais

Seja  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$  tal que

$$\mathbf{X} = \begin{bmatrix} \mathbf{X}_1 \\ \mathbf{X}_2 \end{bmatrix} \text{ e } \boldsymbol{\mu} = \begin{bmatrix} \boldsymbol{\mu}_1 \\ \boldsymbol{\mu}_2 \end{bmatrix}, \text{ com dimensões } \begin{bmatrix} q \times 1 \\ (p-q) \times 1 \end{bmatrix}, \text{ e}$$

$$\boldsymbol{\Sigma} = \begin{bmatrix} \boldsymbol{\Sigma}_{11} & \boldsymbol{\Sigma}_{12} \\ \boldsymbol{\Sigma}_{21} & \boldsymbol{\Sigma}_{22} \end{bmatrix}, \text{ com dimensões } \begin{bmatrix} q \times q & q \times (p-q) \\ (p-q) \times q & (p-q) \times (p-q) \end{bmatrix}.$$

em que  $p > q \geq 1$ . Então,

$$\mathbf{X}_1 \sim \mathcal{N}_q(\boldsymbol{\mu}_1, \boldsymbol{\Sigma}_{11}),$$

$$\mathbf{X}_2 \sim \mathcal{N}_{(p-q)}(\boldsymbol{\mu}_2, \boldsymbol{\Sigma}_{22}), \text{ e}$$

$$(\mathbf{X}_1 | \mathbf{X}_2 = \boldsymbol{\xi}) \sim \mathcal{N}_q(\tilde{\boldsymbol{\mu}}, \tilde{\boldsymbol{\Sigma}}),$$

em que

$$\tilde{\boldsymbol{\mu}} = \boldsymbol{\mu}_1 + \boldsymbol{\Sigma}_{12} \boldsymbol{\Sigma}_{22}^{-1} (\boldsymbol{\xi} - \boldsymbol{\mu}_2)$$

$$\tilde{\boldsymbol{\Sigma}} = \boldsymbol{\Sigma}_{11} - \boldsymbol{\Sigma}_{12} \boldsymbol{\Sigma}_{22}^{-1} \boldsymbol{\Sigma}_{21}.$$

## Algumas propriedades

- Se  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ , então  $(\mathbf{X} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu}) \sim \chi_p^2$ .
- Calculando uma região (elipse) que tenha probabilidade  $\eta$ , i.é.,

$$\Pr((\mathbf{X} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu}) \leq \chi_\eta^2) = \eta,$$

com  $\chi_\eta^2$  representando o  $\eta$ -ésimo percentil da distribuição  $\chi_p^2$ .

- Sejam  $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$  uma AAS de  $\mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$  e  $T^2 = n(\bar{\mathbf{X}} - \boldsymbol{\mu})^\top \mathbf{S}^{-1}(\bar{\mathbf{X}} - \boldsymbol{\mu})$  (estatística  $T^2$  de Hotelling). Então,

$$\frac{n-p}{(n-1)p} T^2 \sim \mathcal{F}_{p, n-p}.$$

Natureza da estatística  $T^2$ :

$$\underbrace{\sqrt{n}(\bar{\mathbf{X}} - \boldsymbol{\mu})^\top}_{\mathcal{N}_p(\mathbf{0}, \boldsymbol{\Sigma})} \left\{ \overbrace{\frac{1}{n-1} \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^\top}^{\left( \frac{\mathcal{W}_p(n-1, \boldsymbol{\Sigma})}{n-1} \right)} \right\}^{-1} \underbrace{\sqrt{n}(\bar{\mathbf{X}} - \boldsymbol{\mu})^\top}_{\mathcal{N}_p(\mathbf{0}, \boldsymbol{\Sigma})}.$$

## Exercício

- Escreva um programa no R para gerar números do vetor aleatório normalmente distribuído com variâncias todas iguais a 1 (um) e correlações todas iguais a um dos valores no conjunto  $\{-0,9; 0,0; 0,9; 0,99\}$ , e varie o tamanho das amostras geradas para cada matriz de covariâncias (por exemplo, 100, 1000, etc). Considere uma normal trivariada para este exercício. Faça diversos gráficos para analisar os conjuntos de dados gerados. Por exemplo *box-plot*, histogramas, *qq-plot*, dispersão dois a dois, etc. É importante neste exercício que você trabalhe com a matriz de dados de acordo com a parte teórica da matéria, ou seja, salve os dados gerados em uma matriz  $p \times n$ .

Dica: verificar as funções `boxplot`, `hist`, `qqnorm` e `plot` do programa R.

## Análise em Componentes Principais - ACP

O propósito principal da ACP é substituir as variáveis originais por um número menor de variáveis que são função das variáveis originais.

A ACP consiste na determinação de uma transformação ortogonal das variáveis originais para um novo conjunto de variáveis não correlacionadas que são obtidas em ordem decrescente de importância.

As novas variáveis, chamadas de componentes, são combinações lineares das variáveis originais.

De forma a estudar as relações entre um conjunto de  $p$  variáveis correlacionadas, pode ser útil transformar o conjunto de dados em um novo conjunto de variáveis não correlacionadas chamados **Componentes Principais**. Estas novas variáveis são combinações lineares das variáveis originais e são obtidas em ordem decrescente de importância tal que, por exemplo, o primeiro componente principal conta com a maior parte possível da variação total nos dados originais. A transformação ortogonal é, de fato, uma rotação espaço original das  $p$  variáveis.

A técnica de obtenção desta transformação é chamada **Análise em Componentes Principais** (ACP). A ACP é apropriada quando as variáveis sob investigação são todas de mesma natureza, de modo que não tenhamos, por exemplo, uma variável dependente e várias variáveis explicativas como no caso da regressão múltipla.

O objetivo usual da análise é verificar se os primeiros poucos componentes principais (CP's) compreendem a maior parte da variação total dos dados originais. Se for este o caso, então coloca-se a questão de que a dimensionalidade efetiva dos dados é menor do que  $p$ .

Em outras palavras, se algumas das variáveis originais são altamente correlacionadas, elas estão, efetivamente, “dizendo a mesma coisa” e podem portanto, existir restrições lineares sobre estas variáveis. Neste caso é esperado que os primeiros componentes sejam significativos, nos ajudem a compreender melhor os dados e sejam úteis nas análises subsequentes, nas quais poderemos lidar com um número menor de variáveis.

Na prática, em diversos casos, não é fácil interpretar os componentes e, portanto, sua principal utilidade está na redução da dimensionalidade dos dados de forma a simplificar as análises posteriores.

## Componentes Principais Populacionais

Algebricamente, CP's são combinações lineares particulares das variáveis  $X_1, X_2, \dots, X_p$ .

Geometricamente, estas combinações lineares representam a seleção de um novo sistema de coordenadas rotacionando-se o sistema original. Os novos eixos representam as direções com variabilidade máxima e fornecem descrição mais simples e parcimoniosa da estrutura de covariância.

Como veremos, os CP's dependem somente da estrutura de covariância  $\Sigma$  (ou da matriz de correlações  $\mathcal{R}$ ). Seu desenvolvimento não exige que os dados sejam normais multivariados. Por outro lado, componentes principais obtidos de populações normais multivariadas têm interpretações úteis em função dos elipsóides de densidade constante. Além disso, inferências podem ser feitas a partir dos componentes amostrais, quando a população é normal multivariada.

Suponha  $\mathbf{X}$  uma v.a.  $p$ -dimensional,  $\mathbf{X}^\top = (X_1, X_2, \dots, X_p)$ , com média  $\mu$  e matriz de covariâncias  $\Sigma$ . Nosso problema aqui é encontrar um novo conjunto de variáveis, a saber  $\mathbf{Y}^\top = (Y_1, Y_2, \dots, Y_p)$ , tal que  $Y_1, Y_2, \dots, Y_p$  são não correlacionadas e cujas variâncias decrescem da primeira para a última.

Cada  $Y_j$  é tomada como uma combinação linear dos componentes do vetor  $\mathbf{X}$  tal que

$$Y_j = a_{j1}X_1 + a_{j2}X_2 + \dots + a_{jp}X_p = \mathbf{a}_j^\top \mathbf{X} \quad (1)$$

em que  $\mathbf{a}_j^\top = (a_{j1}, a_{j2}, \dots, a_{jp})$  é um vetor de constantes.

A Equação 1 contém um fator de escala arbitrário. Portanto, impomos a condição  $\mathbf{a}_j^\top \mathbf{a}_j = \sum_{k=1}^p a_{jk}^2 = 1$ . Veremos que este procedimento particular de normalização assegura que a transformação global é ortogonal. Em outras palavras, assegura que distâncias e ângulos no novo espaço são preservados.

O primeiro componente principal,  $Y_1$ , é encontrado escolhendo-se  $\mathbf{a}_1$  tal que  $Y_1$  tenha a maior variância possível. Em outras palavras, escolhemos  $\mathbf{a}_1$  de forma a

$$\begin{array}{ll} \text{Maximizar} & \text{Var}(\mathbf{a}_1^\top \mathbf{X}) \\ \text{Sujeito à} & \mathbf{a}_1^\top \mathbf{a}_1 = 1. \end{array}$$

O segundo componente principal é encontrado escolhendo-se  $\mathbf{a}_2$  tal que  $Y_2$  tenha a (segunda) maior variância possível para todos os compostos definidos pela equação 1 que são não correlacionados com  $Y_1$ .

Similarmente, derivamos os componentes  $Y_3, Y_4, \dots, Y_p$  de forma que eles sejam não correlacionados com os anteriores e tenham variâncias decrescentes.

**Proposição 1:** Seja  $\Sigma$  a matriz de covariâncias de  $\mathbf{X}$ . Sejam  $(\lambda_1, \mathbf{v}_1), (\lambda_2, \mathbf{v}_2), \dots, (\lambda_p, \mathbf{v}_p)$  os pares de autovalores e autovetores (ortonormalizados) de  $\Sigma$  tais que  $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p \geq 0$ .

O  $j$ -ésimo componente principal é dado por

$$Y_j = \mathbf{v}_j^\top \mathbf{X}, \quad j = 1, 2, \dots, p.$$

Com estas escolhas, tem-se:

$$\text{Var}(Y_j) = \mathbf{v}_j^\top \Sigma \mathbf{v}_j = \lambda_j, \quad j = 1, 2, \dots, p$$

e

$$\text{Cov}(Y_j, Y_k) = \mathbf{v}_j^\top \Sigma \mathbf{v}_k = 0, \quad j \neq k.$$

Se alguns  $\lambda_j$  forem repetidos, as escolhas dos correspondentes autovetores  $\mathbf{v}_j$  e, portanto de  $Y_j$ , não será única.

**Proposição 2:** Seja  $\mathbf{X}$  vetor aleatório de dimensão  $p \times 1$  com matriz de covariâncias  $\mathbf{\Sigma}$  tal que  $(\lambda_1, \mathbf{v}_1), (\lambda_2, \mathbf{v}_2), \dots, (\lambda_p, \mathbf{v}_p)$  são os pares de autovalores e autovetores (ortonormalizados) de  $\mathbf{\Sigma}$  com  $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p \geq 0$ . Então,

$$\sum_{j=1}^p \text{Var}(X_j) = \sum_{j=1}^p \sigma_{jj} = \lambda_1 + \lambda_2 + \dots + \lambda_p = \sum_{j=1}^p \text{Var}(Y_j).$$

**Dem.:** Temos que  $\sum_{j=1}^p \sigma_{jj} = \text{tr}(\mathbf{\Sigma})$ . Mas  $\mathbf{\Sigma} = \mathbf{P}\mathbf{\Delta}\mathbf{P}^\top$  com  $\mathbf{\Delta} = \text{diag}\{\lambda_1, \lambda_2, \dots, \lambda_p\}$  e  $\mathbf{P} = [\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_p]$ , matriz ortogonal.

Logo,

$$\text{tr}(\mathbf{\Sigma}) = \text{tr}(\mathbf{P}\mathbf{\Delta}\mathbf{P}^\top) = \text{tr}(\mathbf{P}^\top \mathbf{P}\mathbf{\Delta}) = \text{tr}(\mathbf{\Delta}) = \sum_{j=1}^p \lambda_j.$$

A Proposição 2 diz que a variância total da população é dada pela soma dos autovalores da matriz de covariâncias populacional,  $\mathbf{\Sigma}$ .

Consequentemente, a **proporção da variação total devida ao  $j$ -ésimo componente** é

$$\frac{\lambda_j}{\lambda_1 + \lambda_2 + \dots + \lambda_p}, \quad \text{para } j = 1, 2, \dots, p.$$

Se a maior parte (por exemplo 80% a 90%) da variação total, para  $p$  grande, pode ser atribuído aos dois ou três primeiros componentes, então estes componentes podem “substituir” as  $p$  variáveis originais sem muita perda de informação.

Cada entrada dos vetores que definem os componentes

$\mathbf{v}_j^\top = [v_{j1}, v_{j2}, \dots, v_{jp}]$  também deve ser olhada. A magnitude de  $v_{jk}$  indica a importância da  $k$ -ésima variável original no  $j$ -ésimo componente, sem olhar as demais variáveis.

Em particular,  $v_{jk}$  é proporcional ao coeficiente de correlação entre  $Y_j$  e  $X_k$ .

**Proposição 3:** Se  $Y_j = \mathbf{v}_j^\top \mathbf{X}$ ,  $j = 1, 2, \dots, p$  são os componentes principais obtidos a partir da matriz de covariâncias  $\mathbf{\Sigma}$  do vetor aleatório  $\mathbf{X}$ , então

$$\rho_{Y_j, X_k} = \frac{v_{jk} \sqrt{\lambda_j}}{\sqrt{\sigma_{kk}}}, \quad j, k = 1, 2, \dots, p$$

são os coeficientes de correlação entre  $Y_j$  e  $X_k$ .

$(\lambda_1, \mathbf{v}_1), (\lambda_2, \mathbf{v}_2), \dots, (\lambda_p, \mathbf{v}_p)$  são os pares de autovalores e autovetores (ortonormalizados) de  $\mathbf{\Sigma}$  com  $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p \geq 0$ .

(Ver exemplo\_10.r)

**Exemplo 1:** Suponha um vetor aleatório de dimensão 3 cuja matriz de covariâncias é dada por (Ver exemplo\_05.r)

$$\Sigma = \begin{bmatrix} 1 & -2 & 0 \\ -2 & 5 & 0 \\ 0 & 0 & 2 \end{bmatrix}.$$

Os pares de autovalores e autovetores associados a  $\Sigma$  são:

$$\begin{aligned} \lambda_1 &\simeq 5,83 & \text{e} & \mathbf{v}_1^\top \simeq [-0,383 \quad 0,924 \quad 0]; \\ \lambda_2 &\simeq 2,00 & \text{e} & \mathbf{v}_2^\top \simeq [0 \quad 0 \quad 1]; & \text{e} \\ \lambda_3 &\simeq 0,17 & \text{e} & \mathbf{v}_3^\top \simeq [0,924 \quad 0,383 \quad 0]. \end{aligned}$$

Portanto, os componentes principais são:

$$\begin{aligned} Y_1 &= -0,383X_1 + 0,924X_2 \\ Y_2 &= X_3 \\ Y_3 &= 0,924X_1 + 0,383X_2. \end{aligned}$$

A variável  $X_3$  é um dos componentes principais, pois ela é não correlacionada com as outras duas variáveis.

A proporção da variação total explicada pelos dois primeiros componentes é  $7,83/8 \simeq 0,98$  de modo que podemos explicar estes dados usando apenas  $Y_1$  e  $Y_2$ , sem perder quase informação.

As correlações entre os componentes e as respectivas variáveis são dadas na tabela a seguir.

**Tabela 1:** Correlações entre os componentes e as respectivas variáveis.

| CP    | $X_1$  | $X_2$ | $X_3$ |
|-------|--------|-------|-------|
| $Y_1$ | -0,924 | 0,997 | 0     |
| $Y_2$ | 0      | 0     | 1     |
| $Y_3$ | 0,383  | 0,071 | 0     |

É informativo considerar componentes principais obtidos de populações normais multivariadas. Suponha  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ . Vimos que a densidade de  $\mathbf{X}$  é constante para os pontos  $\mathbf{x}$  em  $\mathbb{R}^p$  que satisfazem

$$(\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu}) = c^2.$$

As soluções desta equação correspondem ao elipsóide centrado em  $\boldsymbol{\mu}$  cujos semi-eixos são dados por  $\pm c\sqrt{\lambda_j} \mathbf{v}_j$ ,  $j = 1, 2, \dots, p$  em que  $(\lambda_1, \mathbf{v}_1), (\lambda_2, \mathbf{v}_2), \dots, (\lambda_p, \mathbf{v}_p)$  são os pares de autovalores e autovetores (ortonormalizados) de  $\boldsymbol{\Sigma}$ .

Um ponto caindo no  $j$ -ésimo eixo do elipsóide de densidade constante terá coordenadas proporcionais às coordenadas do vetor

$$\mathbf{v}_j^\top = [v_{j1}, v_{j2}, \dots, v_{jp}],$$

no sistema de coordenadas que tem origem em  $\boldsymbol{\mu}$  e eixos que são paralelos aos eixos das variáveis originais  $X_1, X_2, \dots, X_p$ .

Será conveniente aqui considerar  $\boldsymbol{\mu} = \mathbf{0}$ , sem perda de generalidade.

Vimos que os pares de autovalores e autovetores de  $\boldsymbol{\Sigma}^{-1}$  são dados por  $(1/\lambda_1, \mathbf{v}_1), (1/\lambda_2, \mathbf{v}_2), \dots, (1/\lambda_p, \mathbf{v}_p)$ .

Com  $\boldsymbol{\mu} = \mathbf{0}$ , temos

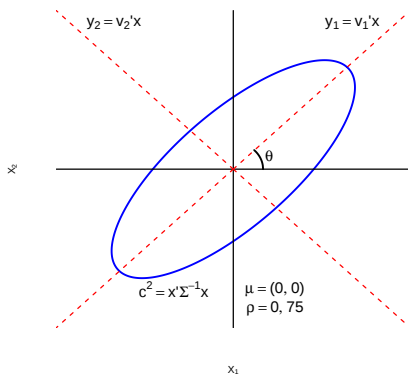
$$\mathbf{x}^\top \boldsymbol{\Sigma}^{-1} \mathbf{x} = \sum_{j=1}^p \frac{1}{\lambda_j} (\mathbf{v}_j^\top \mathbf{x})^2 = \sum_{j=1}^p \frac{1}{\lambda_j} y_j^2 = \frac{1}{\lambda_1} y_1^2 + \frac{1}{\lambda_2} y_2^2 + \dots + \frac{1}{\lambda_p} y_p^2 = c^2,$$

e esta equação define um elipsóide (pois  $\lambda_1, \lambda_2, \dots, \lambda_p$  são positivos) em um sistema de coordenadas  $y_1, y_2, \dots, y_p$  definido pelos eixos cujas direções são  $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_p$  respectivamente.

Se  $\lambda_1$  é o maior autovalor, então o eixo maior está na direção do autovetor  $\mathbf{v}_1$ . Os eixos menores restantes estão nas direções definidas pelos autovetores  $\mathbf{v}_2, \dots, \mathbf{v}_p$ .

**Resumindo:** os componentes principais caem nas direções dos eixos de um elipsóide de densidade constante. Portanto, qualquer ponto sobre o  $j$ -ésimo eixo do elipsóide tem coordenadas proporcionais ao  $j$ -ésimo autovetor.

Quando  $\mu \neq \mathbf{0}$ , seu componente principal centrado na média é  $y_j = \mathbf{v}_j^\top (\mathbf{x} - \mu)$  que tem média zero e está na direção  $\mathbf{v}_j$ .



**Figura 4:** Elipse de densidade constante e os componentes principais para um vetor aleatório normal bivariado com  $\mu = \mathbf{0}$  e  $\rho = 0,75$ .

É possível ver que os CP's são obtidos pela rotação em relação à origem, dos eixos originais de um ângulo  $\theta$  no sentido anti-horário até que eles coincidam com os eixos da elipse de densidade constante. Este resultado também vale para  $\rho > 2$ .

## Componentes Principais Obtidos das Variáveis Padronizadas

Os componentes também podem ser obtidos via variáveis padronizadas. Na notação matricial  $\mathbf{Z} = (\Psi^{1/2})^{-1}(\mathbf{X} - \mu)$  com  $\Psi^{1/2} = \text{diag}\{\sqrt{\sigma_{11}}, \sqrt{\sigma_{22}}, \dots, \sqrt{\sigma_{pp}}\}$ .

Claramente  $E(\mathbf{Z}) = \mathbf{0}$  e  $\text{Var}(\mathbf{Z}) = \mathcal{R}$ .

Os componentes principais de  $\mathbf{Z}$  podem ser obtidos da matriz de correlações  $\mathcal{R}$  de  $\mathbf{X}$ .

Os resultados anteriores todos se aplicam, com algumas simplificações, pois agora a variância total é  $p$ , pois na diagonal de  $\mathcal{R}$  os elementos são todos unitários.

Nesse caso, a proporção total da variação devida ao  $j$ -ésimo componente será dada por  $\frac{\delta_j}{p}$ , para  $j = 1, 2, \dots, p$ , sendo  $\delta_j$  o  $j$ -ésimo autovalor da matriz de correlações  $\mathcal{R}$ .

Porém, os componentes obtidos a partir da matriz de correlações serão diferentes dos componentes obtidos pela matriz de covariâncias.

**Exemplo 2:** Considere a seguinte matriz de covariâncias  $\Sigma = \begin{bmatrix} 1 & 4 \\ 4 & 100 \end{bmatrix}$   
e a correspondente matriz de correlações  $\mathcal{R} = \begin{bmatrix} 1 & 0,4 \\ 0,4 & 1 \end{bmatrix}$ .

(Ver exemplo\_06.r)

Os pares de autovalores e autovetores de  $\Sigma$  são dados por

$$\begin{aligned} \lambda_1 &= 100,16 & \text{e} & \mathbf{v}_1^\top = [0,040 \quad -0,999] \\ \lambda_2 &= 0,839 & \text{e} & \mathbf{v}_2^\top = [0,999 \quad 0,040]. \end{aligned}$$

Os pares de autovalores e autovetores de  $\mathcal{R}$  são dados por

$$\begin{aligned} \delta_1 &= 1,4 & \text{e} & \mathbf{w}_1^\top = [0,707 \quad 0,707] \\ \delta_2 &= 0,6 & \text{e} & \mathbf{w}_2^\top = [-0,707 \quad 0,707]. \end{aligned}$$

Os respectivos componentes principais são, então,

- relativos a  $\Sigma$ :

$$Y_1 = 0,040X_1 - 0,999X_2, \quad \text{e}$$

$$Y_2 = 0,999X_1 + 0,040X_2$$

com  $Y_1$  correspondendo a 99,17% da variação total.

- relativos a  $\mathcal{R}$ :

$$Y_1^* = 0,707 \frac{(X_1 - \mu_1)}{1} + 0,707 \frac{(X_2 - \mu_2)}{10} = 0,707Z_1 + 0,707Z_2, \quad \text{e}$$

$$Y_2^* = -0,707 \frac{(X_1 - \mu_1)}{1} + 0,707 \frac{(X_2 - \mu_2)}{10} = -0,707Z_1 + 0,707Z_2$$

com  $Y_1^*$  correspondendo a 70% da variação total.

Devido à variância de  $X_2$  ser muito maior que a de  $X_1$ , esta variável domina o primeiro componentes principal determinado pela matriz  $\Sigma$ .

Quando as variáveis  $X_1$  e  $X_2$  são padronizadas, as variáveis resultantes contribuem igualmente para os componentes principais determinados pela matriz  $\mathcal{R}$  e as correlações de  $Z_1$  e  $Z_2$  com o primeiro componente são 0,837.

Mais surpreendentemente, vemos que a importância relativa das variáveis para o primeiro componente é fortemente afetada pela padronização. Quando o primeiro componente obtido de  $\mathcal{R}$  é expresso em função de  $X_1$  e  $X_2$ , as magnitudes relativas dos pesos que são 0,707 e  $0,707/10=0,0707$  estão em oposição direta aos pesos do primeiro componente relativo a  $\Sigma$  dados por 0,040 e 0,999, respectivamente.

Este exemplo demonstra que componentes principais obtidas de  $\Sigma$  são diferentes daqueles obtidos de  $\mathcal{R}$ . Além disso, um conjunto de componentes principais não é uma função simples do outro. Isto sugere que a padronização não é inconsequente.

As variáveis devem provavelmente ser padronizadas se elas são medidas sobre escalas cujas variações são muito diferentes ou se as unidades de medida não são comensuráveis. Por exemplo, se  $X_1$  representa vendas anuais de dez mil a 350 mil dólares e  $X_2$  é a razão (renda bruta anual)/(ativos totais) que varia entre 0,01 a 0,60, então a variação total será quase que exclusivamente devida às vendas em dólar. Neste caso, esperaríamos um único componente principal importante com um peso alto em  $X_1$ . Alternativamente, se ambas as variáveis são padronizadas, suas magnitudes subsequentes serão da mesma ordem, e  $X_2$  (ou  $Z_2$ ) representarão um papel maior na construção dos componentes principais.

## Componentes Principais Amostrais

Agora temos a estrutura necessária para estudar o problema de resumir a variação nas  $n$  observações sobre  $p$  variáveis em poucas combinações lineares criteriosamente escolhidas.

Suponha que os dados  $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$  representam  $n$  observações independentes de alguma população  $p$ -dimensional com vetor de médias  $\boldsymbol{\mu}$  e matriz de covariâncias  $\boldsymbol{\Sigma}$ . Estes dados produzem um vetor de médias amostral  $\bar{\mathbf{x}}$ , a matriz de covariâncias amostral  $\mathbf{S}$  e a matriz de correlações amostral  $\mathbf{R}$ .

O objetivo aqui será construir combinações lineares não correlacionadas a partir das características medidas que levam em conta o máximo da variação na amostra. As combinações não correlacionadas com as maiores variâncias serão chamadas **componentes principais amostrais**.

**Proposição 4:** Sejam  $\mathbf{S}$  a matriz de covariância amostral e  $(\hat{\lambda}_j, \hat{\mathbf{v}}_j)$ ,  $j = 1, 2, \dots, p$  seus correspondentes pares de autovalores e autovetores, com  $\hat{\lambda}_1 \geq \hat{\lambda}_2 \geq \dots \geq \hat{\lambda}_p \geq 0$ .

Seja  $\mathbf{x}$  uma observação qualquer.

Variância amostral  $\widehat{\text{Var}}(\hat{y}_j) = \hat{\lambda}_j$ , para  $j = 1, 2, \dots, p$ .

Covariância amostral  $\widehat{\text{Cov}}(\hat{y}_j, \hat{y}_k) = 0$ , para  $j \neq k$  e  $j, k = 1, 2, \dots, p$ .

Variância amostral total  $\sum_{j=1}^p s_{jj} = \sum_{j=1}^p \hat{\lambda}_j$ .

Correlação do  $j$ -ésimo componente principal amostral com a  $k$ -ésima variável

$$r_{\hat{y}_j, x_k} = \frac{\hat{v}_{jk} \sqrt{\hat{\lambda}_j}}{\sqrt{s_{kk}}}, \quad j, k = 1, 2, \dots, p.$$

Vamos denotar

$$\hat{\mathbf{\Lambda}} = \text{diag}(\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_p) \quad \text{e} \quad \hat{\mathbf{\Lambda}}^{1/2} = \text{diag} \left( \sqrt{\hat{\lambda}_1}, \sqrt{\hat{\lambda}_2}, \dots, \sqrt{\hat{\lambda}_p} \right).$$

**Exemplo 3:** Um censo forneceu informações sobre 5 variáveis socioeconômicas em 14 setores censitários na área de Madison, Wisconsin. Os dados de 61 setores censitários estão disponíveis no arquivo T8-5.dat. As variáveis observadas foram  $X_1$  - população (milhares),  $X_2$  - grau profissional,  $X_3$  - população empregada acima de 16 anos,  $X_4$  - funcionários públicos e  $X_5$  - valor mediano da residência.

**PROBLEMA:** A variação amostral pode ser resumida por um ou dois componentes principais? (Ver exemplo\_07.r)

**Vetor de médias amostral das variáveis:**

| Variável             | pop. | prof. | emprego | func. pub. | casa |
|----------------------|------|-------|---------|------------|------|
| $\bar{\mathbf{x}}^T$ | 4,47 | 3,96  | 71,42   | 26,91      | 1,64 |

### Matriz de covariâncias amostrais das variáveis:

|            | pop.   | prof.  | emprego | func. pub. | casa   |
|------------|--------|--------|---------|------------|--------|
| pop.       | 3,397  | -1,102 | 4,306   | -2,078     | 0,027  |
| prof.      | -1,102 | 9,673  | -1,513  | 10,953     | 1,203  |
| emprego    | 4,306  | -1,513 | 55,626  | -28,938    | -0,044 |
| func. pub. | -2,078 | 10,953 | -28,938 | 89,067     | 0,957  |
| casa       | 0,027  | 1,203  | -0,044  | 0,957      | 0,319  |

### Matriz de correlações amostrais das variáveis:

|            | pop.  | prof. | emprego | func. pub. | casa  |
|------------|-------|-------|---------|------------|-------|
| pop.       | 1,00  | -0,19 | 0,31    | -0,12      | 0,03  |
| prof.      | -0,19 | 1,00  | -0,07   | 0,37       | 0,69  |
| emprego    | 0,31  | -0,07 | 1,00    | -0,41      | -0,01 |
| func. pub. | -0,12 | 0,37  | -0,41   | 1,00       | 0,18  |
| casa       | 0,03  | 0,69  | -0,01   | 0,18       | 1,00  |

No R, a função `prcomp` do pacote `stats` roda a ACP.

Desvios padrão:

$$\hat{\Lambda}^{1/2} = [10,3448 \quad 6,2986 \quad 2,8932 \quad 1,6935 \quad 0,3933]$$

Rotação:

|            | $\hat{\mathbf{v}}_1$ | $\hat{\mathbf{v}}_2$ | $\hat{\mathbf{v}}_3$ | $\hat{\mathbf{v}}_4$ | $\hat{\mathbf{v}}_5$ |
|------------|----------------------|----------------------|----------------------|----------------------|----------------------|
| pop.       | 0,0389               | -0,0711              | 0,1879               | 0,9771               | -0,0577              |
| prof.      | -0,1053              | -0,1298              | -0,9610              | 0,1714               | -0,1386              |
| emprego    | 0,4924               | -0,8644              | 0,0458               | -0,0910              | 0,0050               |
| func. pub. | -0,8631              | -0,4803              | 0,1532               | -0,0297              | 0,0067               |
| casa       | -0,0091              | -0,0147              | -0,1250              | 0,0817               | 0,9886               |

Importância das componentes:

|                          | CP1    | CP2   | CP3   | CP4   | CP5   |
|--------------------------|--------|-------|-------|-------|-------|
| $\sqrt{\hat{\lambda}_j}$ | 10,345 | 6,299 | 2,893 | 1,693 | 0,393 |
| Prop. Variância          | 0,677  | 0,251 | 0,053 | 0,018 | 0,001 |
| Prop. Acumulada          | 0,677  | 0,928 | 0,981 | 0,999 | 1,000 |

Os dois primeiros componentes principais, juntos, explicam 92,8% da variância amostral total. Assim, podemos dizer que a variação amostral é bem resumida por dois componentes principais e uma redução nos dados de 61 observações de 5 variáveis para 61 observações de dois componentes é razoável. Dados os coeficientes dos vetores que definem os componentes, o primeiro parece ser essencialmente uma diferença ponderada entre as variáveis funcionário público e empregados maiores de 16 anos. O segundo CP parece ser uma soma ponderada dos dois.

### Correlações entre os componentes e as variáveis

|       | $X_1$  | $X_2$  | $X_3$  | $X_4$  | $X_5$  |
|-------|--------|--------|--------|--------|--------|
| $Y_1$ | 0,218  | -0,350 | -0,683 | -0,946 | -0,167 |
| $Y_2$ | -0,243 | -0,263 | -0,730 | -0,321 | -0,165 |

**Observação:** A resposta ao problema é sim, desde que de fato as variáveis  $X_3$  e  $X_4$  sejam mais importantes. Caso contrário, é melhor padronizar as variáveis.

## O número de componentes principais

A questão de quantos componentes reter está sempre presente. Não há uma resposta definitiva para ela. Aspectos a serem considerados incluem a quantidade total da variância explicada, o tamanho relativo dos autovalores (as variâncias dos componentes amostrais), e as interpretações dos componentes. Em adição, como veremos adiante, um componente associado com um autovalor próximo de zero e, assim, considerados sem importância, podem indicar dependências lineares insuspeitas nos dados.

Uma ajuda visual útil para determinar um número apropriado de componentes é um *scree plot*. Com os autovalores ordenados em ordem decrescente, um *scree plot* é um gráfico de  $j$  versus  $\hat{\lambda}_j$ . Para determinar o número apropriado de componentes, buscamos por um “cotovelo” (dobra) no *scree plot*.

O número de componentes é tomado como o ponto a partir do qual os autovalores restantes são relativamente pequenos e todos têm aproximadamente o mesmo tamanho.

### Exemplo 3: (Continuação)

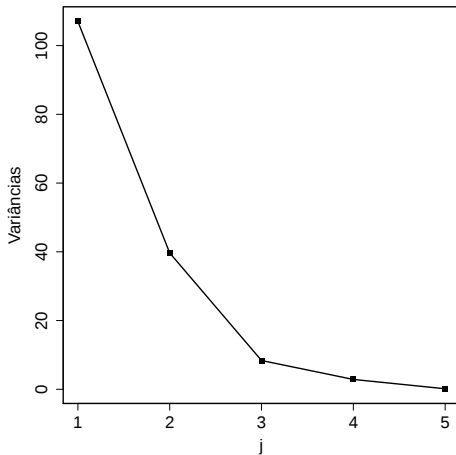


Figura 5: Autovalores ordenados de forma não crescente.

## Componentes principais usando as variáveis padronizadas

Os componentes principais também podem ser obtidos via matriz de correlações amostral  $\mathbf{R}$  e os resultados com suas devidas adaptações serão:

- Seja  $(\hat{\delta}_j, \hat{\mathbf{w}}_j)$ ,  $j = 1, 2, \dots, p$ , os pares de autovalores e autovetores de  $\mathbf{R}$ , com  $\hat{\delta}_1 \geq \hat{\delta}_2 \geq \dots \geq \hat{\delta}_p \geq 0$ . De maneira similar, temos  $\hat{\Delta}$  e  $\hat{\Delta}^{1/2}$ .
- Seja  $\mathbf{z}$  uma observação qualquer padronizada.
- O valor do  $j$ -ésimo CP para esta observação é  $\hat{y}_j^* = \hat{\mathbf{w}}_j^\top \mathbf{z}$ .
- Variância amostral  $\widehat{\text{Var}}(\hat{y}_j^*) = \hat{\delta}_j$ ,  $j = 1, 2, \dots, p$ .
- Covariância amostral  $\widehat{\text{Cov}}(\hat{y}_j^*, \hat{y}_k^*) = 0$ ,  $j \neq k$ ,  $j, k = 1, 2, \dots, p$ .
- Variância amostral total  $\text{tr}(\mathbf{R}) = p = \sum_{j=1}^p \hat{\delta}_j$ .
- Correlação do  $j$ -ésimo componente amostral com a  $k$ -ésima variável

$$r_{\hat{y}_j^*, X_k} = \hat{w}_{jk} \sqrt{\hat{\delta}_j}, \quad j, k = 1, 2, \dots, p.$$

Lembrem-se de que em geral os componentes obtidos via matriz de correlações serão diferentes daqueles obtidos via matriz de covariâncias.

#### Exemplo 4: (Componentes principais amostrais via dados padronizados)

As taxas semanais de retorno para cinco ações (JP Morgan, Citibank, Wells Fargo, Royal Dutch Shell, and ExxonMobil) listadas na Bolsa de Nova York foram registradas para o período de janeiro de 2004 a dezembro de 2005.

As taxas semanais de retorno são definidas como (preço de fechamento da semana corrente menos preço de fechamento da semana anterior)/(preço de fechamento da semana anterior), ajustada para *stock splits* e dividendos. Os dados estão disponíveis no arquivo T8-4.dat.

As observações sobre 103 semanas sucessivas parecem ser independentemente distribuídas, mas as taxas de retorno das ações são correlacionadas, pois como é de se esperar, ações tendem a mover-se junto em resposta à conjuntura econômica. (Ver exemplo\_08.r)

#### Vetor de médias amostral das variáveis:

| Variável             | JP      | Citi    | Fargo   | Royal   | Exxon   |
|----------------------|---------|---------|---------|---------|---------|
| $\bar{\mathbf{x}}^T$ | 0,00106 | 0,00066 | 0,00163 | 0,00405 | 0,00404 |

### Matriz de correlações amostrais das variáveis:

|       | JP    | Citi  | Fargo | Royal | Exxon |
|-------|-------|-------|-------|-------|-------|
| JP    | 1,000 | 0,632 | 0,510 | 0,115 | 0,154 |
| Citi  | 0,632 | 1,000 | 0,574 | 0,322 | 0,213 |
| Fargo | 0,510 | 0,574 | 1,000 | 0,182 | 0,146 |
| Royal | 0,115 | 0,322 | 0,182 | 1,000 | 0,683 |
| Exxon | 0,154 | 0,213 | 0,146 | 0,683 | 1,000 |

### Importância das componentes:

|                          | CP1    | CP2    | CP3    | CP4     | CP5     |
|--------------------------|--------|--------|--------|---------|---------|
| $\sqrt{\hat{\lambda}_j}$ | 1,5612 | 1,1862 | 0,7075 | 0,63248 | 0,50514 |
| Prop. Variância          | 0,4874 | 0,2814 | 0,1001 | 0,08001 | 0,05103 |
| Prop. Acumulada          | 0,4874 | 0,7689 | 0,8690 | 0,94897 | 1,00000 |

## Rotação:

|       | $\hat{\mathbf{v}}_1$ | $\hat{\mathbf{v}}_2$ | $\hat{\mathbf{v}}_3$ | $\hat{\mathbf{v}}_4$ | $\hat{\mathbf{v}}_5$ |
|-------|----------------------|----------------------|----------------------|----------------------|----------------------|
| JP    | -0,4691              | 0,3680               | -0,6043              | 0,3630               | 0,3841               |
| Citi  | -0,5324              | 0,2365               | -0,1361              | -0,6292              | -0,4962              |
| Fargo | -0,4652              | 0,3152               | 0,7718               | 0,2890               | 0,0712               |
| Royal | -0,3873              | -0,5850              | 0,0934               | -0,3813              | 0,5947               |
| Exxon | -0,3607              | -0,6058              | -0,1088              | 0,4934               | -0,4976              |

- O primeiro componente é uma soma ponderada das cinco ações ou um “índice”. Este componente pode ser chamado de componente geral do mercado de ações, ou simplesmente, componente de mercado.
- O segundo componente representa um contraste entre ações de bancos (JP Morgan, Citibank, Wells Fargo) e ações de petróleo (Royal Dutch Shell, ExxonMobil). Ele pode ser chamado de componente industrial.

Assim, vemos que a maior parte da variação nestes retornos de ações (77%) é devida à atividade de mercado e a atividade não correlacionada de indústria.

Os componentes restantes não são fáceis de interpretar e, coletivamente, representam variação que é provavelmente específica a cada ação.

As correlações entre os três primeiros CP's e as variáveis é dada a seguir.

Tabela 2: Correlações entre os componentes e as variáveis.

|       | $X_1$  | $X_2$  | $X_3$  | $X_4$  | $X_5$  |
|-------|--------|--------|--------|--------|--------|
| $Y_1$ | -0,732 | -0,831 | -0,726 | -0.605 | -0,563 |
| $Y_2$ | 0,437  | 0,280  | 0,374  | -0,694 | -0,719 |
| $Y_3$ | -0,428 | -0,096 | 0,546  | 0,066  | -0,077 |

A seguir mostramos o *scree plot*.

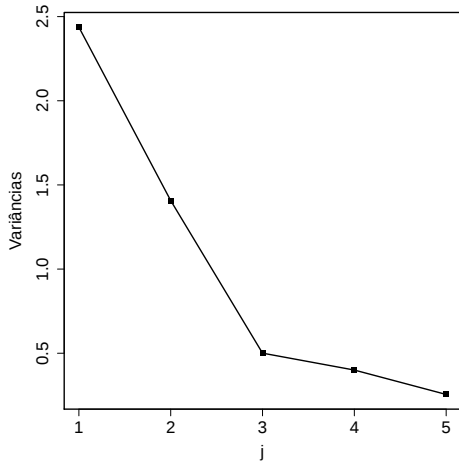


Figura 6: Autovalores ordenados de forma não crescente.

## Gráficos dos componentes principais

Gráficos dos CP's podem revelar observações suspeitas bem como desvios da suposição de normalidade. Como os CP's são combinações lineares das variáveis originais, é razoável esperar que eles comportem-se de acordo com uma distribuição normal, se estivermos supondo  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ .

Em geral é necessário verificar que os primeiros poucos CP's são aproximadamente normais, quando eles são usados como dados de entrada em análises posteriores.

Os últimos CP's podem ajudar a identificar observações suspeitas. Cada observação  $\mathbf{x}_i$  pode ser expressa como uma combinação linear da seguinte forma:

$$\mathbf{x}_i = (\mathbf{x}_i^\top \hat{\mathbf{v}}_1) \hat{\mathbf{v}}_1 + (\mathbf{x}_i^\top \hat{\mathbf{v}}_2) \hat{\mathbf{v}}_2 + \cdots + (\mathbf{x}_i^\top \hat{\mathbf{v}}_p) \hat{\mathbf{v}}_p, \quad i = 1, 2, \dots, n.$$

De fato,  $\mathbf{Y} = \mathbf{P}^\top \mathbf{X}$ , com  $\mathbf{P}$  matriz ortogonal cujas colunas são os autovetores ortonormalizados. Logo, pré-multiplicando pela matriz  $\mathbf{P}$  ambos os lados, obtemos  $\mathbf{X} = \mathbf{P}\mathbf{Y}$ , ou seja, as variáveis originais também podem, obviamente, ser representadas como combinações lineares dos CP's.

Assim, as magnitudes dos últimos CP's determinarão quão bem os poucos primeiros ajustam as observações. Isto é,

$$\sum_{j=1}^{q-1} \hat{y}_{ij} \hat{\mathbf{v}}_j \text{ difere de } \mathbf{x}_i \text{ por } \sum_{j=q}^p \hat{y}_{ij} \mathbf{v}_j.$$

O módulo quadrado deste último vetor é dado por

$$\hat{y}_{iq}^2 + \hat{y}_{i(q+1)}^2 + \cdots + \hat{y}_{ip}^2.$$

Observações suspeitas frequentemente serão tais que pelo menos uma das coordenadas  $\hat{y}_{ij}$ ,  $j = q, q+1, \dots, p$  é grande.

### Recomendações:

1. Para ajudar na verificação da normalidade, construa diagramas de dispersão para os pares dos poucos primeiros CP's que explicam uma proporção alta da variação total. Também construa os *qq-plots* dos valores amostrais gerados por cada componente.
2. Construa os diagramas de dispersão 2 a 2 e os *qq-plots* para os últimos (poucos) CP's. Eles ajudam a identificar observações suspeitas.

## Inferências para grandes amostras

Vimos que a decomposição espectral da matriz de covariâncias (correlações) é a essência da ACP.

Os autovetores determinam as direções de variabilidade máxima e, os autovalores especificam as variâncias.

Quando os primeiros poucos autovalores são muito maiores que os demais, a maior parte da variação pode ser “explicada” por um número menor do que  $p$  de variáveis.

Na prática, decisões olhando a qualidade da aproximação dos CP's devem ser tomadas com base nos pares de autovalores e autovetores associados de  $\mathbf{S}$  ou  $\mathbf{R}$ ,  $(\hat{\lambda}_j, \hat{\mathbf{v}}_j)$  ou  $(\hat{\delta}_j, \hat{\mathbf{w}}_j)$ , para  $j = 1, 2, \dots, p$ .

Devido à variação inerente à amostragem, estes autovalores e autovetores serão diferentes de suas contrapartidas populacionais.

As distribuições amostrais  $\hat{\lambda}_j$  e  $\hat{\mathbf{v}}_j$  são difíceis de se obter e estão além dos objetivos desta disciplina.

No entanto, se a suposição de normalidade multivariada é válida para o vetor de observações  $\mathbf{X}$ , podemos estabelecer propriedades sob grandes amostras.

Suponha  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ , com  $\boldsymbol{\Sigma}$  positiva definida e tal que  $\lambda_1 > \lambda_2 > \cdots > \lambda_p > 0$  são seus autovalores. A única exceção a este caso ocorre quando algum autovalor tem multiplicidade maior que 1.

Geralmente as conclusões para o caso de autovalores distintos são aplicáveis a menos que existam fortes razões para acreditar que  $\boldsymbol{\Sigma}$  tenha uma estrutura especial que produza autovalores iguais. Mesmo quando a suposição de normalidade é violada, os intervalos de confiança descritos a seguir ainda fornecerão alguma indicação da incerteza sobre  $\hat{\lambda}_j$  e componentes de  $\hat{\mathbf{v}}_j$ .

Suponha  $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ . Seja  $\boldsymbol{\Lambda} = \text{diag}\{\lambda_1, \lambda_2, \dots, \lambda_p\}$ , com  $\lambda_j$ ,  $j = 1, 2, \dots, p$ , os autovalores de  $\boldsymbol{\Sigma}$ . Faça  $\boldsymbol{\lambda}^\top = (\lambda_1, \lambda_2, \dots, \lambda_p)$ . Defina

$$\mathbf{v}_j = \lambda_j \sum_{i=1, i \neq j}^p \frac{\lambda_i}{(\lambda_i - \lambda_j)^2} \mathbf{v}_i \mathbf{v}_i^\top.$$

Então

(P1)  $\sqrt{n}(\hat{\boldsymbol{\lambda}} - \boldsymbol{\lambda}) \stackrel{a}{\sim} \mathcal{N}_p(\mathbf{0}, 2\boldsymbol{\Lambda}^2).$

(P2)  $\sqrt{n}(\hat{\mathbf{v}}_j - \mathbf{v}_j) \stackrel{a}{\sim} \mathcal{N}_p(\mathbf{0}, \mathbf{V}_j).$

(P3) Cada  $\hat{\lambda}_j$  é independentemente distribuído dos elementos de seu autovetor  $\hat{\mathbf{v}}_j$ .

Logo, (P1) implica que para  $n$  grande os  $\hat{\lambda}_j$ 's são independentemente distribuídos, isto é

$$\hat{\lambda}_j \stackrel{ind, a}{\sim} \mathcal{N}(\lambda_j, 2\lambda_j^2).$$

Usando esta distribuição, podemos construir intervalos de confiança assintóticos  $100(1-\alpha)\%$  para os  $\lambda_j$ , a saber,

$$P\left(|\hat{\lambda}_j - \lambda_j| \leq z_{(1-\alpha/2)} \lambda_j \sqrt{\frac{2}{n}}\right) \simeq 1 - \alpha$$

tal que

$$IC_{100(1-\alpha)\%}(\lambda_j) = \left( \frac{\hat{\lambda}_j}{1 + z_{(1-\alpha/2)} \sqrt{\frac{2}{n}}}, \frac{\hat{\lambda}_j}{1 - z_{(1-\alpha/2)} \sqrt{\frac{2}{n}}} \right).$$

Intervalos simultâneos assintóticos para os  $\lambda_j$ 's via Bonferroni também podem ser construídos. Suponha que estejamos interessados em  $m < p$  intervalos. Então, basta substituir o quantil acumulado da distribuição normal padrão  $z_{(1-\alpha/2)}$  por  $z_{(1-\alpha/(2m))}$ .

(P2) implica que os  $\hat{\mathbf{v}}_j$  são normalmente distribuídos para  $n$  grande.

Os elementos de cada  $\hat{\mathbf{v}}_j$  são correlacionados e a correlação depende em grande extensão da separação dos autovalores  $\lambda_1, \lambda_2, \dots, \lambda_p$  (que é desconhecida) e do tamanho amostral  $n$ .

Erros padrão aproximados para os coeficientes  $\hat{v}_{kj}$  são dados pelas raízes quadradas dos elementos diagonais da matriz  $\frac{1}{n} \hat{\mathbf{E}}_j$  em que os elementos de  $\hat{\mathbf{E}}_j$  são obtidos a partir de  $\mathbf{E}_j$  substituindo suas entradas por suas estimativas  $\hat{\lambda}_j$  e  $\hat{\mathbf{v}}_j$ .

**Exemplo 4:** (Continuação) Voltando ao exemplo das ações pedem-se intervalos de confiança assintóticos simultâneos de 95% para os dois primeiros autovalores. Usar Bonferroni.

Neste caso  $\alpha = 0,05$ ,  $m = 2$ , implicando  $1 - \frac{\alpha}{2m} = 0,9875$  e  $Z_{(0,9875)} \simeq 2,24$ .

Temos também  $n = 103$ ,  $1 + z_{\left(1 - \frac{\alpha}{2m}\right)} \sqrt{\frac{2}{n}} \simeq 1,31$  e  $1 + z_{\left(1 - \frac{\alpha}{2m}\right)} \sqrt{\frac{2}{n}} \simeq 0,69$ .

Além disso,  $\hat{\lambda}_1 \simeq 2,44$  e  $\hat{\lambda}_2 \simeq 1,41$ .

Assim, os IC's simultâneos são dados por

$$IC_{95\%}(\lambda_1) = \left( \frac{2,44}{1,31}; \frac{2,44}{0,69} \right) \simeq (1,86; 3,54)$$

$$IC_{95\%}(\lambda_2) = \left( \frac{1,41}{1,31}; \frac{1,41}{0,69} \right) \simeq (1,08; 2,04)$$

## Análise de Correlação Canônica - ACC

A ACC visa identificar e quantificar as associações entre dois conjuntos de variáveis.

A ACC foca na correlação entre uma combinação linear das variáveis de um conjunto e uma combinação linear das variáveis em outro conjunto.

Primeiro determina-se o par de combinações lineares tendo a maior correlação.

Depois, determina-se o par de combinações lineares tendo a maior correlação dentre todos os pares não correlacionados com o par inicial, e assim por diante.

Os pares de combinações lineares são chamados de **variáveis canônicas** e suas correlações de **correlações canônicas**.

Suponha dois grupos de variáveis:

$$\mathbf{X}^{(1)} (p \times 1) \text{ e } \mathbf{X}^{(2)} (q \times 1) \text{ com } p \leq q,$$

tal que

$$\begin{aligned} E(\mathbf{X}^{(1)}) &= \boldsymbol{\mu}^{(1)}, & E(\mathbf{X}^{(2)}) &= \boldsymbol{\mu}^{(2)}, \\ \text{Var}(\mathbf{X}^{(1)}) &= \boldsymbol{\Sigma}_{11}, & \text{Var}(\mathbf{X}^{(2)}) &= \boldsymbol{\Sigma}_{22}, \quad \text{e;} \\ \text{Cov}(\mathbf{X}^{(1)}, \mathbf{X}^{(2)}) &= \boldsymbol{\Sigma}_{12} = \boldsymbol{\Sigma}_{21}^{\top}. \end{aligned}$$

Então, podemos organizar as variáveis da seguinte forma:

$$\begin{aligned} \mathbf{X} &= \begin{bmatrix} \mathbf{X}^{(1)} \\ \mathbf{X}^{(2)} \end{bmatrix}, \\ E(\mathbf{X}) &= \boldsymbol{\mu} = \begin{bmatrix} \boldsymbol{\mu}^{(1)} \\ \boldsymbol{\mu}^{(2)} \end{bmatrix}, \quad \text{e} \\ \text{Cov}(\mathbf{X}) &= \boldsymbol{\Sigma} = \begin{bmatrix} \boldsymbol{\Sigma}_{11} & \boldsymbol{\Sigma}_{12} \\ \boldsymbol{\Sigma}_{21} & \boldsymbol{\Sigma}_{22} \end{bmatrix}. \end{aligned}$$

A ACC visa resumir as associações entre os conjuntos de variáveis  $\mathbf{X}^{(1)}$  e  $\mathbf{X}^{(2)}$  em termos de poucas covariâncias (correlações) escolhidas com cuidado.

Sejam  $U = \mathbf{a}^\top \mathbf{X}^{(1)}$  e  $V = \mathbf{b}^\top \mathbf{X}^{(2)}$  tal que

$$\text{Var}(U) = \mathbf{a}^\top \text{Cov}(\mathbf{X}^{(1)}) \mathbf{a} = \mathbf{a}^\top \boldsymbol{\Sigma}_{11} \mathbf{a},$$

$$\text{Var}(V) = \mathbf{b}^\top \text{Cov}(\mathbf{X}^{(2)}) \mathbf{b} = \mathbf{b}^\top \boldsymbol{\Sigma}_{22} \mathbf{b}, \quad \text{e}$$

$$\text{Cov}(U, V) = \mathbf{a}^\top \text{Cov}(\mathbf{X}^{(1)}, \mathbf{X}^{(2)}) \mathbf{b} = \mathbf{a}^\top \boldsymbol{\Sigma}_{12} \mathbf{b}.$$

O interesse está nos valores de  $\mathbf{a}$  e  $\mathbf{b}$  que maximizem

$$\text{Cor}(U, V) = \frac{\mathbf{a}^\top \boldsymbol{\Sigma}_{12} \mathbf{b}}{\sqrt{\mathbf{a}^\top \boldsymbol{\Sigma}_{11} \mathbf{a}} \sqrt{\mathbf{b}^\top \boldsymbol{\Sigma}_{22} \mathbf{b}}}.$$

## Definições

- O primeiro par de variáveis canônicas é o par de combinações lineares  $U_1$  e  $V_1$  tendo variâncias unitárias e que maximiza  $\text{Cor}(U, V)$ .
- O segundo par de variáveis canônicas é o par de combinações lineares  $U_2$  e  $V_2$  tendo variâncias unitárias e que maximiza  $\text{Cor}(U, V)$  dentre todas as escolhas não correlacionadas com  $U_1$  e  $V_1$ .
- O raciocínio prosegue.

Suponha que  $\Sigma$  tenha posto completo. Então,

$$\max_{\mathbf{a}, \mathbf{b}} \text{Cor}(U, V) = \rho_1^*$$

é dado por

$$U_1 = \mathbf{e}_1^\top \Sigma_{11}^{-1/2} \mathbf{X}^{(1)} \quad \text{e} \quad V_1 = \mathbf{f}_1^\top \Sigma_{22}^{-1/2} \mathbf{X}^{(2)}$$

em que  $\mathbf{a}_1^\top = \mathbf{e}_1^\top \Sigma_{11}^{-1/2}$  e  $\mathbf{b}_1^\top = \mathbf{f}_1^\top \Sigma_{22}^{-1/2}$ .

O  $k$ -ésimo par de variáveis canônicas,  $k = 1, 2, \dots, p$ ,

$$U_k = \mathbf{e}_k^\top \Sigma_{11}^{-1/2} \mathbf{X}^{(1)} \quad \text{e} \quad V_k = \mathbf{f}_k^\top \Sigma_{22}^{-1/2} \mathbf{X}^{(2)}$$

maximiza

$$\text{Cor}(U_k, V_k) = \rho_k^*.$$

Seja  $\mathbf{A}$  uma matriz simétrica e positiva definida. Então, pela decomposição do valor singular (SVD) temos que  $\mathbf{A} = \mathbf{P}\mathbf{\Lambda}\mathbf{P}^\top$ . Logo,  $\mathbf{A}^{1/2} = \mathbf{P}\mathbf{\Lambda}^{1/2}\mathbf{P}^\top$  chama-se raiz quadrada de  $\mathbf{A}$  e consequentemente  $\mathbf{A}^{-1/2} = \mathbf{P}\mathbf{\Lambda}^{-1/2}\mathbf{P}^\top$ .

Vale ressaltar que  $\rho_1^{*2} \geq \rho_2^{*2} \geq \dots \geq \rho_p^{*2}$  são os autovalores de  $\Sigma_{11}^{-1/2} \Sigma_{12} \Sigma_{22}^{-1} \Sigma_{21} \Sigma_{11}^{-1/2}$  e  $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_p$  os autovetores associados.

Tem-se que  $\rho_1^{*2} \geq \rho_2^{*2} \geq \dots \geq \rho_p^{*2}$  são também os  $p$  maiores autovalores de  $\Sigma_{22}^{-1/2} \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12} \Sigma_{22}^{-1/2}$  e  $\mathbf{f}_1, \mathbf{f}_2, \dots, \mathbf{f}_p$  os autovetores associados. Cada  $\mathbf{f}_i$  é proporcional a  $\Sigma_{22}^{-1/2} \Sigma_{21} \Sigma_{11}^{-1} \mathbf{e}_i$ .

Tem-se também que

$$\begin{aligned} \text{Var}(U_k) &= \text{Var}(V_k) = 1 \\ \text{Cov}(U_k, U_\ell) &= \text{Cor}(U_k, U_\ell) = 0, \quad k \neq \ell \\ \text{Cov}(V_k, V_\ell) &= \text{Cor}(V_k, V_\ell) = 0, \quad k \neq \ell \\ \text{Cov}(U_k, V_\ell) &= \text{Cor}(U_k, V_\ell) = 0, \quad k \neq \ell, \end{aligned}$$

para  $k, \ell = 1, \dots, p$ .

Se as variáveis  $\mathbf{X}^{(1)}$  e  $\mathbf{X}^{(2)}$  forem padronizadas em  $\mathbf{Z}^{(1)} = \mathbf{V}^{-1/2}(\mathbf{X}^{(1)} - \boldsymbol{\mu}^{(1)})$  e  $\mathbf{Z}^{(2)} = \mathbf{V}^{-1/2}(\mathbf{X}^{(2)} - \boldsymbol{\mu}^{(2)})$ , então a correlação canônica não se altera.

Além disso, se  $\mathbf{a}_k$  é o vetor de coeficientes para a  $k$ -ésima variável canônica, então  $\mathbf{a}_k^\top \mathbf{V}^{1/2}$  é o vetor de coeficientes da  $k$ -ésima variável canônica construída das variáveis padronizadas  $\mathbf{Z}^{(1)}$ .

## Exemplo

Suponha que  $\mathbf{Z}^{(1)} = (Z_1^{(1)}, Z_2^{(1)})^\top$  e  $\mathbf{Z}^{(2)} = (Z_1^{(2)}, Z_2^{(2)})^\top$  sejam variáveis padronizadas tal que  $\mathbf{Z} = (\mathbf{Z}^{(1)\top}, \mathbf{Z}^{(2)\top})^\top$ ,

$$\text{Cor}(\mathbf{Z}) = \left[ \begin{array}{c|c} \mathcal{R}_{11} & \mathcal{R}_{12} \\ \hline \mathcal{R}_{21} & \mathcal{R}_{22} \end{array} \right] = \left[ \begin{array}{cc|cc} 1,0 & 0,4 & 0,5 & 0,6 \\ 0,4 & 1,0 & 0,3 & 0,4 \\ \hline 0,5 & 0,3 & 1,0 & 0,2 \\ 0,6 & 0,4 & 0,2 & 1,0 \end{array} \right].$$

Então,

$$\mathcal{R}_{11}^{-1/2} = \begin{bmatrix} 1,0681 & -0,2229 \\ -0,2229 & 1,0681 \end{bmatrix}, \quad \mathcal{R}_{22}^{-1} = \begin{bmatrix} 1,0417 & -0,2083 \\ -0,2083 & 1,0417 \end{bmatrix} \quad \text{e}$$

$$\mathcal{R}_{11}^{-1/2} \mathcal{R}_{12} \mathcal{R}_{22}^{-1} \mathcal{R}_{21} \mathcal{R}_{11}^{-1/2} = \begin{bmatrix} 0,4370 & 0,2178 \\ 0,2178 & 0,1097 \end{bmatrix}.$$

Para  $\mathcal{R}_{11}^{-1/2}\mathcal{R}_{12}\mathcal{R}_{22}^{-1}\mathcal{R}_{21}\mathcal{R}_{11}^{-1/2}$ , tem-se que os autovalores são dados por  $\rho_1^{*2} = 0,54572$  e  $\rho_1^{*2} = 0,00091$ , e o primeiro autovetor por  $\mathbf{e}_1 = (0,8947; 0,4468)^\top$ .

Assim,  $\mathbf{a}_1 = \mathcal{R}_{11}^{-1/2}\mathbf{e}_1 = (0,8560; 0,2777)^\top$ .

Tem-se que  $\mathbf{f}_1 \propto \mathcal{R}_{22}^{-1/2}\mathcal{R}_{21}\mathcal{R}_{11}^{-1/2}\mathbf{e}_1$  e  $\mathbf{b}_1 = \mathcal{R}_{22}^{-1/2}\mathbf{f}_1$ . Logo,

$$\mathbf{b}_1 \propto \mathcal{R}_{22}^{-1}\mathcal{R}_{21}\mathbf{a}_1 = (0,4025; 0,5442)^\top.$$

Agora, devemos padronizar  $\mathbf{b}_1$  para que

$$\text{Var}(V_1) = \text{Var}(\mathbf{b}_1^\top \mathbf{Z}^{(2)}) = \mathbf{b}_1^\top \mathcal{R}_{22}\mathbf{b}_1 = 1.$$

Assim,

$$\begin{bmatrix} 0,4025 & 0,5442 \end{bmatrix} \begin{bmatrix} 1,0 & 0,2 \\ 0,2 & 1,0 \end{bmatrix} \begin{bmatrix} 0,4025 \\ 0,5442 \end{bmatrix} = 0,5457.$$

Usando  $\sqrt{0,5457} = 0,7387$ , tem-se que

$$\mathbf{b}_1 = \frac{1}{0,7387} \begin{bmatrix} 0,4025 \\ 0,5442 \end{bmatrix} = \begin{bmatrix} 0,5448 \\ 0,7366 \end{bmatrix}.$$

O primeiro par de variáveis canônicas é dado por

$$\begin{aligned}U_1 &= \mathbf{a}_1^\top \mathbf{Z}^{(1)} = 0,8560Z_1^{(1)} + 0,2777Z_2^{(1)} \\V_1 &= \mathbf{b}_1^\top \mathbf{Z}^{(2)} = 0,5448Z_1^{(2)} + 0,7366Z_2^{(2)}\end{aligned}$$

e a correlação canônica

$$\rho_1^* = \sqrt{\rho_1^{*2}} = \sqrt{0,54572} = 0,7387.$$

Esta é a maior correlação possível entre combinações linear dos conjuntos  $\mathbf{Z}^{(1)}$  e  $\mathbf{Z}^{(2)}$ .

A segunda correlação canônica,  $\rho_2^* = \sqrt{0,00092} = 0,03$ , é muito pequena (não tem informação relevante).

## Variáveis Canônicas Populacionais

- É melhor trabalhar com variáveis padronizadas para evitar problemas de interpretação das variáveis canônica.
- Variáveis padronizadas são livres de escala.
- As variáveis canônicas podem ser “identificadas” em termos de variáveis por assunto.
- O cálculo da correlação entre as variáveis canônicas e as originais auxiliam o trabalho de interpretação (mas é necessário cuidado devido à unidimensionalidade).

Sejam  $\mathbf{A}_{(p \times p)} = (\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_p)^\top$  e  $\mathbf{B}_{(q \times q)} = (\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_q)^\top$  tal que

$$\mathbf{U} = \mathbf{A}\mathbf{X}^{(1)} \quad \text{e} \quad \mathbf{V} = \mathbf{B}\mathbf{X}^{(2)}$$

com foco nas  $p$  primeiras variáveis canônicas de  $\mathbf{V}$ . Logo,

$$\text{Cov}(\mathbf{U}, \mathbf{X}^{(1)}) = \text{Cov}(\mathbf{A}\mathbf{X}^{(1)}, \mathbf{X}^{(1)}) = \mathbf{A}\Sigma_{11}$$

$$\text{Cor}(\mathbf{U}, \mathbf{X}^{(1)}) = \text{Cov}(\mathbf{U}, \mathbf{V}_{11}^{-1/2} \mathbf{X}^{(1)}) = \text{Cov}(\mathbf{A}\mathbf{X}^{(1)}, \mathbf{V}_{11}^{-1/2} \mathbf{X}^{(1)}) = \mathbf{A}\Sigma_{11} \mathbf{V}_{11}^{-1/2},$$

$$\text{Cor}(\mathbf{U}, \mathbf{X}^{(2)}) = \mathbf{A}\Sigma_{12} \mathbf{V}_{22}^{-1/2},$$

$$\text{Cor}(\mathbf{V}, \mathbf{X}^{(1)}) = \mathbf{B}\Sigma_{21} \mathbf{V}_{11}^{-1/2}, \quad \text{e}$$

$$\text{Cor}(\mathbf{V}, \mathbf{X}^{(2)}) = \mathbf{B}\Sigma_{22} \mathbf{V}_{22}^{-1/2}.$$

No caso de variáveis padronizadas  $\mathbf{Z}^{(1)}$  e  $\mathbf{Z}^{(2)}$ , tem-se que

$$\text{Cor}(\mathbf{U}, \mathbf{Z}^{(1)}) = \mathbf{A}_z \mathcal{R}_{11}, \quad \text{Cor}(\mathbf{U}, \mathbf{Z}^{(2)}) = \mathbf{A}_z \mathcal{R}_{12},$$

$$\text{Cor}(\mathbf{V}, \mathbf{Z}^{(1)}) = \mathbf{B}_z \mathcal{R}_{21} \quad \text{e} \quad \text{Cor}(\mathbf{V}, \mathbf{Z}^{(2)}) = \mathbf{B}_z \mathcal{R}_{22}.$$

Vale ressaltar que a padronização não afeta a correlação. Portanto,

$\text{Cor}(\mathbf{U}, \mathbf{X}^{(1)}) = \text{Cor}(\mathbf{U}, \mathbf{Z}^{(1)})$ , e assim por diante.

## Exemplo (Continuação)

Tem-se que

$$\mathcal{R}_{11} = \begin{bmatrix} 1,0 & 0,4 \\ 0,4 & 1,0 \end{bmatrix}, \quad \mathcal{R}_{22} = \begin{bmatrix} 1,0 & 0,2 \\ 0,2 & 1,0 \end{bmatrix} \quad \text{e} \quad \mathcal{R}_{12} = \begin{bmatrix} 0,5 & 0,6 \\ 0,3 & 0,4 \end{bmatrix}.$$

Com  $p = 1$ ,  $\mathbf{A}_Z = [0,8560 \quad 0,2777]$  e  $\mathbf{B}_Z = [0,5448 \quad 0,7366]$  tal que

$$\text{Cor}(U_1, \mathbf{Z}^{(1)}) = [0,8560 \quad 0,2777] \begin{bmatrix} 1,0 & 0,4 \\ 0,4 & 1,0 \end{bmatrix} = [0,9671 \quad 0,6201] \quad \text{e}$$

$$\text{Cor}(V_1, \mathbf{Z}^{(2)}) = [0,5448 \quad 0,7366] \begin{bmatrix} 1,0 & 0,2 \\ 0,2 & 1,0 \end{bmatrix} = [0,6921 \quad 0,8456].$$

- A primeira variável em  $\mathbf{Z}^{(1)}$  é mais relacionada a variável canônica  $U_1$ .
- A segunda variável em  $\mathbf{Z}^{(2)}$  é mais relacionada a variável canônica  $V_1$ .
- Em ambos os casos, as correlações reforçam a informação dada em  $\mathbf{A}_Z$  e  $\mathbf{B}_Z$ .
- Note que as correlações aumentam a importância de  $Z_2^{(1)}$  em  $\mathbf{Z}^{(1)}$ , e de  $Z_1^{(2)}$  em  $\mathbf{Z}^{(2)}$  porque ignoram a contribuição das variáveis restantes em cada grupo.

## Variáveis Canônicas Amostrais e Correlações Canônicas Amostrais

Sejam  $\mathbf{X}^{(1)}$  e  $\mathbf{X}^{(2)}$  matrizes de amostras aleatórias de dois grupos de variáveis tal que

$$\bar{\mathbf{x}} = \begin{bmatrix} \bar{\mathbf{x}}^{(1)} \\ \bar{\mathbf{x}}^{(2)} \end{bmatrix} \quad \text{e} \quad \mathbf{S} = \begin{bmatrix} \mathbf{S}_{11} & \mathbf{S}_{12} \\ \mathbf{S}_{21} & \mathbf{S}_{22} \end{bmatrix}.$$

As combinações lineares

$$\hat{U} = \hat{\mathbf{a}}^\top \mathbf{x}^{(1)} \quad \text{e} \quad \hat{V} = \hat{\mathbf{b}}^\top \mathbf{x}^{(2)}$$

têm correlação amostral dada por

$$\text{Cor}(\hat{U}, \hat{V}) = \frac{\hat{\mathbf{a}}^\top \mathbf{S}_{12} \hat{\mathbf{b}}}{\sqrt{\hat{\mathbf{a}}^\top \mathbf{S}_{11} \hat{\mathbf{a}}} \sqrt{\hat{\mathbf{b}}^\top \mathbf{S}_{22} \hat{\mathbf{b}}}}.$$

O primeiro par de variáveis canônicas amostrais é o par de combinações lineares  $\hat{U}$  e  $\hat{V}$  que maximiza a correlação acima.

Sejam  $\hat{\rho}_1^{*2} \geq \hat{\rho}_2^{*2} \geq \dots \geq \hat{\rho}_p^{*2}$  e  $\hat{\mathbf{e}}_1, \hat{\mathbf{e}}_2, \dots, \hat{\mathbf{e}}_p$ , os autovalores e autovetores, respectivamente, de  $\mathbf{S}_{11}^{-1/2} \mathbf{S}_{12} \mathbf{S}_{22}^{-1} \mathbf{S}_{21} \mathbf{S}_{11}^{-1/2}$  com  $p \leq q$ .

Sejam  $\hat{\mathbf{f}}_1, \hat{\mathbf{f}}_2, \dots, \hat{\mathbf{f}}_p$ , autovetores de  $\mathbf{S}_{22}^{-1/2} \mathbf{S}_{21} \mathbf{S}_{11}^{-1} \mathbf{S}_{12} \mathbf{S}_{22}^{-1/2}$  podendo ser calculado como

$$\hat{\mathbf{f}}_k = (1/\hat{\rho}_k^{*2}) \mathbf{S}_{22}^{-1/2} \mathbf{S}_{21} \mathbf{S}_{11}^{-1} \hat{\mathbf{e}}_k, \quad k = 1, 2, \dots, p.$$

Então, o  $k$ -ésimo par de variáveis canônicas é dado por

$$\hat{U}_k = \hat{\mathbf{e}}_k^\top \mathbf{S}_{11}^{-1/2} \mathbf{x}^{(1)} \quad \text{e} \quad \hat{V}_k = \hat{\mathbf{f}}_k^\top \mathbf{S}_{22}^{-1/2} \mathbf{x}^{(2)}$$

com a  $k$ -ésima correlação canônica amostral dada por

$$\text{Cor}(\hat{U}_k, \hat{V}_k) = \hat{\rho}_k^*.$$

As variáveis canônicas amostrais têm variâncias amostrais unitárias, e as correlações amostrais entre elas é zero:

$$s_{\hat{U}_k, \hat{U}_k} = s_{\hat{V}_k, \hat{V}_k} = 1 \quad \text{e} \quad s_{\hat{U}_k, \hat{U}_\ell} = s_{\hat{V}_k, \hat{V}_\ell} = s_{\hat{U}_k, \hat{V}_\ell} = 0.$$

Sejam  $\widehat{\mathbf{A}}_{(p \times p)} = (\widehat{\mathbf{a}}_1, \widehat{\mathbf{a}}_2, \dots, \widehat{\mathbf{a}}_p)^\top$  e  $\widehat{\mathbf{B}}_{(q \times q)} = (\widehat{\mathbf{b}}_1, \widehat{\mathbf{b}}_2, \dots, \widehat{\mathbf{b}}_q)^\top$  tal que

$$\widehat{\mathbf{U}} = \widehat{\mathbf{A}}\mathbf{x}^{(1)} \quad \text{e} \quad \widehat{\mathbf{V}} = \widehat{\mathbf{B}}\mathbf{x}^{(2)}.$$

$$\begin{aligned} \text{Cor}(\widehat{\mathbf{U}}, \mathbf{x}^{(1)}) &= \widehat{\mathbf{A}}\mathbf{S}_{11}\mathbf{D}_{11}^{-1/2}, \\ \text{Cor}(\widehat{\mathbf{U}}, \mathbf{x}^{(2)}) &= \widehat{\mathbf{A}}\mathbf{S}_{12}\mathbf{D}_{22}^{-1/2}, \\ \text{Cor}(\widehat{\mathbf{V}}, \mathbf{x}^{(1)}) &= \widehat{\mathbf{B}}\mathbf{S}_{21}\mathbf{D}_{11}^{-1/2}, \quad \text{e} \\ \text{Cor}(\widehat{\mathbf{V}}, \mathbf{x}^{(2)}) &= \widehat{\mathbf{B}}\mathbf{S}_{22}\mathbf{D}_{22}^{-1/2}, \end{aligned}$$

em que  $\mathbf{D}_{11} = \text{diag}(\mathbf{S}_{11})$  e  $\mathbf{D}_{22} = \text{diag}(\mathbf{S}_{22})$ .

No caso de variáveis padronizadas  $\mathbf{z}^{(1)}$  e  $\mathbf{z}^{(2)}$ , tem-se que

$$\widehat{\mathbf{U}} = \widehat{\mathbf{A}}_{\mathbf{z}}\mathbf{z}^{(1)} \quad \text{e} \quad \widehat{\mathbf{V}} = \widehat{\mathbf{B}}_{\mathbf{z}}\mathbf{z}^{(2)}$$

com  $\widehat{\mathbf{A}}_{\mathbf{z}} = \widehat{\mathbf{A}}\mathbf{D}_{11}^{1/2}$  e  $\widehat{\mathbf{B}}_{\mathbf{z}} = \widehat{\mathbf{B}}\mathbf{D}_{22}^{1/2}$ .

(Ver exemplo\_10.r)

Técnicas multivariadas que dizem respeito à “separação” de conjuntos distintos de objetos (ou observações) e à alocação de novos objetos (observações) a grupos previamente definidos.

Podemos enumerar os principais objetivos aqui como:

**Discriminação:** Descrever grafica e algebricamente os aspectos que diferenciam os grupos de objetos (observações). Determinar “discriminantes” entre grupos.

**Classificação:** Alocar objetos em classes previamente definidas. A ênfase aqui está na construção de uma regra que pode ser usada para designar de forma ótima um novo objeto às classes existentes.

**Exemplo 1: (Diagnóstico médico)** Suponha que dispõe-se de uma amostra de  $n$  fichas de pacientes para os quais foram registrados  $p$  sintomas que podem ser representados por um vetor  $\mathbf{x}$  e cujo diagnóstico foi uma entre  $k$  doenças possíveis. Um novo paciente apresenta vetor de sintomas  $\mathbf{x}_0$ . Como utilizar a informação amostral para diagnosticar a doença do novo paciente?

Uma função que separa objetos pode servir algumas vezes como “alocadora” e, reciprocamente, uma regra que aloca objetos pode sugerir um procedimento de discriminação.

Na prática, os dois principais objetivos se sobrepõem e a distinção entre separação e alocação fica obscurecida.

## Separação e Classificação para o Caso de Duas Populações

Sejam  $\pi_1$  e  $\pi_2$  as duas populações.

Os objetos são separados ou classificados com base em  $p$  medidas  $\mathbf{X}^T = (X_1, X_2, \dots, X_p)$ . Os valores observados  $\mathbf{x}$  diferem de alguma forma de uma população para outra.

Se  $\mathbf{x}$  vem da população  $\pi_1$  dizemos que a distribuição caracterizada pela densidade de probabilidade conjunta de  $\mathbf{X}$  é dada por  $f_1(\mathbf{x})$ , caso contrário, a densidade é dada por  $f_2(\mathbf{x})$ .

As regras de classificação costumam ser desenvolvidas a partir de amostras de aprendizagem: amostras para as quais a classificação de todos os elementos é conhecida.

Essencialmente, o conjunto de todos os resultados amostrais é dividido em duas regiões complementares,  $R_1$  e  $R_2$  tal que se uma nova observação cair em  $R_1$  ela será classificada em  $\pi_1$  e, caso contrário, será classificada em  $\pi_2$ .

Regras de classificação não fornecem métodos livres de erro. Muitas vezes não é clara a distinção entre as medidas observadas de cada população: os grupos podem se sobrepor de alguma forma. Logo, será possível classificar um objeto de  $\pi_2$  em  $\pi_1$  e vice-versa.

Um bom procedimento de classificação deve resultar numa taxa de erro de classificação pequena.

### Probabilidades a Priori

Pode ocorrer que uma população tenha verossimilhança maior do que a outra porque uma população é muito maior. A regra de classificação deve levar em conta estas “probabilidades” a priori de cada população.

Notação: Sejam  $p_j$ ,  $j = 1, 2$  tal que  $p_j > 0$ ,  $j = 1, 2$  e  $p_1 + p_2 = 1$  tais probabilidades.

## Custos de Classificação Incorreta

Também pode ocorrer que classificar um objeto de  $\pi_1$  em  $\pi_2$  represente um erro muito mais sério do que o recíproco. A regra de classificação deve levar em conta os custos de classificação incorreta.

Seja  $\Omega$  o espaço amostral - isto é, a coleção de todos os valores possíveis do vetor  $\mathbf{x}$ . Sejam  $R_1 \subset \Omega$  o conjunto de valores para o qual classificamos o objeto com sendo de  $\pi_1$  e,  $R_2 = \Omega \setminus R_1$  o conjunto de valores para o qual classificamos o objeto como sendo de  $\pi_2$ .

Se  $p = 2$ , podemos representar graficamente esta situação.

A probabilidade condicional  $\Pr(2|1)$  de classificar um objeto de  $\pi_1$  em  $\pi_2$  é

$$\Pr(\mathbf{X} \in R_2|\pi_1) = \int_{R_2} f_1(\mathbf{x})d\mathbf{x} = \Pr(2|1).$$

Similarmente, a probabilidade condicional  $\Pr(1|2)$  de classificar um objeto de  $\pi_2$  em  $\pi_1$  é

$$\Pr(\mathbf{X} \in R_1|\pi_2) = \int_{R_1} f_2(\mathbf{x})d\mathbf{x} = \Pr(1|2).$$

Desse modo, a probabilidade global de classificação incorreta pode ser obtida como a soma do produto das probabilidades condicionais por suas probabilidades a priori:

$$P(\text{classificar incorretamente em } \pi_1) = \Pr(\mathbf{X} \in R_1 | \pi_2) \Pr(\pi_2) = \Pr(1|2)p_2.$$

$$P(\text{classificar incorretamente em } \pi_2) = \Pr(\mathbf{X} \in R_2 | \pi_1) \Pr(\pi_1) = \Pr(2|1)p_1.$$

Os esquemas de classificação costumam ser avaliados em função de suas probabilidades de classificação incorreta, mas observe que estas probabilidades ignoram os custos de classificação incorreta.

Suponha a seguinte tabela de custos de classificação:

| População Real | Classificado em $\pi_1$ | Classificado em $\pi_2$ |
|----------------|-------------------------|-------------------------|
| $\pi_1$        | 0                       | $C(2 1)$                |
| $\pi_2$        | $C(1 2)$                | 0                       |

Para qualquer regra, o custo esperado de classificação incorreta (CECI) é dado por

$$\text{CECI} = C(2|1)\Pr(2|1)p_1 + C(1|2)\Pr(1|2)p_2.$$

Uma regra de classificação razoável deve ter um CECI tão pequeno quanto possível.

**Proposição 1:** As regiões  $R_1$  e  $R_2$  que minimizam o CECI são definidas pelos valores de  $\mathbf{x}$  para os quais valem:

$$R_1 : \frac{f_1(\mathbf{x})}{f_2(\mathbf{x})} \geq \frac{C(1|2) p_2}{C(2|1) p_1}.$$

(Razão de densidades)  $\geq$  (Razão de custos)  $\times$  (Razão de probabilidades a priori).

$$R_2 : \frac{f_1(\mathbf{x})}{f_2(\mathbf{x})} < \frac{C(1|2) p_2}{C(2|1) p_1}.$$

A implementação dessa regra requer o conhecimento da razão de densidades para uma nova observação  $\mathbf{x}_0$ , da razão de custos e da razão de probabilidades a priori. É, em geral, mais simples atribuir valores para as razões do lado direito da desigualdade acima do que atribuir um valor para cada probabilidade a priori e custo de classificação incorreta.

Casos especiais da regra estabelecida pela Proposição 1:

(1a) probabilidades a priori iguais,

$$R_1 : \frac{f_1(\mathbf{x})}{f_2(\mathbf{x})} \geq \frac{C(1|2)}{C(2|1)};$$

(1b) custos de classificação incorreta iguais,

$$R_1 : \frac{f_1(\mathbf{x})}{f_2(\mathbf{x})} \geq \frac{p_2}{p_1};$$

(1c) probabilidades a priori iguais e custos iguais,

$$R_1 : \frac{f_1(\mathbf{x})}{f_2(\mathbf{x})} \geq 1.$$

No caso especial (1c), observe que a regra reduz-se à comparação de densidades tal que se  $f_1(\mathbf{x}_0) \geq f_2(\mathbf{x}_0)$ ,  $\mathbf{x}_0$  é classificado em  $\pi_1$ . Caso, contrário,  $\mathbf{x}_0$  é classificado em  $\pi_2$ .

Outros critérios podem ser usados para obter uma regra de classificação ótima. Por exemplo, podemos ignorar os custos de classificação incorreta e escolher  $R_1$  e  $R_2$  que minimizam a probabilidade total de classificação incorreta (PTCI).

$$\begin{aligned}\text{PTCI} &= \Pr(\text{classificar incorretamente uma observação}) \\ &= p_1 \int_{R_2} f_1(\mathbf{x}) d\mathbf{x} + p_2 \int_{R_1} f_2(\mathbf{x}) d\mathbf{x}.\end{aligned}$$

Matematicamente, este problema é equivalente a minimizar o custo esperado de classificação incorreta quando os custos de classificação incorreta são iguais (1b).

Poderíamos também alocar uma nova observação  $\mathbf{x}_0$  a população com maior probabilidade a posteriori  $\Pr(\pi_i|\mathbf{x}_0)$ ,  $i = 1, 2$ .

Pelo teorema de Bayes

$$\begin{aligned}\Pr(\pi_1|\mathbf{x}_0) &= \frac{\Pr(\pi_1 \text{ ocorrer e observarmos } \mathbf{x}_0)}{\Pr(\text{observarmos } \mathbf{x}_0)} \\ &= \frac{\Pr(\mathbf{x}_0|\pi_1)p_1}{\Pr(\mathbf{x}_0|\pi_1)p_1 + \Pr(\mathbf{x}_0|\pi_2)p_2} \\ &= \frac{p_1 f_1(\mathbf{x}_0)}{p_1 f_1(\mathbf{x}_0) + p_2 f_2(\mathbf{x}_0)}, \quad \text{e} \\ \Pr(\pi_2|\mathbf{x}_0) &= \frac{p_2 f_2(\mathbf{x}_0)}{p_1 f_1(\mathbf{x}_0) + p_2 f_2(\mathbf{x}_0)}.\end{aligned}$$

Se

$$\Pr(\pi_1|\mathbf{x}_0) \geq \Pr(\pi_2|\mathbf{x}_0)$$

classificamos  $\mathbf{x}_0$  em  $\pi_1$ . Caso contrário, classificamos  $\mathbf{x}_0$  em  $\pi_2$ . Observe que essa regra é equivalente a regra (1b), que considera os custos de classificação incorreta iguais.

# Classificação em Uma de Duas Populações Normais Multivariadas

Primeiro, suponha que as populações tenham matrizes de covariâncias iguais,  $\Sigma_1 = \Sigma_2 = \Sigma$ .

$$f_j(\mathbf{x}) = (2\pi)^{-p/2} |\Sigma|^{-1/2} \exp \left\{ -\frac{1}{2} (\mathbf{x} - \mu_j)^\top \Sigma^{-1} (\mathbf{x} - \mu_j) \right\}, \quad j = 1, 2,$$

com  $\mu_1$ ,  $\mu_2$  e  $\Sigma$  desconhecidos.

**Proposição 2:** A regra do CECI mínimo é dada por

$$\left\{ \begin{array}{l} R_1 : \left\{ \mathbf{x} \mid (\mu_1 - \mu_2)^\top \Sigma^{-1} \mathbf{x} - \frac{1}{2} (\mu_1 - \mu_2)^\top \Sigma^{-1} (\mu_1 + \mu_2) \geq \ln \left( \frac{C(1|2)p_2}{C(2|1)p_1} \right) \right\}; \\ R_2 : \text{ caso contrário.} \end{array} \right.$$

**Observação:** Nas aplicações  $\mu_1$ ,  $\mu_2$  e  $\Sigma$  são desconhecidos, por essa razão, a regra dada pela proposição 2 deve ser modificada.

O estimador de  $\Sigma$  é dado por

$$\mathbf{S}_c = \frac{(n_1 - 1)\mathbf{S}_1 + (n_2 - 1)\mathbf{S}_2}{(n_1 + n_2 - 2)}.$$

A regra resultante da substituição pelos vetores de média amostral e a matriz  $\mathbf{S}_c$  é:

$$\left\{ \begin{array}{l} R_1 : \left\{ \mathbf{x} \mid (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} \mathbf{x} - \frac{1}{2}(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} (\bar{\mathbf{x}}_1 + \bar{\mathbf{x}}_2) \geq \ln \left( \frac{C(1|2)p_2}{C(2|1)p_1} \right) \right\} ; \\ R_2 : \text{ caso contrário.} \end{array} \right.$$

Se as probabilidades a priori são iguais e os custos de classificação incorreta também são iguais, então a regra acima se simplifica para

$$R_1 : \left\{ \mathbf{x} \mid (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} \mathbf{x} \geq \frac{1}{2}(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} (\bar{\mathbf{x}}_1 + \bar{\mathbf{x}}_2) \right\}.$$

Faça

$$\begin{aligned}\hat{y} &= \hat{\mathbf{a}}^\top \mathbf{x}, \quad \text{com} \\ \hat{\mathbf{a}}^\top &= (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1}, \\ \hat{m} &= \frac{1}{2}(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1}(\bar{\mathbf{x}}_1 + \bar{\mathbf{x}}_2) = \frac{1}{2}(\bar{y}_1 + \bar{y}_2) \quad \text{com} \\ \bar{y}_j &= \hat{\mathbf{a}}^\top \bar{\mathbf{x}}_j, \quad j = 1, 2.\end{aligned}$$

Resumindo, a regra estimada do CECI mínimo é equivalente a criar duas populações normais univariadas para os valores  $y$ , tomando-se uma combinação linear apropriada das observações de  $\pi_1$  e  $\pi_2$  e, então, designar  $\mathbf{x}_0$  a  $\pi_1$  ou a  $\pi_2$  dependendo se  $\hat{y}_0 = \hat{\mathbf{a}}^\top \mathbf{x}_0$  cai à direita ou à esquerda do ponto médio  $\hat{m}$  entre as duas médias amostrais  $\bar{y}_1$  e  $\bar{y}_2$ .

Como os parâmetros são substituídos por suas estimativas não se pode mais assegurar que a regra resultante minimize o custo esperado de classificação incorreta em uma particular aplicação.

Porém, parece razoável esperar que ela deva comportar-se bem para tamanhos amostrais grandes.

Resumindo: se os dados parecem ser normais multivariados, a estatística de classificação do lado esquerdo

$$(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} \mathbf{x} - \frac{1}{2}(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1}(\bar{\mathbf{x}}_1 + \bar{\mathbf{x}}_2)$$

pode ser calculada para cada nova observação  $\mathbf{x}_0$ . Essas observações são classificadas comparando-se os valores da estatística com o valor de

$$\ln \left( \frac{C(1|2)p_2}{C(2|1)p_1} \right).$$

**Exemplo 2:** Um biólogo obteve medidas sobre  $n = 25$  lagartos conhecidos cientificamente como *Cophosaurus texanus*. O peso (*mass*) é dados em gramas, enquanto que o comprimento da abertura do focinho (*svl*) e a extensão dos membros posteriores (*hls*) são dados em milímetros. Os dados estão disponíveis no arquivo T1-3.DAT. Além das três medidas, o biólogo identificou o gênero de cada lagarto *m*-macho, *f*-fêmea. Construir uma regra de classificação de gênero a partir das três medidas usando os dados disponíveis. (Ver exemplo\_11.r)

Utilizaremos a função `lda` do R disponível no pacote `mass`.

Tabela 3: Probabilidade a prior dos grupos.

| Grupo 1 | Grupo 2 |
|---------|---------|
| 0,48    | 0,52    |

Como nada foi dito na função do  $\mathbb{R}$ , o mesmo adota prioris iguais às proporções amostrais.

Tabela 4: Grupo de médias.

| Grupo | <i>mass</i> | <i>svl</i> | <i>hls</i> |
|-------|-------------|------------|------------|
| Fêmea | 7,012       | 63,042     | 118,000    |
| Macho | 10,233      | 73,346     | 139,769    |

Tabela 5: Coeficientes de discriminação linear.

| Variável    | $LD_1$  |
|-------------|---------|
| <i>mass</i> | -0,5723 |
| <i>svl</i>  | -0,0908 |
| <i>hls</i>  | 0,2949  |

Tabela 6: Erros de classificação.

| Grupo | Classificação |       |
|-------|---------------|-------|
|       | Fêmea         | Macho |
| Fêmea | 11            | 1     |
| Macho | 0             | 13    |

Observação: Especificando prioris iguais, a tabela acima não apresentará erros de classificação reaplicada a amostra de aprendizagem.

## Escala

Para qualquer constante  $c \neq 0$ , o vetor  $\hat{c\mathbf{a}} = c\mathbf{S}_c^{-1}(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)$  também servirá como coeficientes discriminantes. O vetor  $\hat{\mathbf{a}}$  é frequentemente “normalizado” para facilitar a interpretação de seus elementos. Duas das normalizações mais comuns são apresentadas a seguir.

1. Faça

$$\hat{\mathbf{a}}^* = \frac{\hat{\mathbf{a}}}{\sqrt{\hat{\mathbf{a}}^\top \hat{\mathbf{a}}}}$$

tal que  $\hat{\mathbf{a}}^*$  tenha comprimento unitário.

2. Faça

$$\hat{\mathbf{a}}^* = \frac{\hat{\mathbf{a}}}{\hat{a}_1}$$

tal que o primeiro elemento de  $\hat{\mathbf{a}}^*$  seja 1.

Em ambos os casos,  $\hat{\mathbf{a}}^*$  é da forma  $\hat{c\mathbf{a}}$ .

# Abordagem de Fisher para Classificação em Uma de Duas Populações

Fisher de fato chegou a estatística de classificação

$$R_1 : \left\{ \mathbf{x} \mid (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} \mathbf{x} \geq \frac{1}{2} (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} (\bar{\mathbf{x}}_1 + \bar{\mathbf{x}}_2) \right\} \quad \text{ou}$$

$$R_1 : \left\{ y \mid \hat{y} \geq (\bar{y}_1 + \bar{y}_2)/2 \right\}, \quad \left( \hat{y} = \mathbf{a}^\top \mathbf{x}, \quad \mathbf{a} = (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1}, \quad \bar{y}_j = \hat{\mathbf{a}}^\top \bar{\mathbf{x}}_j \right),$$

usando um argumento completamente diferente. A ideia de Fisher foi transformar as observações multivariadas  $\mathbf{x}$  em observações univariadas  $y$  tal que os  $y$ 's provenientes de  $\pi_1$  e de  $\pi_2$  sejam tão separados quanto possível.

Fisher sugeriu tomar combinações lineares de  $\mathbf{x}$  para criar os  $y$ 's, porque elas são funções simples de  $\mathbf{x}$  e podem ser manipuladas facilmente.

A abordagem de Fisher não assume que as populações sejam normais. No entanto, implicitamente, assume que as matrizes de covariâncias das populações sejam iguais, porque uma estimativa combinada da matriz de covariâncias é usada.

Uma combinação linear fixada de  $\mathbf{x}$  toma os valores  $y_{11}, y_{12}, \dots, y_{1n_1}$  para as observações de  $\pi_1$  e os valores  $y_{21}, y_{22}, \dots, y_{2n_2}$  para as observações de  $\pi_2$ .

A separação destes dois conjuntos de valores univariados é avaliada em função da diferença entre as médias amostrais  $\bar{y}_1$  e  $\bar{y}_2$  expressa em unidades de desvio padrão.

$$\text{Separação} = \frac{|\bar{y}_1 - \bar{y}_2|}{s_y} \quad \text{com}$$

$$s_y^2 = \frac{1}{n_1 + n_2 - 2} \left[ \sum_{i=1}^{n_1} (y_{1i} - \bar{y}_1)^2 + \sum_{i=1}^{n_2} (y_{2i} - \bar{y}_2)^2 \right].$$

O objetivo é seleccionar a combinação linear de  $\mathbf{x}$  que alcança a separação máxima das médias amostrais  $\bar{y}_1$  e  $\bar{y}_2$ .

### Proposição 3: A combinação linear

$$y = \hat{\mathbf{a}}^\top \mathbf{x} = (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} \mathbf{x} \quad \text{maximiza a razão}$$

$$\frac{\text{distância quadrada entre médias}}{\text{variância amostral de } y} = \frac{(\bar{y}_1 - \bar{y}_2)^2}{s_y^2} = \frac{(\hat{\mathbf{a}}^\top \mathbf{d})^2}{\hat{\mathbf{a}}^\top \mathbf{S}_c \hat{\mathbf{a}}}.$$

### Regra de Alocação: Função Discriminante (Linear) de Fisher

Sejam

$$\hat{y}_0 = (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} \mathbf{x}_0$$

$$\hat{m} = \frac{1}{2}(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1}(\bar{\mathbf{x}}_1 + \bar{\mathbf{x}}_2) = \frac{1}{2}(\bar{y}_1 + \bar{y}_2).$$

Aloque  $\mathbf{x}_0$  a  $\pi_1$  se  $\hat{y}_0 \geq \hat{m}$ .

Caso contrário, aloque  $\mathbf{x}_0$  a  $\pi_2$ .

## Classificação É Uma Boa Ideia?

Para duas populações, a separação máxima relativa que pode ser obtida considerando-se combinações lineares das observações multivariadas é igual a distância

$$D^2 = (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2).$$

Isto é conveniente, porque  $D^2$  pode ser usada, em certas situações, para testar se as médias das populações  $\pi_1$  e  $\pi_2$  diferem significativamente. Consequentemente, um teste de diferença entre vetores de média pode ser visto como um teste para a “significância” da separação que pode ser alcançada.

Suponha que as populações  $\pi_1$  e  $\pi_2$  sejam normais multivariadas com uma matriz de covariâncias comum  $\Sigma$ . Então, vimos que um teste de  $H_0 : \mu_1 = \mu_2$  versus  $H_1 : \mu_1 \neq \mu_2$  usa a estatística

$$\left( \frac{n_1 + n_2 - p - 1}{(n_1 + n_2 - 2)p} \right) \left( \frac{n_1 n_2}{n_1 + n_2} \right) D^2,$$

que sob  $H_0$  tem distribuição  $\mathcal{F}_{p, n_1 + n_2 - p - 1}$ .

Se  $H_0$  é rejeitada, podemos concluir que a separação entre as duas populações é significativa.

**Observação:** Separação significativa não necessariamente implicará em boa classificação. A eficácia de um procedimento de classificação pode ser avaliada independentemente de qualquer teste de separação. Em contraste, se a separação não é significativa, a busca por uma regra de classificação útil será, provavelmente, infrutífera.

## Classificação de Populações Normais - Caso $\Sigma_1 \neq \Sigma_2$

As regras de classificação são mais complicadas quando as matrizes de covariâncias das populações são desiguais. Considere novamente a razão das densidades normais multivariadas, agora considerando as covariâncias desiguais. Neste caso, os fatores fora do termo exponencial não simplificam e não é possível colocar o termo dentro da exponencial em evidência.

$$\frac{f_1(\mathbf{x})}{f_2(\mathbf{x})} = \left( \frac{|\Sigma_2|}{|\Sigma_1|} \right)^{1/2} \exp \left\{ -\frac{1}{2} \left[ (\mathbf{x} - \mu_1)^\top \Sigma_1^{-1} (\mathbf{x} - \mu_1) + (\mathbf{x} - \mu_2)^\top \Sigma_2^{-1} (\mathbf{x} - \mu_2) \right] \right\}.$$

Nesse caso, as regiões de classificação, segundo o critério do custo esperado de classificação incorreta mínimo, serão dadas por (na escala logaritmo natural):

$$\begin{cases} R_1 : & R_1 : -\frac{1}{2} \mathbf{x}^\top (\Sigma_1^{-1} - \Sigma_2^{-1}) \mathbf{x} + (\mu_1^\top \Sigma_1^{-1} - \mu_2^\top \Sigma_2^{-1}) \mathbf{x} - k \geq \ln \left( \frac{C(1|2) p_2}{C(2|1) p_1} \right); \\ R_2 : & \text{caso contrário,} \end{cases}$$

com

$$k = \frac{1}{2} \ln \left( \frac{|\Sigma_2|}{|\Sigma_1|} \right) + \frac{1}{2} (\mu_1^\top \Sigma_1^{-1} \mu_1 - \mu_2^\top \Sigma_2^{-1} \mu_2).$$

As regiões de classificação são quadráticas em  $\mathbf{x}$ . Quando  $\Sigma_1 = \Sigma_2$ , o termo quadrático  $\mathbf{x}^\top (\Sigma_1^{-1} - \Sigma_2^{-1}) \mathbf{x}$  se anula, e as regiões resultantes são aquelas obtidas anteriormente no caso de variâncias iguais.

**Proposição 4:** Sob normalidade multivariada com covariâncias desiguais, aloque  $\mathbf{x}_0$  a  $\pi_1$  se

$$-\frac{1}{2} \mathbf{x}_0^\top (\Sigma_1^{-1} - \Sigma_2^{-1}) \mathbf{x}_0 + (\mu_1^\top \Sigma_1^{-1} - \mu_2^\top \Sigma_2^{-1}) \mathbf{x}_0 - k \geq \ln \left( \frac{C(1|2) p_2}{C(2|1) p_1} \right).$$

Caso contrário, aloque  $\mathbf{x}_0$  a  $\pi_2$ .

Na prática, a regra de classificação acima é implementada substituindo-se os parâmetros populacionais por estimativas  $\bar{\mathbf{x}}_1$ ,  $\bar{\mathbf{x}}_2$  e  $\mathbf{S}_1$  e  $\mathbf{S}_2$ .

## Regra de Classificação Quadrática

Populações normais, covariâncias desiguais:

Aloque  $\mathbf{x}_0$  a  $\pi_1$  se

$$-\frac{1}{2} \mathbf{x}_0^\top (\mathbf{S}_1^{-1} - \mathbf{S}_2^{-1}) \mathbf{x}_0 + (\bar{\mathbf{x}}_1^\top \mathbf{S}_1^{-1} - \bar{\mathbf{x}}_2^\top \mathbf{S}_2^{-1}) \mathbf{x}_0 - k \geq \ln \left( \frac{C(1|2) p_2}{C(2|1) p_1} \right).$$

Caso contrário, aloque  $\mathbf{x}_0$  a  $\pi_2$ .

Classificação com funções quadráticas é bem complicada quando se tem mais de duas medidas e pode levar a resultados estranhos. Isto é particularmente verdade quando os dados não são (essencialmente) normais multivariados.

As regiões de classificação podem ser uma união de regiões disjuntas do espaço amostral.

Em muitas aplicações, a cauda inferior da distribuição de  $\pi_1$  será menor do que a prescrita por uma distribuição normal e a regra quadrática poderá levar a altas taxas de erro de classificação. Uma desvantagem séria da regra quadrática é que ela é bem sensível a desvios da normalidade.

Se os dados não são normais multivariados, duas opções estão disponíveis. A primeira, envolve transformar os dados não normais, e depois testar a igualdade das matrizes de covariâncias para verificar se é a regra linear ou a quadrática que devem ser usadas.

Os testes usuais para homogeneidade das covariâncias são fortemente afetados sob não normalidade. A conversão de dados não normais para dados normais deve sempre ser feita antes de realizar tais testes.

Como segunda opção, podemos usar uma regra linear (ou quadrática) sem nos preocuparmos com a forma das distribuições populacionais e esperar que elas irão funcionar razoavelmente bem.

Estudos mostraram, porém, que existem casos não normais para os quais uma função de classificação linear tem uma performance ruim, mesmo se as matrizes de covariâncias das duas populações são iguais.

Moral da história: sempre verificar a performance de qualquer procedimento de classificação. Em último caso, isto deve ser feito com o conjunto de dados usado para construir a regra. O ideal é que se tenha uma quantidade de dados suficientemente grande que podem ser repartidos em amostras de treinamento e de validação. A amostra de treinamento/aprendizagem é usada para construir a regra, e a amostra de validação é usada para avaliar a performance da regra construída.

## Avaliação das Funções de Classificação

A avaliação envolve calcular taxas de erro ou probabilidades de classificação incorreta.

Como as densidades são em geral desconhecidas, concentraremos-nos sobre as taxas de erro associadas à função de classificação amostral.

Taxa de Erro Ótima (TEO) - regra de classificação segundo o critério da probabilidade total de classificação incorreta (PTCI) mínima.

$$\text{TEO} = p_1 \int_{R_2} f_1(\mathbf{x}) d\mathbf{x} + p_2 \int_{R_1} f_2(\mathbf{x}) d\mathbf{x}$$

**Exemplo 3:** Suponha duas populações normais multivariadas com matrizes de covariâncias iguais,  $p_1 = p_2 = 1/2$  e também  $C(2|1) = C(1|2)$  tal que  $\ln \left( \frac{C(1|2)}{C(2|1)} \frac{p_2}{p_1} \right) = 0$ .

Neste caso,

$$R_1 : \left\{ \mathbf{x} \mid (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)^\top \boldsymbol{\Sigma}^{-1} \mathbf{x} \geq \frac{1}{2} (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)^\top \boldsymbol{\Sigma}^{-1} (\boldsymbol{\mu}_1 + \boldsymbol{\mu}_2) \right\} \quad \text{ou}$$

$$R_1 : \left\{ \mathbf{x} \mid \mathbf{a}^\top \mathbf{x} \geq \frac{1}{2} \mathbf{a}^\top (\boldsymbol{\mu}_1 + \boldsymbol{\mu}_2) \right\}.$$

Fazendo  $Y = \mathbf{a}^\top \mathbf{X}$  teremos  $\sigma_Y^2 = \mathbf{a}^\top \boldsymbol{\Sigma} \mathbf{a} = \delta^2$ .

$$PTCI = \frac{1}{2} \int_{R_2} f_1(\mathbf{x}) d\mathbf{x} + \frac{1}{2} \int_{R_1} f_2(\mathbf{x}) d\mathbf{x} = \Phi \left( -\frac{\delta}{2} \right).$$

Se  $\delta^2 = 2,56$ , teremos  $PTCI = \Phi(-0,8) = 0,2119$ .

A regra de classificação aqui irá alocar cerca de 21% dos itens incorretamente.

Este exemplo ilustra como a TEO pode ser calculada quando as funções de densidade são conhecidas. Como em geral os parâmetros populacionais são desconhecidos, eles deverão ser estimados e a avaliação da taxa de erro não será tão direta.

A performance da função de classificação amostral pode, em princípio, ser avaliada calculando-se a taxa de erro real (TER).

$$\text{TER} = p_1 \int_{\hat{R}_2} f_1(\mathbf{x}) d\mathbf{x} + p_2 \int_{\hat{R}_1} f_2(\mathbf{x}) d\mathbf{x}.$$

$\hat{R}_1$  e  $\hat{R}_2$  são as regiões de classificação determinadas pelas amostras de tamanhos  $n_1$  e  $n_2$ , respectivamente.

A TER indica como a função de classificação amostral se comportará em amostras futuras. Como a TEO, geralmente ela não poderá ser calculada, pois depende das densidades  $f_1$  e  $f_2$ . Porém, uma estimativa de uma quantidade relacionada a TER pode ser calculada e será apresentada aqui.

A taxa de erro real aparente (TERA) pode ser calculada a partir da matriz de “confusão” (tabela de dupla entrada indicando as frequências de classificações corretas e incorretas em cada população (grupo)).

| População | Classificação em        |                         |
|-----------|-------------------------|-------------------------|
|           | $\pi_1$                 | $\pi_2$                 |
| $\pi_1$   | $n_{1c}$                | $n_{1M} = n_1 - n_{1c}$ |
| $\pi_2$   | $n_{2M} = n_2 - n_{2c}$ | $n_{2c}$                |

$$\text{TERA} = \frac{n_{1M} + n_{2M}}{n_1 + n_2}.$$

Observe que a TERA nada mais é do que a proporção amostral de classificações incorretas considerando-se a amostra de treinamento.

A TERA é uma medida intuitiva e simples, mas tem um viés: tende a subestimar a TER, a menos que  $n_1$  e  $n_2$  sejam suficientemente grandes.

Estimativas de taxas de erro melhores do que a TERA e que não exigem a suposição das distribuições populacionais podem ser construídas.

Um procedimento é dividir a amostra total em uma amostra de treinamento e outra de validação. A amostra de treinamento é usada para construir a função de classificação e, a de validação, para avaliar a função obtida. A taxa de erro é determinada pela proporção amostral de classificações incorretas na amostra de validação.

Apesar deste método superar o problema do viés, ele padece de dois defeitos:

1. requer amostras muito grandes;
2. a função avaliada não é a função de interesse. Em última análise, quase todos os dados devem ser usados para construir a regra. Caso contrário, informação importante pode estar sendo desperdiçada.

Uma segunda abordagem, que parece funcionar bem, é chamada procedimento de validação “reter um fora” (*holdout*) de Lachenbruch.

1. Comece em  $\pi_1$ . Omita uma de suas observações e desenvolva a função de classificação com as restantes  $n_1 - 1 + n_2$ .
2. Classifique a observação omitida com a função obtida.
3. Repita os passos (1) e (2) até que todas as observações de  $\pi_1$  sejam classificadas. Defina  $n_{1M}^{(H)}$  como o número de classificações incorretas neste grupo.
4. Repita os passos (1), (2) e (3) para as observações de  $\pi_2$  e defina  $n_{2M}^{(H)}$  como o número de classificações incorretas neste grupo.

$$\widehat{\Pr(2|1)} = \frac{n_{1M}^{(H)}}{n_1}, \quad \widehat{\Pr(1|2)} = \frac{n_{2M}^{(H)}}{n_2} \quad \text{e} \quad \widehat{PTCI} = \frac{n_{1M}^{(H)} + n_{2M}^{(H)}}{n_1 + n_2}.$$

Para amostras moderadas  $\widehat{PTCI}$  é uma estimativa não viesada do valor esperado da TERA (taxa de erro aparente).

Deve ser intuitivamente claro que classificação boa (taxas de erro pequenas) dependerá da separação dos grupos. O mais separados são os grupos, mais provavelmente uma regra de classificação útil será desenvolvida.

Como veremos, regras de alocação apropriadas para o caso envolvendo probabilidades a priori iguais e custos de classificação incorreta iguais correspondem às funções designadas para populações separadas o máximo possível. É nesta situação que começamos a perder a distinção entre classificação e separação.

## Classificação em Uma de $g$ Populações ( $g > 2$ )

Pelo menos em teoria, a extensão para a classificação em um de  $g$  grupos,  $g > 2$  é imediata. Porém, não muito é conhecido sobre as propriedades das funções de classificação amostrais correspondentes, e em particular, sobre suas taxas de erro investigadas.

A “robustez” da estatística linear de classificação em dois grupos para, por exemplo, covariâncias desiguais ou distribuições não normais pode ser estudada a partir de experimentos simulados. Para mais de duas populações, esta abordagem não leva a conclusões gerais, porque as propriedades dependem sobre onde as populações estão localizadas, e existem muitas configurações para serem convenientemente estudadas.

Como antes, a abordagem aqui será desenvolver regras ótimas teóricas e, então indicar as modificações exigidas para as aplicações reais.

## Regra do Custo Esperado de Classificação Incorreta Mínimo

Notação:

- $f_k(\mathbf{x})$  - função de densidade de probabilidade conjunta para o  $k$ -ésimo grupo,  $k = 1, 2, \dots, g$ .
- $p_1, p_2, \dots, p_g$  - probabilidades a priori de cada grupo tais que  $p_k > 0, \forall k$  e  $\sum_{k=1}^g p_k = 1$ .
- $C(k|j)$  - custo de classificação incorreta de uma observação de  $\pi_j$  em  $\pi_k, \forall j, k = 1, 2, \dots, g$  e  $j \neq k$ . Se  $j = k$ , então  $c(k|k) = 0$ .
- $R_k$  - região de classificação em  $\pi_k$  tal que  $\cup_{k=1}^g R_k = \Omega, R_j \cap R_k = \emptyset$  para  $j \neq k$ .

A probabilidade de classificar uma observação de  $\pi_j$  em  $\pi_k$  é

$$\Pr(k|j) = \int_{R_k} f_j(\mathbf{x}) d\mathbf{x} \quad \text{para } k \in \{1, 2, \dots, g\}, \quad k \neq j \quad \text{e}$$

$$\Pr(j|j) = 1 - \sum_{k=1, k \neq j}^g \Pr(k|j).$$

O custo esperado de classificação incorreta de uma observação proveniente de  $\pi_1$  será dado por

$$\begin{aligned}\text{CECI}(1) &= \Pr(2|1)C(2|1) + \Pr(3|1)C(3|1) + \cdots + \Pr(g|1)C(g|1) \\ &= \sum_{k=2}^g P(k|1)c(k|1).\end{aligned}$$

Este custo esperado condicional ocorre com probabilidade  $p_1$ , a probabilidade a priori de  $\pi_1$ .

De maneira similar, podemos obter os custos esperados de classificação incorreta condicionais  $\text{CECI}(2)$ ,  $\text{CECI}(3)$ ,  $\dots$ ,  $\text{CECI}(g)$ .

Multiplicando os custos condicionais pelas respectivas probabilidades a priori temos o custo esperado de classificação incorreta dado por

$$\text{CECI} = \sum_{j=1}^g p_j \left( \sum_{k=1, k \neq j}^g \Pr(k|j)C(k|j) \right).$$

**Proposição 5:** As regiões de classificação que minimizam o custo esperado de classificação incorreta são definidas por

- Aloque  $\mathbf{x}$  a  $\pi_j$ ,  $j = 1, 2, \dots, g$  na qual

$$\sum_{j=1, j \neq k}^g p_j f_j(\mathbf{x}) C(k|j) \quad \text{é um mínimo.}$$

- Se os custos de classificação incorreta são todos iguais a unidade, observe que a regra alocará  $\mathbf{x}$  à população  $\pi_k$ ,  $k = 1, 2, \dots, g$  para a qual,

$$\sum_{j=1, j \neq k}^g p_j f_j(\mathbf{x}) \quad \text{é um mínimo.}$$

Observe que esta soma será um mínimo se o termo deixado de fora,  $p_k f_k(\mathbf{x})$ , for um máximo.

## Regra do CECI Mínimo para Custos de Classificação Incorreta Iguais

- Aloque  $\mathbf{x}_0$  à população  $\pi_k$  se

$$p_k f_k(\mathbf{x}_0) > p_j f_j(\mathbf{x}_0), \quad \forall j \neq k,$$

ou equivalentemente,

- Aloque  $\mathbf{x}_0$  à população  $\pi_k$  se

$$\ln(p_k f_k(\mathbf{x}_0)) > \ln(p_j f_j(\mathbf{x}_0)), \quad \forall j \neq k.$$

Esta regra é equivalente à regra que maximiza a probabilidade a posteriori  $\Pr(\pi_k|\mathbf{x}_0)$ .

Deve-se ter em mente que as regras do CECI mínimo têm três componentes: probabilidades a priori, custos de classificação incorreta e funções de densidade. Estes componentes devem ser especificados (ou estimados) antes da regra poder ser implementada.

#### Exemplo 4: (Classificação de nova observação em uma de três populações conhecidas)

Suponha os seguintes custos de classificação incorreta, probabilidades a priori e densidades avaliadas em  $\mathbf{x}_0$  uma nova observação.

| População           | Classificação em           |                            |                         |
|---------------------|----------------------------|----------------------------|-------------------------|
|                     | $\pi_1$                    | $\pi_2$                    | $\pi_3$                 |
| $\pi_1$             | $C(1 1) = 0$               | $C(2 1) = 10$              | $C(3 1) = 50$           |
| $\pi_2$             | $C(1 2) = 500$             | $C(2 2) = 0$               | $C(3 2) = 200$          |
| $\pi_3$             | $C(1 3) = 100$             | $C(2 3) = 50$              | $C(3 3) = 0$            |
| Prioris             | $p_1 = 0,05$               | $p_2 = 0,60$               | $p_3 = 0,35$            |
| $f_j(\mathbf{x}_0)$ | $f_1(\mathbf{x}_0) = 0,01$ | $f_2(\mathbf{x}_0) = 0,85$ | $f_3(\mathbf{x}_0) = 2$ |

Classificar  $\mathbf{x}_0$  em uma das três populações.

Usando a regra do CECI mínimo, alocaremos  $\mathbf{x}_0$  a  $\pi_k$ ,  $k = 1, 2, 3$  para a qual

$$\sum_{j=1, j \neq k}^3 p_j f_j(\mathbf{x}) C(k|j) \quad \text{é um mínimo.}$$

| $k$      | $\sum_{j=1, j \neq k}^3 p_j f_j(\mathbf{x}) C(k j)$ |
|----------|-----------------------------------------------------|
| 1        | 325                                                 |
| <b>2</b> | <b>35,055</b>                                       |
| 3        | 102,025                                             |

Como o menor valor ocorre para  $k = 2$ , alocamos  $\mathbf{x}_0$  a  $\pi_2$ .

Se os custos de classificação incorreta fossem todos iguais, designaríamos  $\mathbf{x}_0$  a  $\pi_k$ ,  $k = 1, 2, 3$  na qual  $p_k f_k(\mathbf{x}_0) > p_j f_j(\mathbf{x}_0)$ ,  $\forall j \neq k$ .

| $k$      | $p_k f_k(\mathbf{x}_0)$ |
|----------|-------------------------|
| 1        | 0,0005                  |
| 2        | 0,5100                  |
| <b>3</b> | <b>0,7000</b>           |

Nesse caso,  $\mathbf{x}_0$  seria alocada a  $\pi_3$ .

## Classificação com Populações Normais

$$f_k(\mathbf{x}) = (2\pi)^{-p/2} |\Sigma_k|^{-1/2} \exp\left\{-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu}_k)^\top \Sigma_k^{-1}(\mathbf{x} - \boldsymbol{\mu}_k)\right\}, \quad k = 1, 2, \dots, g.$$

Se considerarmos todos os custos iguais a unidade, a regra resultante será:

- Aloque  $\mathbf{x}_0$  a  $\pi_k$  se

$$\begin{aligned} \ln(p_k f_k(\mathbf{x}_0)) &= \ln p_k - \frac{p}{2} \ln(2\pi) - \frac{1}{2} \ln |\Sigma_k| - \frac{1}{2}(\mathbf{x} - \boldsymbol{\mu}_k)^\top \Sigma_k^{-1}(\mathbf{x} - \boldsymbol{\mu}_k) \\ &= \max_{1 \leq j \leq g} \ln(p_j f_j(\mathbf{x}_0)) \end{aligned}$$

A constante  $p \ln(2\pi)/2$  pode ser desprezada, pois ela é igual para todas as populações. Portanto, podemos definir um escore discriminante quadrático para a  $k$ -ésima população dado por

$$d_k^Q(\mathbf{x}) = \ln p_k - \frac{1}{2} \ln |\Sigma_k| - \frac{1}{2}(\mathbf{x} - \boldsymbol{\mu}_k)^\top \Sigma_k^{-1}(\mathbf{x} - \boldsymbol{\mu}_k), \quad k = 1, 2, \dots, g.$$

O escore quadrático  $d_k^Q(\mathbf{x})$  é composto pelas contribuições da variância generalizada  $|\Sigma_k|$ , da probabilidade a priori  $p_k$ , e da distância quadrada de  $\mathbf{x}$  à média da população  $\pi_k$ .

## Regra da Probabilidade Total de Classificação Incorreta Mínima - Populações Normais, Covariâncias Desiguais

- Aloque  $\mathbf{x}_0$  a  $\pi_k$  se o escore quadrático

$$d_k^Q(\mathbf{x}_0) = \max_{1 \leq j \leq g} \{d_j^Q(\mathbf{x}_0)\}.$$

Na prática  $\mu_k$  e  $\Sigma_k$  são desconhecidas para todo  $k = 1, 2, \dots, g$ , mas um conjunto de treinamento cujas classificações corretas das observações são conhecidas está em geral disponível para a construção de estimativas. As quantidades amostrais relevantes para a população  $\pi_k$  são  $\bar{\mathbf{x}}_k$  e  $\mathbf{S}_k$  com  $n_k$  o número de observações da  $k$ -ésima população.

$$\hat{d}_k^Q(\mathbf{x}) = \ln p_k - \frac{1}{2} \ln |\mathbf{S}_k| - \frac{1}{2} (\mathbf{x} - \bar{\mathbf{x}}_k)^\top \mathbf{S}_k^{-1} (\mathbf{x} - \bar{\mathbf{x}}_k), \quad k = 1, 2, \dots, g.$$

Assim, a regra estimada é alocar  $\mathbf{x}_0$  a  $\pi_k$  se o escore quadrático estimado

$$\hat{d}_k^Q(\mathbf{x}_0) = \max_{1 \leq j \leq g} \{\hat{d}_j^Q(\mathbf{x}_0)\}.$$

Uma simplificação aqui é possível para o caso em que

$$\mathbf{\Sigma}_1 = \mathbf{\Sigma}_2 = \cdots = \mathbf{\Sigma}_g = \mathbf{\Sigma}.$$

Neste caso os escores discriminantes passam a ser lineares em  $\mathbf{x}$  e simplificam para

$$d_k(\mathbf{x}) = \ln p_k + \boldsymbol{\mu}_k^\top \mathbf{\Sigma}^{-1} \mathbf{x} - \frac{1}{2} \boldsymbol{\mu}_k^\top \mathbf{\Sigma}^{-1} \boldsymbol{\mu}_k, \quad k = 1, 2, \dots, g.$$

Uma estimativa de  $d_k(\mathbf{x})$  é baseada em

$$\mathbf{s}_c = \frac{1}{n_1 + n_2 + \cdots + n_g - g} \sum_{k=1}^g (n_k - 1) \mathbf{s}_k$$

e é dada por

$$\hat{d}_k(\mathbf{x}) = \ln p_k + \bar{\mathbf{x}}_k^\top \mathbf{s}_c^{-1} \mathbf{x} - \frac{1}{2} \bar{\mathbf{x}}_k^\top \mathbf{s}_c^{-1} \bar{\mathbf{x}}_k, \quad k = 1, 2, \dots, g.$$

Consequentemente, temos a regra estimada dada por “aloque  $\mathbf{x}_0$  a  $\pi_k$  se

$$\hat{d}_k(\mathbf{x}_0) = \max_{1 \leq j \leq g} \{\hat{d}_j(\mathbf{x}_0)\}”.$$

## Método de Discriminação de Fisher para Várias Populações

O propósito principal na análise discriminante de Fisher (ADF) é separar populações (grupos). No entanto, como veremos, o produto final pode levar a uma regra de classificação. Na ADF não é necessário supor a normalidade das  $g$  populações, embora supõe-se que as covariâncias sejam iguais.

Denote  $\bar{\mu}$  o vetor de médias combinado das  $g$  populações, isto é,

$$\bar{\mu} = \frac{1}{g} \sum_{k=1}^g \mu_k.$$

Denote a matriz  $\mathbf{B}_\mu$  a matriz  $p \times p$  de somas de quadrados e produtos cruzados

$$\mathbf{B}_\mu = \sum_{k=1}^g (\mu_k - \bar{\mu})(\mu_k - \bar{\mu})^\top.$$

Faça  $Y = \mathbf{a}^\top \mathbf{X}$  tal que  $E(Y|\pi_k) = \mathbf{a}^\top \mu_k$  e  $\text{Var}(Y|\pi_k) = \mathbf{a}^\top \Sigma \mathbf{a}$ ,  $k = 1, 2, \dots, g$ .

Consequentemente  $\mu_{kY} = \mathbf{a}^\top \mu_k$  depende da população na qual  $\mathbf{X}$  foi observada.

Média Global:  $\bar{\mu}_Y = \mathbf{a}^\top \frac{1}{g} \sum_{k=1}^g \mu_k = \mathbf{a}^\top \bar{\mu}$ .

Medida de separação dos grupos:

$$\begin{aligned}\frac{\sum_{k=1}^g (\mu_{kY} - \bar{\mu}_Y)^2}{\sigma_Y^2} &= \frac{\sum_{k=1}^g (\mathbf{a}^\top \boldsymbol{\mu}_k - \mathbf{a}^\top \bar{\boldsymbol{\mu}})^2}{\mathbf{a}^\top \boldsymbol{\Sigma} \mathbf{a}} \\ &= \frac{\mathbf{a}^\top \mathbf{B}_\mu \mathbf{a}}{\mathbf{a}^\top \boldsymbol{\Sigma} \mathbf{a}}.\end{aligned}$$

Esta razão mede a variabilidade entre os grupos de valores de  $Y$  relativa à variabilidade comum dentro dos grupos.

Podemos seleccionar  $\mathbf{a}$  que maximiza esta razão.

Em geral  $\boldsymbol{\Sigma}$  e  $\boldsymbol{\mu}_k$  são desconhecidos, mas dispõe-se de uma amostra de treinamento a partir da qual podemos estimar estas quantidades.

Sejam  $\bar{\mathbf{x}}_k$ ,  $\mathbf{S}_k$  as estimativas em cada grupo e

$$\begin{aligned}\bar{\mathbf{x}} &= \frac{1}{g} \sum_{k=1}^g \bar{\mathbf{x}}_k, \\ \hat{\mathbf{B}}_\mu &= \sum_{k=1}^g (\bar{\mathbf{x}}_k - \bar{\mathbf{x}})(\bar{\mathbf{x}}_k - \bar{\mathbf{x}})^\top \quad \text{e} \\ \hat{\mathbf{W}} &= \sum_{k=1}^g \sum_{j=1}^{n_k} (\mathbf{x}_{jk} - \bar{\mathbf{x}}_k)(\mathbf{x}_{jk} - \bar{\mathbf{x}}_k)^\top.\end{aligned}$$

## Discriminantes Lineares Amostrais de Fisher (DALF)

Sejam  $\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_s$ ,  $s \leq \min\{g-1, p\}$  autovalores não-nulos de  $\widehat{\mathbf{W}}^{-1} \widehat{\mathbf{B}}_\mu$  e  $\hat{\mathbf{v}}_1, \hat{\mathbf{v}}_2, \dots, \hat{\mathbf{v}}_s$  os autovetores correspondentes tal que  $(\hat{\mathbf{v}}_k^\top \mathbf{S}_c^{-1} \hat{\mathbf{v}}_k = 1)$ .

Então, o vetor  $\hat{\mathbf{a}}$  que maximiza a razão  $\frac{\mathbf{a}^\top \widehat{\mathbf{B}}_\mu \mathbf{a}}{\mathbf{a}^\top \widehat{\mathbf{W}} \mathbf{a}}$  é dado por  $\hat{\mathbf{a}}_1 = \hat{\mathbf{v}}_1$ .

A combinação linear  $\hat{\mathbf{a}}_1^\top \mathbf{x}$  é chamada primeiro discriminante amostral.

A escolha  $\hat{\mathbf{a}}_2 = \hat{\mathbf{v}}_2$  produz o segundo discriminante amostral e continuamos até obter o  $k$ -ésimo discriminante amostral  $\hat{\mathbf{a}}_k = \hat{\mathbf{v}}_k$ ,  $k \leq s$ .

## Regressão Logística

As funções de classificação discutidas até aqui são baseadas em variáveis quantitativas. A regressão logística é uma abordagem apropriada para classificação quando algumas ou todas as variáveis são qualitativas. Na sua configuração mais simples, a variável resposta  $Y$  está restrita a dois valores. Por exemplo,  $Y$  pode representar gênero: macho/fêmea, ou empregado/desempregado, aprovado/reprovado, etc.

Quando a resposta assume apenas dois valores possíveis é comum codificá-la como 0 ou 1 e, o interesse passa a ser estimar a probabilidade da variável assumir o valor 1 dado o vetor de covariáveis  $\mathbf{x}$ , que representa a proporção na população codificada com o valor 1.

Esta modelagem pode então ser usada para fins de classificação em um de dois grupos, e a ideia pode ser estendida para vários grupos, substituindo a distribuição binomial pela multinomial.

## Inclusão de Variáveis Qualitativas

Neste parte assumimos que as variáveis de discriminação  $X_1, X_2, \dots, X_p$  são contínuas. Com frequência, uma variável qualitativa ou categórica pode ser útil como variável discriminante (classificadora). Esta situação é frequentemente contornada criando-se uma variável  $X$  cujo valor numérico é 1 se o objeto possui a tal característica e zero, caso contrário. A variável é, então, tratada como uma variável de medida nos procedimentos de classificação e discriminação usuais.

Exceto para classificação logística, há pouca teoria disponível para lidar com o caso em que algumas variáveis são contínuas e outras são qualitativas. Experimentos de simulação indicaram que a função discriminante linear de Fisher pode comportar-se tanto pobremente como satisfatoriamente, dependendo das correlações entre as variáveis contínuas e qualitativas. Krzanowski: “Uma correlação baixa em uma população, mas uma correlação alta na outra, ou uma mudança no sinal das correlações entre as duas populações poderiam indicar condições desfavoráveis à função discriminante linear de Fisher”. Esta é uma área problemática e que precisa de mais estudo.

## Árvores de Classificação

Uma abordagem de classificação completamente diferente dos métodos discutidos aqui foi desenvolvida. (Breiman, L., 1. Friedman, R Olshen, and C. Stone. Classification and Regression Trees. Belmont, CA: Wadsworth, Inc., 1984.) Ela é computacionalmente intensiva. A abordagem, chamada árvore de classificação e regressão (CART), é proximamente relacionada com as técnicas de conglomeração divisivas. Inicialmente, todos os objetos são considerados em um único grupo. O grupo é então dividido em dois subgrupos, usando, por exemplo, altos valores de uma variável para um grupo e baixos valores dessa mesma variável para o outro grupo. Os dois subgrupos são então cada um dividido novamente, agora usando valores de uma segunda variável. O processo de divisão continua até que um ponto de parada adequado seja atingido. Os valores das variáveis divisoras podem ser categorias ordenados ou não. É este aspecto que torna o CART tão geral.

## Redes Neurais

Uma rede neural é um procedimento computacional intensivo para transformar entradas em saídas programadas usando redes altamente conectadas de unidades de processamento relativamente simples (neurônios ou nós). Suas três características essenciais são as unidades básicas de computação (neurônios ou nós), a arquitetura da rede descrevendo as conexões entre as unidades de computação, e o algoritmo de treinamento usado para encontrar valores dos parâmetros da rede (pesos) para realizar uma tarefa particular.

As unidades de computação são conectadas umas às outras no sentido de que a saída de uma unidade pode servir como entrada para outra unidade. Cada unidade de computação transforma uma entrada em uma saída usando alguma função pré-especificada que é tipicamente monótona, mas de alguma forma arbitrária. Esta função depende de constantes (parâmetros) cujos valores devem ser determinados com um conjunto de treinamento de entradas e saídas.

Arquitetura da rede é a organização das unidades computacionais e os tipos de conexão permitidos. Em aplicações estatísticas, as unidades computacionais são arrumadas em uma série de camadas com conexões entre nós em camadas diferentes, mas não entre nós da mesma camada. A camada que recebe as entradas iniciais é chamada camada de entrada. A camada final é chamada camada de saída. Todas as camadas entre as camadas de entrada e saída são chamadas camadas ocultas.

Redes Neurais podem ser usadas para discriminação e classificação. Quando elas são usadas com este fim, as variáveis de entrada são as medidas  $X_1, X_2, \dots, X_p$ , e a variável de saída é a variável categórica que indica de qual grupo veio a observação de entrada. A experiência indica que redes neurais apropriadamente construídas comportam-se tão bem quanto à regressão logística e as funções discriminantes discutidas aqui. Os autores sugerem a seguinte referência para uma boa discussão do uso de redes neurais em aplicações da estatística: Stem, H. S. Neural Networks in Applied Statistics. *Technometrics*, 38, (1996), 205-214.

## Seleção de Variáveis

Em algumas aplicações da análise discriminante, os dados estão disponíveis para um grande número de variáveis. Mucciardi e Gose (A Comparison of Seven Techniques for Choosing Subsets of Pattern Recognition Properties. *IEEE Trans. Computers*, C20 (1971), 1023-1031.) estudaram uma análise discriminante baseada em 157 variáveis. Neste caso, seria obviamente desejável selecionar um subconjunto menor de variáveis que contivesse quase toda a informação original para efeitos da classificação. Este é o propósito da análise discriminante passo-a-passo *stepwise*, e vários programas de computador dispõem destas funções de seleção de variável.

Se uma análise discriminante *stepwise* (ou qualquer outro método de seleção) é empregado, os resultados devem ser interpretados com cautela. (Veja Murray, G. D. A Cautionary Note on Selection of Variables in Discriminant Analysis. *Applied Statistics*, 26, no. 3 (1977), 246-250.) Não há garantia de que o subconjunto selecionado seja o “melhor”, sem olhar o critério usado para fazer a seleção. Por exemplo, subconjuntos selecionados com base na minimização da taxa de erro aparente ou maximização do “poder de discriminação” podem comportar-se pobremente em amostras futuras. Problemas associados com procedimentos de seleção de variáveis são ampliados se existem correlações altas entre as variáveis ou entre combinações lineares das variáveis.

A escolha de um subconjunto de variáveis que parece ser ótima para um dado conjunto de dados é especialmente preocupante se a classificação é o objetivo. No mínimo, a função de classificação obtida deve ser avaliada com uma amostra de validação. Como Murray (1977) sugeriu, uma ideia melhor pode ser dividir a amostra em um número de lotes e determinar o “melhor” subconjunto para cada lote. O número de vezes que uma dada variável aparece nos melhores subconjuntos fornece uma medida do valor dessa variável para classificações futuras.