

ANÁLISE ESTATÍSTICA MULTIVARIADA

Ralph dos Santos Silva

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

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

Programa da disciplina

- Distribuições de vetores e matrizes aleatórias:
 - ◆ Normal multivariada;
 - ◆ t -Student multivariada; e
 - ◆ Wishart (inversa Wishart).
- Inferência sobre os parâmetros.
- Análise de variância multivariada.
- Regressão linear multivariada.
- Discriminação e classificação.
- Análise de componentes principais.
- Análise fatorial.
- Análise de conglomerados.
- Escalonamento multidimensional.

Referências

- Johnson e Wichern (2007). *Applied Multivariate Statistical Analysis*, 6th Edition. Prentice-Hall.
- Chatfield e Collins (1980). *Introduction to Multivariate Statistical Analysis*. Chapman and Hall.
- Anderson (1958). *An Introduction to Statistical Analysis*. John Wiley and Sons.

Organização de dados

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 denotado por $\mathbf{x}_{.i}$ 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 $p \times n$, isto é,

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

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

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

em que $\mathbf{x}_{j.}^\top = (x_{j1}, x_{j2}, \dots, x_{jn})$ é o vetor das n observações da j -ésima variável, $j = 1, 2, \dots, p$.

Também podemos representar a matriz $\mathbf{X}$ através de suas colunas.

Adotando 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

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

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,

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

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

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

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

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

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

Agora, podemos representar de forma organizada as informações sobre variabilidade através da **matriz de 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: $\mathbf{S} = \mathbf{S}^\top$.

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

$$\begin{aligned}\mathbf{S} &= \frac{1}{n-1}(\mathbf{X} - \bar{\mathbf{x}}\mathbf{1}^\top)(\mathbf{X} - \bar{\mathbf{x}}\mathbf{1}^\top)^\top \\ &= \frac{1}{n-1} \sum_{i=1}^n (\mathbf{x}_{.i} - \bar{\mathbf{x}})(\mathbf{x}_{.i} - \bar{\mathbf{x}})^\top.\end{aligned}$$

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: $\mathbf{R} = \mathbf{R}^\top$.

Algebricamente, a matriz de correlações pode ser escrita na forma

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

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:

$$\begin{aligned}w_{jj} &= \sum_{i=1}^n (x_{ji} - \bar{x}_j)^2, \quad j = 1, 2, \dots, p, \quad \text{e} \\w_{jk} &= \sum_{i=1}^n (x_{ji} - \bar{x}_j)(x_{ki} - \bar{x}_k), \quad j, k = 1, 2, \dots, p \quad \text{e} \quad j \neq k.\end{aligned}$$

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

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

Ver arquivo exemplo_01.r

Representação gráfica

Podemos utilizar gráficos de dispersão para variáveis duas a duas.

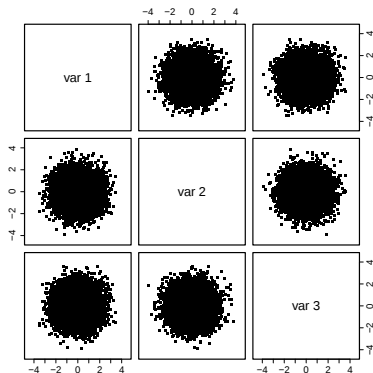


Figura 1: Exemplo de gráfico de dispersão de variáveis duas a duas.

Ver arquivo exemplo_02.r

Exemplo

O papel é fabricado em folhas contínuas com vários metros de largura. Por causa da orientação das fibras dentro do papel, ele tem uma força diferente quando medido na direção produzida pela máquina do que quando medido em ângulos retos para a direção da máquina.

Temos as seguintes variáveis:

- x_1 : densidade (gramas/centímetro cúbico);
- x_2 : força (libras) na direção da máquina; e
- x_3 : força (libras) na direção cruzada.

Ver arquivo exemplo_03.r

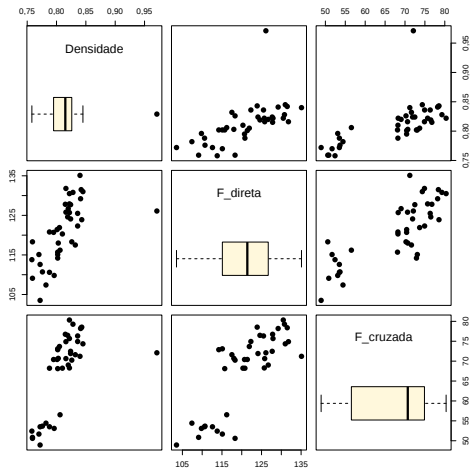


Figura 2: Gráficos de dispersão e *boxplots* das variáveis densidade, força na direção da máquina e força na direção cruzada referentes à qualidade do papel.

Distâncias

A maior parte das técnicas multivariadas se baseia 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.

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 ($\sqrt{\sigma_{jj}}$) e calculando a distância euclidiana

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

em que $\boldsymbol{\Omega}$ é uma matriz diagonal com elementos $\sigma_{11}, \sigma_{22}, \dots, \sigma_{pp}$, $\mathbf{x} = (x_1, x_2, \dots, x_p)^\top$ 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

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

em que $\boldsymbol{\Sigma}$ é a matriz de covariâncias.

Distribuição normal

Seja X uma variável aleatória contínua com distribuição normal com média μ e variância σ^2 . Então,

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

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

Temos que

$$\Pr(|X - \mu| \leq \sigma) = \Pr(\mu - \sigma \leq X \leq \mu + \sigma) \simeq 0,683.$$

$$\Pr(|X - \mu| \leq 2\sigma) = \Pr(\mu - 2\sigma \leq X \leq \mu + 2\sigma) \simeq 0,954.$$

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

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

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

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

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

em que a matriz de pesos é 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

Consideramos a seguir a particularização da normal multivariada para o caso bivariado.

Considere o seguinte:

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

e

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

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

Ver arquivo exemplo_04.r

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 Σ

Podemos representar a distribuição normal multivariada como uma transformação linear de variáveis aleatórias normais univariadas?

Podemos representar a distribuição normal multivariada como uma transformação linear de variáveis aleatórias normais **padrão** univariadas?

Podemos ver a distribuição normal multivariada como uma elipse p -dimensional?

Sim é a resposta para as três perguntas acima.

Podemos, por exemplo, empregar a decomposição em valores singulares (SVD) da matriz de covariâncias Σ para obter

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

em que $\mathbf{V}$ é uma matriz $p \times p$ ortonormal (as colunas da matriz são ortogonais e cada uma tem norma 1) e $\mathbf{D}$ é uma matriz $p \times p$ diagonal com elementos $d_j, j = 1, \dots, p$, positivos. Em geral, os elementos d_j estão ordenados de forma não crescente.

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 μ_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 $\mathbf{L} = \mathbf{U}^\top$.

Ver arquivo exemplo_05.r

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

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

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

Revisão

Seja X uma variável aleatória contínua com função de densidade de probabilidade $f(x)$.

Momentos ordinários:

$$E(X^\ell) = \int_{-\infty}^{\infty} x^\ell f(x) dx.$$

Definimos $\mu_x = E(X)$ como a **média**.

Momentos centrais:

$$E([X - \mu_x]^\ell) = \int_{-\infty}^{\infty} (x - \mu_x)^\ell f(x) dx.$$

Definimos $\sigma_x^2 = E([X - \mu_x]^2)$ como a **variância**.

Definimos $\sigma_x = \sqrt{\sigma_x^2}$ como **desvio padrão**.

O **coeficiente de assimetria** é definido por

$$S(X) = E \left(\left[\frac{X - \mu_X}{\sigma_X} \right]^3 \right) = \frac{E([X - \mu_X]^3)}{\sigma_X^3}$$

e o **coeficiente de curtose** por

$$K(X) = E \left(\left[\frac{X - \mu_X}{\sigma_X} \right]^4 \right) = \frac{E([X - \mu_X]^4)}{\sigma_X^4}.$$

A quantidade $K(X) - 3$ é chamada de **excesso de curtose** porque $K(X) = 3$ para a distribuição normal¹.

Uma distribuição com excesso de curtose positivo é dita ter **caudas pesadas** (leptocúrtica).

Uma distribuição com excesso de curtose negativo é dita ter **caudas leves** (platicúrtica).

¹Se $X \sim \mathcal{N}(\mu, \sigma^2)$, então $K(X) = 3$.

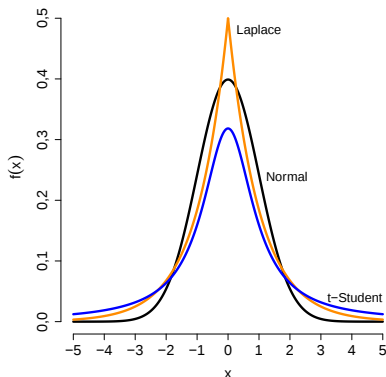


Figura 3: Comparação das funções de densidades de probabilidades da normal, Laplace e t -Student com 1 grau de liberdade.

- A curtose da Laplace é maior que a curtose da normal.
- A curtose da t -Student² é maior que a curtose da Laplace.

²Distribuição t -Student com 1 grau de liberdade ou distribuição de Cauchy.

Momentos amostrais (unidimensionais)

Suponha que $\{X_1, X_2, \dots, X_n\}$ seja uma amostra aleatória de X com n observações.

- Média amostral:

$$\hat{\mu}_x = \bar{X} = \frac{1}{n} \sum_{i=1}^n X_i.$$

- Variância amostral:

$$\hat{\sigma}_x^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2.$$

- O coeficiente de assimetria amostral:

$$\hat{S}(X) = \frac{M^3}{\hat{\sigma}_x^3} \quad \text{com} \quad M^3 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^3.$$

- O coeficiente de curtose amostral:

$$\hat{K}(X) = \frac{M^4}{\hat{\sigma}_x^4} \quad \text{com} \quad M^4 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^4.$$

M^ℓ , $\ell = 1, 2, \dots$, podem ser chamados de **momentos centrais amostrais**.

Testes de assimetria e curtose

Sob a hipótese de normalidade (**que a amostra aleatória é proveniente de uma distribuição normal**), temos que

$$\hat{S}(X) \approx \mathcal{N}(0, 6/n) \quad \text{e} \quad \hat{K}(X) \approx \mathcal{N}(3, 24/n)$$

para n “suficientemente grande”.

Dado uma amostra aleatória $\{X_1, X_2, \dots, X_n\}$, para testar a assimetria da distribuição dos dados, consideramos

$$\begin{cases} H_0 : & S(X) = 0 \\ H_1 : & S(X) \neq 0. \end{cases}$$

A estatística da razão t -Student da assimetria é

$$T = \frac{\hat{S}(X)}{\sqrt{6/n}} \approx \mathcal{N}(0, 1).$$

Rejeitamos H_0 ao nível α de significância se $|T_{obs}| > z_{1-\alpha/2}$, em que $z_{1-\alpha/2}$ é o percentil $100(1 - \alpha/2)$ da distribuição normal padrão.

Para testar o excesso de curtose da distribuição dos dados, consideramos

$$\begin{cases} H_0 : & K(X) - 3 = 0 \\ H_1 : & K(X) - 3 \neq 0. \end{cases}$$

A estatística da razão t -Student do excesso de curtose é

$$T = \frac{\hat{K}(X) - 3}{\sqrt{24/n}} \approx \mathcal{N}(0, 1).$$

Rejeitamos H_0 ao nível α de significância se $|T_{obs}| > z_{1-\alpha/2}$, em que $z_{1-\alpha/2}$ é o percentil $100(1 - \alpha/2)$ da distribuição normal padrão.

Teste de normalidade de Jarque-Bera

Considere o seguinte teste de hipóteses:

$$\begin{cases} H_0 : & \text{A amostra aleatória é proveniente de uma dist. normal;} \\ H_1 : & \text{A amostra aleatória não é proveniente de uma dist. normal.} \end{cases}$$

A estatística do teste de normalidade de **Jarque e Bera** é dada por

$$JB = \frac{\hat{S}^2(X)}{6/n} + \frac{(\hat{K}(X) - 3)^2}{24/n} \approx \chi_2^2,$$

para n “suficientemente grande”.

Rejeitamos H_0 (normalidade) se $JB > \chi_{1-\alpha}^*$, em que $\chi_{1-\alpha}^*$ é o percentil $100(1 - \alpha)$ da distribuição χ_2^2 .

Ver arquivos exemplo_06.r

Distribuição gama

Seja X uma variável aleatória com distribuição gama com parâmetros de forma α e de escala β . Então, a função de densidade de probabilidade é dada por

$$f(x|\alpha, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} \exp\{-\beta x\} I_{[0, \infty)}(x), \quad \alpha > 0 \quad \text{e} \quad \beta > 0.$$

Temos que

$$\begin{aligned} E(X) &= \frac{\alpha}{\beta} \\ \text{Var}(X) &= \frac{\alpha}{\beta^2} \end{aligned}$$

Denotaremos a distribuição gama com parâmetros α e β por $\mathcal{G}(\alpha, \beta)$.

Distribuição inversa gama

Seja Y uma variável aleatória com distribuição gama com parâmetros de forma α e de escala β e $X = 1/Y$. Então, a função de densidade de probabilidade de X é dada por

$$f(x|\alpha, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha)} x^{-(\alpha+1)} \exp\left\{-\frac{\beta}{x}\right\} I_{[0, \infty]}(x), \quad \alpha > 0 \quad \text{e} \quad \beta > 0.$$

Temos que

$$E(X) = \frac{\beta}{\alpha - 1}, \quad \text{para } \alpha > 1$$

$$\text{Var}(X) = \frac{\beta^2}{(\alpha - 1)^2(\alpha - 2)}, \quad \text{para } \alpha > 2$$

Denotaremos a distribuição gama invertida com parâmetros α e β por $\mathcal{IG}(\alpha, \beta)$.

Distribuição qui-quadrada

Seja X uma variável aleatória com distribuição qui-quadrada com (parâmetro) n graus de liberdade. Então, a função de densidade de probabilidade é dada por

$$f(x|n) = \frac{1}{2^{n/2}\Gamma\left(\frac{n}{2}\right)} x^{n/2-1} \exp\left\{-\frac{x}{2}\right\} I_{[0,\infty]}(x), \quad \alpha > 0 \quad \text{e} \quad \beta > 0.$$

Temos que

$$\begin{aligned} E(X) &= n \\ \text{Var}(X) &= 2n \end{aligned}$$

Denotaremos a distribuição qui-quadrada com n graus de liberdade por χ_n^2 .

Distribuição t -Student

Seja X uma variável aleatória com distribuição t -Student com parâmetros localização μ , de escala σ^2 e ν graus de liberdade. Então, a função de densidade de probabilidade é dada por

$$\begin{aligned} f(x|\mu, \sigma^2, \nu) &= \frac{\Gamma\left(\frac{\nu+1}{2}\right)}{\Gamma\left(\frac{\nu}{2}\right) \Gamma\left(\frac{1}{2}\right) \nu^{1/2} \sigma} \left(1 + \frac{(x-\mu)^2}{\nu \sigma^2}\right)^{-(\nu+1)/2}, \\ &= \frac{\Gamma\left(\frac{\nu+1}{2}\right)}{\Gamma\left(\frac{\nu}{2}\right) \pi^{1/2} \nu^{1/2} \sigma} \left(1 + \frac{(x-\mu)^2}{\nu \sigma^2}\right)^{-(\nu+1)/2}, \quad x \in \mathbb{R}, \end{aligned}$$

em que $\mu \in \mathbb{R}$, $\nu > 0$ e $\sigma^2 > 0$. Temos que

$$\begin{aligned} E(X) &= \mu, & \text{se } \nu > 1 \\ \text{Var}(X) &= \frac{\nu}{\nu-2} \sigma^2, & \text{se } \nu > 2. \end{aligned}$$

Denotaremos a distribuição t -Student por $t_\nu(\mu, \sigma^2)$.

Distribuição F-Snedecor

Seja X uma variável aleatória com distribuição F-Snedecor com (parâmetros) graus de liberdades n e m . Então, a função de densidade de probabilidade é dada por

$$f(x|n, m) = \left(\frac{n}{m}\right)^{n/2} \frac{\Gamma\left(\frac{m+n}{2}\right)}{\Gamma\left(\frac{n}{2}\right) \Gamma\left(\frac{m}{2}\right)} x^{n/2-1} \left(1 + \frac{n}{m}x\right)^{-(n+m)/2} I_{[0, \infty)}(x),$$

em que $n > 0$ e $m > 0$.

Denotaremos a distribuição F-Snedecor com n e m graus de liberdade por $\mathcal{F}_{n,m}$.

Resultado

Se Z e W são variáveis aleatórias independentes tal que $Z \sim \mathcal{N}(0, 1)$, $W \sim \chi_n^2$ e

$$T = \frac{Z}{\sqrt{W/n}},$$

então $T \sim t_n(0, 1)$.

Prova

$$\begin{aligned} F_T(t) &= \Pr(T \leq t) = \Pr(Z \leq t\sqrt{W/n}) \\ &= \int_0^\infty \Pr(Z \leq t\sqrt{w/n} \mid W = w) f_W(w) dw \\ &= \int_0^\infty \Phi(t\sqrt{w/n}) f_W(w) dw. \end{aligned}$$

$$\begin{aligned}
f_T(t) &= \frac{\partial F_T(t)}{\partial t} = \frac{\partial}{\partial t} \int_0^\infty \Phi(t\sqrt{w/n}) f_W(w) dw \\
&= \int_0^\infty \frac{\partial}{\partial t} \Phi(t\sqrt{w/n}) f_W(w) dw = \int_0^\infty \phi(t\sqrt{w/n}) \sqrt{\frac{w}{n}} f_W(w) dw \\
&= \frac{1}{\Gamma\left(\frac{n}{2}\right) \Gamma\left(\frac{1}{2}\right) 2^{(n+1)/2} n^{1/2}} \int_0^\infty w^{(n+1)/2-1} \exp\left\{-w \frac{1}{2} \left(1 + \frac{t^2}{n}\right)\right\} dw \\
&= \frac{\Gamma\left(\frac{n+1}{2}\right)}{\Gamma\left(\frac{n}{2}\right) \Gamma\left(\frac{1}{2}\right) n^{1/2}} \left(1 + \frac{t^2}{n}\right)^{-(n+1)/2}, \quad t \in \mathbb{R}.
\end{aligned}$$

Portanto, $T \sim t_n(0, 1)$.

Resultado

Se $T \sim t_n(0, 1)$ e $Y = T^2$, então $Y \sim \mathcal{F}_{1,n}$.

Prova

$$\begin{aligned}F_Y(y) &= \Pr(Y \leq y) = \Pr(T^2 \leq y) = \Pr(-\sqrt{y} \leq T \leq \sqrt{y}) \\&= \Pr(T \leq \sqrt{y}) - \Pr(T \leq -\sqrt{y}) = F_T(\sqrt{y}) - F_T(-\sqrt{y}) \\f_Y(y) &= \frac{\partial F_Y(y)}{\partial y} = f_T(\sqrt{y}) \frac{1}{2\sqrt{y}} + f_T(-\sqrt{y}) \frac{1}{2\sqrt{y}} \\&= \frac{\Gamma\left(\frac{n+1}{2}\right)}{\Gamma\left(\frac{n}{2}\right) \Gamma\left(\frac{1}{2}\right) n^{1/2}} y^{1/2-1} \left(1 + \frac{y}{n}\right)^{-(n+1)/2}, \quad y \geq 0.\end{aligned}$$

Portanto, $Y \sim \mathcal{F}_{1,n}$.

Distribuição t -Student multivariada

Seja $\mathbf{X}$ um vetor aleatório p -dimensional com distribuição t -Student multivariada com (parâmetros) **vetor de localização** $\boldsymbol{\mu}$, **matriz de escala** $\boldsymbol{\Sigma}$ e **graus de liberdade** ν . Então, a função de densidade de probabilidade é dada por

$$f(\mathbf{x}|\boldsymbol{\mu}, \boldsymbol{\Sigma}, \nu) = \frac{\Gamma\left(\frac{\nu+p}{2}\right)}{\Gamma\left(\frac{\nu}{2}\right) \pi^{p/2} \nu^{p/2} |\boldsymbol{\Sigma}|^{1/2}} \left(1 + \frac{1}{\nu} (\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu})\right)^{-(\nu+p)/2},$$

para $\mathbf{x} \in \mathbb{R}^p$, com $\boldsymbol{\mu} \in \mathbb{R}^p$, $\boldsymbol{\Sigma}$ uma matriz $p \times p$ simétrica e positiva definida, e $\nu > 0$. Temos que

$$\begin{aligned} E(\mathbf{X}) &= \boldsymbol{\mu}, & \text{se } \nu > 1 \\ \text{Var}(\mathbf{X}) &= \frac{\nu}{\nu-2} \boldsymbol{\Sigma}, & \text{se } \nu > 2. \end{aligned}$$

Denotaremos a distribuição t -Student multivariada por $t_\nu(\boldsymbol{\mu}, \boldsymbol{\Sigma})$.

Distribuição Wishart

Seja $\mathbf{X}$ uma matriz aleatória $p \times p$ com distribuição Wishart com (parâmetros) **matriz de escala** $\mathbf{V}$ e **graus de liberdade** ν . Então, a função de densidade de probabilidade é dada por

$$f(\mathbf{X}|\mathbf{V}, \nu) = \frac{|\mathbf{X}|^{(\nu-p-1)/2} \exp \left\{ -\frac{1}{2} \text{tr}(\mathbf{V}^{-1} \mathbf{X}) \right\}}{2^{(\nu p)/2} |\mathbf{V}|^{\nu/2} \pi^{p(p-1)/4} \prod_{j=1}^p \Gamma \left(\frac{\nu}{2} + \frac{(1-j)}{2} \right)},$$

para $\mathbf{X}$ no espaço das matrizes $p \times p$, simétricas e positiva definida. Além disso, $\mathbf{V}$ é uma matriz $p \times p$ simétrica e positiva definida, e $\nu \geq p$. Temos que

$$\begin{aligned} E(\mathbf{X}) &= \nu \mathbf{V} \\ \text{Var}(X_{jk}) &= \nu (v_{jk}^2 + v_{jj} v_{kk}), \quad j, k = 1, 2, \dots, p. \end{aligned}$$

Denotaremos a distribuição Wishart por $\mathcal{W}(\mathbf{V}, \nu)$.

Resultado

Seja $(\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n)$ uma amostra aleatória de $\mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ tal que $n \geq p$.
Então,

$$\sum_{i=1}^n (\mathbf{X}_i - \boldsymbol{\mu})(\mathbf{X}_i - \boldsymbol{\mu})^\top \sim \mathcal{W}(\boldsymbol{\Sigma}, n).$$

Resultado

Se $\mathbf{X} \sim \mathcal{W}(\mathbf{V}, \nu)$ e $\mathbf{Y} = \mathbf{X}^{-1}$, então $\mathbf{Y} \sim \mathcal{IW}(\mathbf{U}, \nu)$ - chamada de **inversa Wishart** com $\mathbf{U} = \mathbf{V}^{-1}$.

A função de densidade de probabilidade de $\mathbf{Y}$ é dada por

$$f(\mathbf{Y}|\mathbf{U}, \nu) = \frac{|\mathbf{U}|^{\nu/2} |\mathbf{Y}|^{-(\nu+p+1)/2} \exp \left\{ -\frac{1}{2} \text{tr}(\mathbf{U}\mathbf{Y}^{-1}) \right\}}{2^{(\nu p)/2} \pi^{p(p-1)/4} \prod_{j=1}^p \Gamma \left(\frac{\nu}{2} + \frac{(1-j)}{2} \right)},$$

Temos que

$$\mathbb{E}(\mathbf{Y}) = \frac{\mathbf{U}}{\nu - p - 1}, \quad \text{para } \nu > p + 1.$$

Normal multivariada: inferência sobre o vetor de médias

Considere o problema univariado no qual se tem uma amostra aleatória de tamanho n , $(X_1, X_2, \dots, X_n)$, da distribuição $\mathcal{N}(\mu, \sigma^2)$, em que ambos os parâmetros - média e variância - são desconhecidos.

A estatística dada por

$$T = \frac{\bar{X} - \mu}{\sqrt{S^2/n}}$$

tem distribuição t -Student com $n - 1$ graus de liberdade.

$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ e $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$ são as estatísticas utilizadas.

Podemos tomar o quadrado da estatística t na forma

$$T^2 = \sqrt{n}(\bar{X} - \mu)^\top (S^2)^{-1} \sqrt{n}(\bar{X} - \mu).$$

Neste caso, a distribuição T^2 é uma F -Snedecor com 1 e $n - 1$ graus de liberdade.

Notação: $T^2 \sim \mathcal{F}_{1,n-1}$.

Esta última expressão nos sugere a seguinte versão multivariada desta estatística:

$$T^2 = n(\bar{\mathbf{X}} - \boldsymbol{\mu})^\top \mathbf{S}^{-1}(\bar{\mathbf{X}} - \boldsymbol{\mu}),$$

em que

$$\bar{\mathbf{X}} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \quad \text{e} \quad \mathbf{S} = \frac{1}{n-1} \sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{X}})(\mathbf{x}_i - \bar{\mathbf{X}})^\top,$$

com $(\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n)$ sendo uma amostra aleatória da $\mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$.

A estatística T^2 é chamada **Estatística T^2 de Hotelling** em homenagem a Harold Hotelling.

Proposição

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

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

Natureza da estatística T^2 :

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

em que

$$\sqrt{n}(\bar{\mathbf{X}} - \boldsymbol{\mu}) \quad \text{e} \quad \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^\top$$

são independentes.

Teste de hipóteses

Considere o problema de testar as hipóteses

$$\begin{cases} H_0 : \mu = \mu_0 \\ H_1 : \mu \neq \mu_0, \end{cases}$$

quando se tem uma amostra aleatória da distribuição normal multivariada.

A estatística de teste neste caso será dada por

$$T^2 = n(\bar{\mathbf{X}} - \mu)^\top \mathbf{S}^{-1}(\bar{\mathbf{X}} - \mu).$$

Sob H_0 , a estatística T^2 tem distribuição $\frac{(n-1)p}{n-p} \mathcal{F}_{p, n-p}$ de acordo com a Proposição anterior. Assim, ao nível de significância α , rejeitamos H_0 se

$$T_{obs}^2 > \frac{(n-1)p}{n-p} F_{p, n-p, 1-\alpha} \quad \text{para} \quad T_{obs}^2 = n(\bar{\mathbf{X}} - \mu)^\top \mathbf{S}^{-1}(\bar{\mathbf{X}} - \mu),$$

com $F_{p, n-p, 1-\alpha}$ representando o quantil de $100(1 - \alpha)\%$ da distribuição $\mathcal{F}_{p, n-p}$.

Notação: $\mathbf{S}$ representa a matriz aleatória e também sua realização (o valor observado). Em geral, omitiremos $\{obs\}$ de T_{obs}^2 , isto é, utilizaremos só T^2 .

Exemplo

A transpiração de 20 mulheres saudáveis foi analisada produzindo X_1 - taxa de suor, X_2 - conteúdo de sódio e X_3 - conteúdo de potássio. Deseja-se testar, ao nível de significância de 10%, as hipóteses

$$\begin{cases} H_0 : \mu = (4, 50, 10)^\top \\ H_1 : \mu \neq (4, 50, 10)^\top. \end{cases}$$

Lembrem-se antes de verificar a normalidade dos dados!

Após os cálculos, obteve-se $T^2 \simeq 9,74$ e $RC = \{T^2 : T^2 > 8,17\}$.

Logo, ao nível de significância de 10%, rejeitamos a hipótese nula.

Qual é o valor p deste teste? Calcule $T^2(n-p)/(p(n-1))$ e obtenha a probabilidade da cauda superior da distribuição $\mathcal{F}_{p,n-p}$ associada a este valor.

Temos, $p = 0,0649$. Portanto, para qualquer nível de significância menor que 6,5%, H_0 não seria rejeitada.

Ver arquivo [exemplo_07.r](#)

Invariância da estatística T^2

A estatística T^2 é invariante sob transformações de localização e escala.

Defina $\mathbf{Y} = \mathbf{C}\mathbf{X} + \mathbf{d}$, com $\mathbf{X}$ ($p \times 1$), $\mathbf{Y}$ ($p \times 1$) vetores aleatórios, $\mathbf{d}$ vetor ($p \times 1$) de constantes fixadas e $\mathbf{C}$ matriz ($p \times p$) não-singular de constantes fixadas.

Então, $\bar{\mathbf{Y}} = \mathbf{C}\bar{\mathbf{X}} + \mathbf{d}$ e $\mathbf{S}_Y = \mathbf{C}\mathbf{S}_X\mathbf{C}^\top$. Além disso, $\mu_Y = \mathbf{C}\mu_X + \mathbf{d}$.

Assim,

$$\begin{aligned}T_Y^2 &= n(\bar{\mathbf{Y}} - \mu_Y)^\top \mathbf{S}_Y^{-1}(\bar{\mathbf{Y}} - \mu_Y) \\&= n(\mathbf{C}\bar{\mathbf{X}} - \mathbf{C}\mu_X)^\top (\mathbf{C}\mathbf{S}_X\mathbf{C}^\top)^{-1} (\mathbf{C}\bar{\mathbf{X}} - \mathbf{C}\mu_X) \\&= n(\bar{\mathbf{X}} - \mu_X)^\top \underbrace{\mathbf{C}^\top (\mathbf{C}^\top)^{-1}}_{\mathbf{I}_p} \mathbf{S}_X^{-1} \underbrace{\mathbf{C}^{-1} \mathbf{C}}_{\mathbf{I}_p} (\bar{\mathbf{X}} - \mu_X) \\&= n(\bar{\mathbf{X}} - \mu_X)^\top \mathbf{S}_X^{-1}(\bar{\mathbf{X}} - \mu_X) \\&= T_X^2.\end{aligned}$$

A estatística T^2 e o teste de razão de verossimilhanças

Uma metodologia muito usada na construção de testes de hipóteses é conhecida como teste de razão de verossimilhanças (**teste RV**).

Em linhas gerais, se temos uma amostra aleatória de uma distribuição que depende de um vetor de parâmetros θ cuja densidade é $f_n(\mathbf{x}|\theta)$ e desejamos testar

$$\begin{cases} H_0 : & \theta \in \Theta_0 \\ H_1 : & \theta \notin \Theta_0, \end{cases}$$

a estatística do teste RV é dada por

$$\Lambda(\mathbf{x}) = \frac{\max_{\theta \in \Theta_0} L(\theta|\mathbf{x})}{\max_{\theta} L(\theta|\mathbf{x})}$$

com $L(\theta|\mathbf{x})$ a função de verossimilhança.

Observe que a estatística $\Lambda(\mathbf{x})$ é um número entre 0 e 1.

Se o máximo sob H_0 for o máximo global, teremos $\Lambda(\mathbf{x}) = 1$.

Caso contrário, $\Lambda(\mathbf{x}) < 1$ e como $\Lambda(\mathbf{x})$ representa uma razão entre quantidades positivas, segue que $0 < \Lambda(\mathbf{x}) \leq 1$.

Assim, é razoável dizer que a hipótese nula será rejeitada para valores pequenos de $\Lambda(\mathbf{x})$ tal que as regiões críticas nos testes RV são da forma $\Lambda(\mathbf{x}) \leq c$.

Para obter o valor de c é necessário conhecer a distribuição amostral de $\Lambda(\mathbf{x})$. Esta distribuição nem sempre é fácil de ser obtida e muitas vezes são necessários métodos aproximados para avaliar o valor de c para um dado nível de significância.

Distribuição assintótica da estatística Λ

Temos que

$$-2 \ln \Lambda \stackrel{a}{\sim} \chi^2_{\nu - \nu_0}$$

em que ν é a dimensão do espaço de parâmetros e ν_0 é a dimensão do sub-espaço correspondente à hipótese nula.

Calcularemos a estatística $\Lambda(\mathbf{x})$ no contexto do teste das hipóteses

$$\begin{cases} H_0 : & \boldsymbol{\mu} = \boldsymbol{\mu}_0 \\ H_1 : & \boldsymbol{\mu} \neq \boldsymbol{\mu}_0, \end{cases}$$

quando se tem uma amostra aleatória da distribuição normal multivariada.

O máximo global da função de verossimilhança é obtido quando

$$\hat{\mu} = \bar{\mathbf{X}} \quad \text{e} \quad \hat{\Sigma} = \frac{1}{n} \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^\top = \frac{n-1}{n} \mathbf{S}$$

e dado por

$$L(\hat{\mu}, \hat{\Sigma}) = (2\pi)^{-np/2} |\hat{\Sigma}|^{-n/2} e^{-np/2}.$$

Portanto, já temos o denominador da estatística $\Lambda(\mathbf{x})$, neste caso.

Para obter o numerador, observe que sob H_0 há apenas um valor possível para μ dado por μ_0 . Neste caso, usando os resultados apresentados na seção de estimação por máxima verossimilhança, é fácil ver que a matriz que maximiza a verossimilhança sob H_0 é dada por

$$\hat{\Sigma}_0 = \frac{1}{n} \sum_{i=1}^n (\mathbf{X}_i - \mu_0)(\mathbf{X}_i - \mu_0)^\top.$$

Assim, o numerador de $\Lambda(\mathbf{x})$ é

$$L(\mu_0, \hat{\Sigma}_0) = (2\pi)^{-np/2} |\hat{\Sigma}_0|^{-n/2} e^{-np/2}.$$

Logo,

$$\Lambda = \left(\frac{|\hat{\Sigma}|}{|\hat{\Sigma}_0|} \right)^{n/2} = \left(\frac{\left| \sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{X}})(\mathbf{x}_i - \bar{\mathbf{X}})^\top \right|}{\left| \sum_{i=1}^n (\mathbf{x}_i - \mu_0)(\mathbf{x}_i - \mu_0)^\top \right|} \right)^{n/2}.$$

A estatística equivalente $\Lambda^{2/n} = |\hat{\Sigma}|/|\hat{\Sigma}_0|$ é chamada **lambda de Wilks**.

é fácil, por meio de artifícios algébricos, mostrar que as estatísticas T^2 e Λ para o teste aqui considerado satisfazem a relação

$$\Lambda^{2/n} = \left(1 + \frac{T^2}{n-1} \right)^{-1}$$

e, usando esta relação, chegamos a uma outra expressão para o cálculo da estatística T^2 :

$$T^2 = \frac{(n-1)|\hat{\Sigma}_0|}{|\hat{\Sigma}|} - (n-1).$$

Ver arquivo [exemplo_07.r](#)

Regiões de confiança e intervalos simultâneos

Dizemos que $R(\mathbf{X})$ é uma região de $100(1 - \alpha)\%$ de confiança para θ se

$$\Pr(R(\mathbf{X}) \text{ compreender } \theta) = 1 - \alpha.$$

A região de confiança para o vetor de médias μ quando se dispõe de uma amostra aleatória da distribuição $\mathcal{N}_p(\mu, \Sigma)$ é dada por

$$n(\bar{\mathbf{x}} - \mu)^\top \mathbf{S}^{-1}(\bar{\mathbf{x}} - \mu) \leq \frac{(n-1)p}{n-p} F_{p, n-p, 1-\alpha},$$

em que $F_{p, n-p, 1-\alpha}$ é o quantil acumulado de $100(1 - \alpha)\%$ da distribuição $\mathcal{F}_{p, n-p}$.

Este resultado é obtido se utilizando a distribuição amostral da estatística T^2 apresentada anteriormente.

Observe que a região de confiança é dada pelo hiper-elipsóide de eixos determinados pelos autovetores da matriz de covariância amostral $\mathbf{S}$ e cujas medidas são proporcionais às raízes quadradas dos respectivos autovalores.

Para verificar se um dado vetor μ_0 pertence à região de confiança, basta calcular $n(\bar{\mathbf{x}} - \mu_0)^\top \mathbf{S}^{-1}(\bar{\mathbf{x}} - \mu_0)$ e comparar com $\frac{(n-1)p}{n-p} F_{p, n-p, 1-\alpha}$.

Para $p \geq 4$ não é possível representar visualmente a região de confiança. No entanto, podemos calcular as medidas dos eixos do hiper-elipsóide de confiança centrado em $\bar{\mathbf{x}}$:

$$n(\bar{\mathbf{x}} - \mu)^\top \mathbf{S}^{-1}(\bar{\mathbf{x}} - \mu) \leq c^2 = \frac{(n-1)p}{n-p} F_{p, n-p, 1-\alpha}.$$

Os semi-eixos têm medida

$$\sqrt{\lambda_j} \frac{c}{\sqrt{n}} = \sqrt{\lambda_j \frac{(n-1)p}{n(n-p)} F_{p, n-p, 1-\alpha}}.$$

Exemplo

O departamento de controle de qualidade de uma fabricante de fornos de microondas foi cobrado pelo governo federal a monitorar a quantidade de radiação quando as portas dos microondas são fechadas. Observações da radiação emitida através das portas fechadas de $n = 42$ fornos selecionados ao acaso foram feitas. Medidas de radiação também foram feitas com as portas abertas dos 42 fornos selecionados.

Se verificarmos a suposição de normalidade, veremos que esta não é apropriada e uma transformação potência dos dados é buscada.

Sugere-se trabalhar com a potência 0,25, ou seja, a raiz quarta da escala original da medida de radiação.

Pede-se construir uma elipse de 95% de confiança para o vetor de médias, considerando a escala dos dados transformados de modo que a suposição de normalidade seja razoável.

Ver arquivo exemplo_08.r

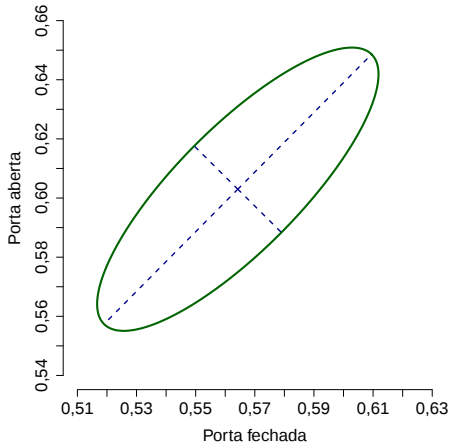


Figura 4: Região de 95% de confiança para o vetor de médias, μ , das radiações de microondas de portas abertas e fechadas, respectivamente.

Intervalos de confiança simultâneos para as componentes do vetor de médias

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

Vimos que se $\mathbf{a}$ é um vetor de constantes em $\mathbb{R}^p$, então

$$Z = \mathbf{a}^\top \mathbf{X} \sim \mathcal{N}(\mathbf{a}^\top \boldsymbol{\mu}, \mathbf{a}^\top \boldsymbol{\Sigma} \mathbf{a}).$$

Logo, se $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$ é uma amostra aleatória da $\mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$, segue que $Z_1, Z_2, \dots, Z_n$, definidos por $Z_i = \mathbf{a}^\top \mathbf{X}_i$, $i = 1, \dots, n$, é uma amostra aleatória da $\mathcal{N}(\underbrace{\mathbf{a}^\top \boldsymbol{\mu}}_{\mu_Z}, \underbrace{\mathbf{a}^\top \boldsymbol{\Sigma} \mathbf{a}}_{\sigma_Z^2})$.

Da teoria normal univariada, temos que um intervalo de $100(1 - \alpha)\%$ de confiança para $\mu_Z = \mathbf{a}^\top \boldsymbol{\mu}$ é dado por

$$\text{IC}_{100(1-\alpha)\%}(\mu_Z) = \left(\mathbf{a}^\top \bar{\mathbf{x}} - t_{n-1, 1-\frac{\alpha}{2}} \sqrt{\frac{\mathbf{a}^\top \mathbf{S} \mathbf{a}}{n}}; \mathbf{a}^\top \bar{\mathbf{x}} + t_{n-1, (1-\frac{\alpha}{2})} \sqrt{\frac{\mathbf{a}^\top \mathbf{S} \mathbf{a}}{n}} \right)$$

Claramente, poderíamos construir vários intervalos de confiança sobre combinações lineares dos componentes do vetor μ , cada um associado com um coeficiente de confiança $1 - \alpha$, escolhendo diferentes vetores de constantes $\mathbf{a}$. Porém, o coeficiente de confiança conjunto do conjunto de intervalos resultantes não seria mais $1 - \alpha$.

É desejável associar um coeficiente de confiança **coletivo** de $1 - \alpha$ aos intervalos de confiança que podem ser gerados para todas as escolhas de $\mathbf{a}$.

Naturalmente, um preço deverá ser pago pela conveniência de uma confiança simultânea grande para todos os intervalos: intervalos que têm amplitudes maiores (mais largos, menos precisos) do que os intervalos apresentados anteriormente via a distribuição amostral t -Student com $n - 1$ graus de liberdade.

Dado o conjunto de dados observados $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ e um $\mathbf{a}$ particular

$$|t| = \frac{|\sqrt{n}(\mathbf{a}^\top \bar{\mathbf{x}} - \mathbf{a}^\top \boldsymbol{\mu})|}{\sqrt{\mathbf{a}^\top \mathbf{S} \mathbf{a}}} \leq t_{n-1, 1-\frac{\alpha}{2}}$$

ou, equivalentemente,

$$t^2 = \frac{n(\mathbf{a}^\top \bar{\mathbf{x}} - \mathbf{a}^\top \boldsymbol{\mu})^2}{\mathbf{a}^\top \mathbf{S} \mathbf{a}} \leq t_{n-1, 1-\frac{\alpha}{2}}^2.$$

Uma região de confiança simultânea é dada para o conjunto de valores $\mathbf{a}^\top \boldsymbol{\mu}$ tais que t^2 é relativamente pequeno para todas as escolhas de $\mathbf{a}$.

Parece razoável esperar que o valor

$$t_{n-1, 1-\frac{\alpha}{2}}^2$$

seja substituído por um valor maior, c^2 , quando afirmações são feitas para muitas escolhas de $\mathbf{a}$.

Considerando os valores de $\mathbf{a}$ para os quais $t^2 \leq c^2$, somos naturalmente levados a

$$\max_{\mathbf{a}} t^2(\mathbf{a}) = \max_{\mathbf{a}} \frac{n(\mathbf{a}^\top \bar{\mathbf{x}} - \mathbf{a}^\top \boldsymbol{\mu})^2}{\mathbf{a}^\top \mathbf{S} \mathbf{a}}.$$

Usando os resultados sobre desigualdades, percebemos que o máximo ocorrerá para

$$\mathbf{a} \propto \mathbf{S}^{-1}(\bar{\mathbf{x}} - \boldsymbol{\mu}).$$

Ora, isto nos levará à estatística T^2 .

Por conveniência, costuma se referir a esses intervalos como intervalos- T^2 .

Em particular, tomando os vetores $\mathbf{a}$'s como os vetores da base canônica do $\mathbb{R}^p$, obtém-se

$$\text{IC}_{100(1-\alpha)\%}(\mu_j) = \left(\bar{x}_j - \xi_{n,p,1-\alpha} \sqrt{\frac{s_{jj}}{n}}; \bar{x}_j + \xi_{n,p,1-\alpha} \sqrt{\frac{s_{jj}}{n}} \right), \quad j = 1, 2, \dots, p,$$

em que

$$\xi_{n,p,1-\alpha} = \sqrt{\frac{p(n-1)}{n-p} F_{p,n-p,1-\alpha}}.$$

Agora os intervalos- T^2 coletivamente têm nível de confiança $1 - \alpha$.

Observe que também podemos construir intervalos de confiança para relações estruturais entre os componentes do vetor μ como, por exemplo, intervalos para as diferenças entre os componentes de μ .

Fazendo

$$\mathbf{a}^\top = (0, \dots, 0, \underbrace{1}_{j\text{-ésima entrada}}, 0, \dots, 0, \underbrace{-1}_{k\text{-ésima entrada}}, 0, \dots, 0),$$

teremos $\mathbf{a}^\top \boldsymbol{\mu} = \mu_j - \mu_k$, $\mathbf{a}^\top \bar{\mathbf{x}} = \bar{x}_j - \bar{x}_k$ e $\mathbf{a}^\top \mathbf{S} \mathbf{a} = s_{jj} + s_{kk} - 2s_{jk}$.

O intervalo para a diferença $\mu_j - \mu_k$ será dado por

$$(\bar{x}_j - \bar{x}_k) \pm \sqrt{\frac{p(n-1)}{n-p} F_{p, n-p, 1-\alpha}} \sqrt{\frac{s_{jj} + s_{kk} - 2s_{jk}}{n}}.$$

Exemplo

Obtivemos uma elipse de 95% de confiança para o vetor μ nos dados sobre radiação em microondas. Pede-se construir os intervalos- T^2 de 95% de confiança para os componentes individuais do vetor μ , identificando-os como as “sombras” da elipse de 95% de confiança sobre os eixos coordenados. Pede-se também construir os intervalos baseados na distribuição t -Student e compará-los com os correspondentes intervalos T^2 .

Ver arquivo exemplo_09.r

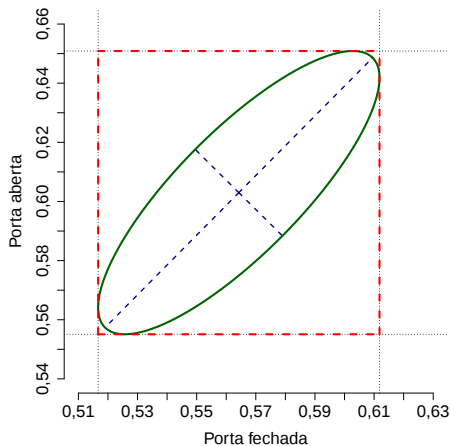


Figura 5: Intervalos T^2 (vermelho) e região (verde) de 95% de confiança para o vetor de médias, μ , das radiações de microondas de portas abertas e fechadas, respectivamente.

Uma comparação entre os intervalos- T^2 e os intervalos separados

A tabela a seguir mostra uma comparação entre os comprimentos dos intervalos de confiança separados e os intervalos “simultâneos” T^2 para alguns valores selecionados de p , n e $\alpha = 0,05$.

n	$t_{n-1}(0,975)$	$\xi_{n;p;0,95}$	
		$p = 4$	$p = 10$
15	2,145	4,14	11,52
25	2,064	3,60	6,39
50	2,010	3,31	5,05
100	1,970	3,19	4,61
∞	1,960	3,08	4,28

Os valores nas duas últimas colunas da tabela correspondem a

$$\xi_{n;p;0,95} = \sqrt{\frac{(n-1)p}{n-p} F_{p,n-p}(0,95)}.$$

A comparação feita é imprópria, pois o nível de confiança associado a qualquer coleção de intervalos- T^2 , para p fixado, é 0,95, e o nível global associado com uma coleção de intervalos separados via distribuição t -Student deve ser menor do que 0,95.

Uma outra abordagem, conhecida como **método de Bonferroni de comparações múltiplas**, será considerada. O método leva este nome devido à desigualdade de Bonferroni.

Seja $A_1, A_2, \dots, A_m$ uma coleção de eventos em um espaço de probabilidade tais que $\Pr(A_\ell) = 1 - \alpha_\ell$, $\ell = 1, \dots, m$. Então,

$$\begin{aligned}\Pr\left(\cap_{\ell=1}^m A_\ell\right) &= 1 - P([\cap_{\ell=1}^m A_\ell]^c) \\ &= 1 - P(\cup_{\ell=1}^m A_\ell^c) \\ &\geq 1 - \sum_{\ell=1}^m \Pr(A_\ell^c) \\ &= 1 - \sum_{\ell=1}^m \alpha_\ell \\ &= 1 - (\alpha_1 + \alpha_2 + \dots + \alpha_m).\end{aligned}$$

A desigualdade de Bonferroni permite ao investigador controlar a taxa de erro

$$\alpha_1 + \alpha_2 + \cdots + \alpha_m,$$

sem olhar a estrutura de correlação por trás dos intervalos de confiança.

Assim, se o problema envolve a construção de m intervalos importantes, a ideia é fazer

$$\alpha_\ell = \alpha/m, \quad \text{para } \ell = 1, \dots, m,$$

e tomar os intervalos separados dados por

$$\mathbf{a}^\top \bar{\mathbf{x}} \pm t_{n-1, 1-\frac{\alpha}{2m}} \sqrt{\frac{\mathbf{a}^\top \mathbf{S} \mathbf{a}}{n}}.$$

Observe que agora vale que o coeficiente coletivo de confiança é pelo menos

$$1 - \underbrace{\left(\frac{\alpha}{m} + \frac{\alpha}{m} + \cdots + \frac{\alpha}{m} \right)}_{m \text{ termos}} = 1 - \alpha.$$

Portanto, com um coeficiente de confiança global de pelo menos $1 - \alpha$, podemos construir os seguintes $m = p$ intervalos para os componentes do vetor μ :

$$\text{IC}_{100(1-\alpha)\%}(\mu_j) = \left(\bar{x}_j - t_{n-1, 1-\frac{\alpha}{2p}} \sqrt{\frac{s_{jj}}{n}}; \bar{x}_j + t_{n-1, 1-\frac{\alpha}{2p}} \sqrt{\frac{s_{jj}}{n}} \right), j = 1, 2, \dots, p.$$

Então, esses intervalos podem, de forma mais apropriada, ser comparados aos intervalos- T^2 .

Exemplo

Usando novamente os dados sobre radiação em fornos de microondas, pede-se comparar os intervalos- T^2 para os componentes do vetor de médias com os intervalos via Bonferroni.

Ver arquivo exemplo_10.r

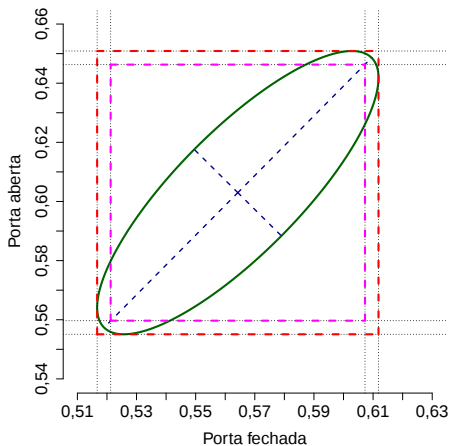


Figura 6: Intervalos T^2 (vermelho), via Bonferroni (magenta) e região (verde) de 95% de confiança para o vetor de médias, μ , das radiações de microondas de portas abertas e fechadas, respectivamente.

A tabela a seguir ilustra uma comparação entre os comprimentos dos intervalos via Bonferroni e T^2 para alguns valores selecionados de p , $m = p$, n e $\alpha = 0,05$. As entradas nas três colunas referentes aos diferentes valores de p selecionados representam a razão entre o comprimento do intervalo via Bonferroni e o comprimento do intervalo- T^2 .

n	p		
	2	4	10
15	0,88	0,69	0,29
25	0,90	0,75	0,48
50	0,91	0,78	0,58
100	0,91	0,80	0,62
∞	0,91	0,81	0,66

Podemos ver desta tabela que **os intervalos via Bonferroni produzem intervalos mais estreitos quando $m = p$.**

Devido à facilidade de aplicação e aos resultados mais eficientes em termos de estimação, geralmente **é preferível usar os intervalos simultâneos via Bonferroni.**

Inferências sobre um vetor de médias para grandes amostras

Quando o tamanho da amostra é grande, testes de hipóteses e regiões de confiança para μ podem ser construídos mesmo que a população subjacente não seja normal.

De fato, desvios fortes de uma população normal podem ser superados para tamanhos amostrais grandes.

Ambos, testes de hipóteses e intervalos de confiança simultâneos, terão níveis nominais aproximados.

As vantagens associadas com grandes amostras podem ser parcialmente compensadas por uma perda de informação amostral causada pelo uso somente das estatísticas sumário $\bar{\mathbf{X}}$ e $\mathbf{S}$.

Por outro lado, como $(\bar{\mathbf{X}}, \mathbf{S})$ é uma estatística suficiente para populações normais, quanto mais próximas da normal multivariada forem as distribuições das populações, mais eficientemente a informação amostral será utilizada ao fazer inferências.

Todas as inferências sobre μ quando se tem grandes amostras são baseadas na distribuição de qui-quadrado.

Proposição

Seja $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ uma amostra aleatória observada de uma população com vetor de médias $\boldsymbol{\mu}$ e matriz de covariâncias positiva definida $\boldsymbol{\Sigma}$. Quando $n - p$ é grande, a hipótese $H_0 : \boldsymbol{\mu} = \boldsymbol{\mu}_0$ é rejeitada em favor de $H_1 : \boldsymbol{\mu} \neq \boldsymbol{\mu}_0$, ao nível de significância α , se

$$n(\bar{\mathbf{x}} - \boldsymbol{\mu}_0)^\top \mathbf{S}^{-1}(\bar{\mathbf{x}} - \boldsymbol{\mu}_0) > \chi_{p, 1-\alpha}^2.$$

Comparando este teste com o obtido via teoria normal, no início destas notas, vemos que a estatística de teste é a mesma, o que muda é o valor crítico.

Um exame mais minucioso revela, porém, que ambos os testes produzirão os mesmos resultados em situações nas quais o teste χ^2 é apropriado (amostras realmente grandes).

De fato, $\frac{(n-1)p}{n-p} F_{p, n-p, 1-\alpha}$ e $\chi_{p, 1-\alpha}^2$ são aproximadamente iguais para $n \gg p$.

Proposição

Seja $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ uma amostra aleatória observada de uma população com vetor de médias $\boldsymbol{\mu}$ e matriz de covariâncias positiva definida $\boldsymbol{\Sigma}$. Se $n - p$ é grande,

$$\mathbf{a}^\top \bar{\mathbf{x}} \pm \sqrt{\chi_{p,1-\alpha}^2} \sqrt{\frac{\mathbf{a}^\top \mathbf{S} \mathbf{a}}{n}}$$

compreenderá $\mathbf{a}^\top \boldsymbol{\mu}$, para todo $\mathbf{a}$ com probabilidade aproximadamente $1 - \alpha$.

Consequentemente, podemos fazer afirmações simultâneas para os p componentes do vetor de médias dadas por

$$\bar{x}_j \pm \sqrt{\chi_{p,1-\alpha}^2} \sqrt{\frac{s_{jj}}{n}}, \quad j = 1, 2, \dots, p.$$

Observação: elipses de confiança para pares de componentes também podem ser facilmente construídas.

A questão de quão grande deve ser o tamanho da amostra não é simples de ser respondida.

Em uma ou duas dimensões, tamanhos amostrais em torno de 30 a 50 podem geralmente ser considerados grandes.

A medida que o número de características torna-se maior, certamente tamanhos amostrais maiores serão exigidos para que as distribuições assintóticas forneçam boas aproximações das verdadeiras distribuições das várias estatísticas de teste.

Na falta de estudos definitivos, propõe-se que $n - p$ deve ser grande, reconhecendo que o caso real pode ser mais complicado do que isso.

Uma aplicação com $p = 2$ e $n = 50$ é muito diferente de uma aplicação com $p = 52$ e $n = 100$ apesar de ambas apresentarem $n - p = 48$.

Deve-se realizar as mesmas verificações exigidas para os métodos baseados na normal.

Apesar de pequenos desvios da normalidade não causarem quaisquer dificuldades para n grande, desvios extremos podem causar problemas.

Especificamente, a taxa de erro verdadeira pode estar bem afastada do nível nominal α .

Se, com base nos *qq-plots* e outros esquemas de investigação pontos extremos (*outliers*) e outras formas de desvios extremos aparecem, ações corretivas apropriadas, incluindo transformações, são desejáveis.

Análise de variância multivariada (MANOVA)

Parte 1: Comparações emparelhadas

Começaremos com uma breve revisão deste problema no caso univariado. O problema aqui pode ser descrito da seguinte forma.

Seja $(X_1, Y_1), (X_2, Y_2), \dots, (X_n, Y_n)$ uma amostra aleatória de tamanho n de uma população bivariada.

Aqui X e Y podem representar medidas antes e depois que um certo “tratamento” foi aplicado sobre a unidade experimental, por exemplo.

O objetivo aqui é comparar se não há diferenças entre tratamentos ou se não há efeito de tratamento, quando uma das observações pode ser considerada controle.

É interessante notar que neste problema, apesar de usarmos um teste t -Student, a suposição de normalidade bivariada não é necessária.

Basta supor $D_i = (Y_i - X_i) \sim \mathcal{N}(\theta, \psi^2)$, $i = 1, 2, \dots, n$.

Sejam $E(X_i) = \mu_1$, $\text{Var}(X_i) = \sigma_1^2$, $E(Y_i) = \mu_2$, $\text{Var}(Y_i) = \sigma_2^2$ e $\text{Cov}(X_i, Y_i) = \sigma_{12}$.

Para unidades experimentais diferentes se tem $\text{Cov}(X_i, Y_r) = 0$, $i \neq r$, pois as diferentes observações são independentes.

Segue que, $\theta = E(D_i) = \mu_2 - \mu_1$ e $\psi^2 = \text{Var}(D_i) = \sigma_1^2 + \sigma_2^2 - 2\sigma_{12}$.

Como $(X_1, Y_1), (X_2, Y_2), \dots, (X_n, Y_n)$ é uma amostra aleatória, segue que $D_1, D_2, \dots, D_n$ é uma amostra aleatória de tamanho n de uma população com média θ e variância ψ^2 . Logo,

$$D_i \stackrel{\text{ind}}{\sim} \mathcal{N}(\theta, \psi^2), \quad i = 1, 2, \dots, n \quad \text{e} \quad t = \frac{\sqrt{n}(\bar{D} - \theta)}{S_D} \sim t_{n-1},$$

$$\text{com } \bar{D} = \frac{1}{n} \sum_{i=1}^n (Y_i - X_i) \quad \text{e} \quad S_D^2 = \frac{1}{n-1} \sum_{i=1}^n (D_i - \bar{D})^2.$$

Se desejamos testar as hipóteses $H_0 : \theta = 0$ versus $H_1 : \theta \neq 0$, rejeitaremos H_0 , ao nível de significância α se

$$|t| = \frac{|\bar{d}|}{S_D/\sqrt{n}} \geq t_{n-1, 1-\alpha/2}.$$

Equivalentemente, um intervalo de $100(1 - \alpha)\%$ de confiança para θ é dado por

$$\text{IC}_{100(1-\alpha)\%}(\theta) = \left(\bar{d} - t_{n-1, 1-\alpha/2} \frac{s_D}{\sqrt{n}}; \bar{d} + t_{n-1, 1-\alpha/2} \frac{s_D}{\sqrt{n}} \right).$$

O que muda no caso multivariado?

Neste caso, os elementos dos pares ordenados que representam a amostra aleatória, em vez de medidas escalares serão vetores em $\mathbb{R}^p$.

Notação adicional para a extensão multivariada é necessária.

Assim, suponha que observamos uma amostra $((\mathbf{x}_1, \mathbf{y}_1), (\mathbf{x}_2, \mathbf{y}_2), \dots, (\mathbf{x}_n, \mathbf{y}_n))$ de $(\mathbf{X}_i, \mathbf{Y}_i)$, em que $\mathbf{X}_i$ e $\mathbf{Y}_i$ são vetores aleatórios em $\mathbb{R}^p$.

Neste caso, podemos definir os vetores aleatórios diferenças

$$\mathbf{D}_i = \mathbf{Y}_i - \mathbf{X}_i, \quad i = 1, 2, \dots, n.$$

Temos que

$$\mathbf{E}(\mathbf{D}_i) = \boldsymbol{\theta} = \begin{bmatrix} \theta_1 \\ \theta_2 \\ \vdots \\ \theta_p \end{bmatrix} = \begin{bmatrix} \mu_{2,1} - \mu_{1,1} \\ \mu_{2,2} - \mu_{1,2} \\ \vdots \\ \mu_{2,p} - \mu_{1,p} \end{bmatrix},$$

em que $\theta_j = \mathbf{E}(D_{ji})$, para $j = 1, 2, \dots, p$, $\mathbf{D}_i^\top = (D_{1i}, D_{2i}, \dots, D_{pi})$, e

$$\text{Var}(\mathbf{D}_i) = \boldsymbol{\Psi}.$$

Se além disso,

$$\mathbf{D}_i \stackrel{\text{ind}}{\sim} \mathcal{N}_p(\boldsymbol{\theta}, \boldsymbol{\Psi}),$$

então, segue que

$$T^2 = n(\bar{\mathbf{D}} - \boldsymbol{\theta})^\top \mathbf{S}_D^{-1}(\bar{\mathbf{D}} - \boldsymbol{\theta}) \sim \frac{(n-1)p}{n-p} \mathcal{F}_{p, n-p},$$

em que $\bar{\mathbf{D}} = \frac{1}{n} \sum_{i=1}^n (\mathbf{Y}_i - \mathbf{X}_i)$ e $\mathbf{S}_D = \frac{1}{n-1} \sum_{i=1}^n (\mathbf{D}_i - \bar{\mathbf{D}})(\mathbf{D}_i - \bar{\mathbf{D}})^\top$.

Logo, no teste das hipóteses $H_0 : \boldsymbol{\theta} = \mathbf{0}$ versus $H_1 : \boldsymbol{\theta} \neq \mathbf{0}$, rejeitaremos H_0 , ao nível de significância α se

$$T^2 = n\bar{\mathbf{d}}^\top \mathbf{S}_D^{-1} \bar{\mathbf{d}} > \frac{(n-1)p}{n-p} F_{p, n-p, 1-\alpha}.$$

Uma região de $100(1 - \alpha)\%$ de confiança para $\boldsymbol{\theta}$ é dada pelo hiper-elipsóide

$$(\bar{\mathbf{d}} - \boldsymbol{\theta})^\top \mathbf{S}_D^{-1} (\bar{\mathbf{d}} - \boldsymbol{\theta}) \leq \frac{(n-1)p}{n(n-p)} F_{p, n-p, 1-\alpha}.$$

Intervalos- T^2 simultâneos de $100(1 - \alpha)\%$ de confiança para os componentes do vetor $\boldsymbol{\theta}$ são dados por

$$\text{IC}_{100(1-\alpha)\%}(\theta_j) = \left(\bar{d}_j - \xi_{n,p,1-\alpha} \sqrt{\frac{s_{Dj}^2}{n}}; \bar{d}_j + \xi_{n,p,1-\alpha} \sqrt{\frac{s_{Dj}^2}{n}} \right), \quad j = 1, 2, \dots, p,$$

em que

$$\xi_{n,p,1-\alpha} = \sqrt{\frac{(n-1)p}{n-p} F_{p, n-p, 1-\alpha}}.$$

Para $n - p$ grande,

$$\frac{(n-1)p}{n-p} F_{p, n-p, 1-\alpha} \simeq \chi_{p, 1-\alpha}^2$$

e a suposição de normalidade não é necessária.

Os intervalos de Bonferroni simultâneos de $100(1 - \alpha)\%$ de confiança para os componentes do vetor de médias são dados por

$$\text{IC}_{100(1-\alpha)\%}(\theta_j) = \left(\bar{d}_j - t_{n-1, 1-\frac{\alpha}{2p}} \sqrt{\frac{s_{Dj}^2}{n}}, \bar{d}_j + t_{n-1, 1-\frac{\alpha}{2p}} \sqrt{\frac{s_{Dj}^2}{n}} \right),$$

para $j = 1, 2, \dots, p$.

Exemplo

As plantas de tratamento de esgoto municipais devem monitorar suas descargas em rios e afluentes regularmente por exigência de lei.

Preocupações com a fidedignidade dos dados de um destes programas de auto-monitoramento levaram a um estudo, no qual amostras de efluentes foram divididas e enviadas a dois laboratórios para teste. Metade de cada amostra foi enviada ao Laboratório de Higiene (LH) do estado de Wisconsin, e a outra metade para um laboratório comercial privado (LC) rotineiramente usado no programa de monitoramento. Medidas do oxigênio bioquímico demandado (obd) e sólidos suspensos (ss) foram obtidas, para $n = 11$ amostras divididas, dos dois laboratórios.

- As análises dos dois laboratórios são compatíveis?
- Se há diferenças, qual é a sua natureza?

Ver arquivo exemplo_11.r

Ao nível de significância de 5%, a hipótese nula de que não há diferença no vetor de médias deve ser rejeitada.

Porém, ao construir os intervalos- T^2 simultâneos de 95% de confiança para os componentes do vetor de médias, verifica-se que o zero pertence a ambos:

$$IC_{95\%}(\xi_1) = (-22, 45; 3, 73) \quad \text{e} \quad IC_{95\%}(\xi_2) = (-5, 70; 32, 25).$$

O que concluir?

A evidência aponta para diferenças reais.

O ponto $\theta = \mathbf{0}$ está fora da região de 95% de confiança (elipse).

O coeficiente de confiança simultâneo de 95% se aplica ao conjunto completo de intervalos que poderiam ser construídos para todas as escolhas possíveis de combinações lineares $\mathbf{a}^\top \boldsymbol{\theta}$ da forma

$$a_1 \theta_1 + a_2 \theta_2.$$

Os intervalos particulares apresentados correspondem às escolhas $\mathbf{a}^\top = (a_1, a_2)^\top = (1, 0)$ e $\mathbf{a}^\top = (a_1, a_2)^\top = (0, 1)$ e contêm o zero.

Outras escolhas de a_1 e a_2 produzirão intervalos simultâneos que provavelmente não contêm o zero.

Se $H_0 : \boldsymbol{\theta} = \mathbf{0}$ não tivesse sido rejeitada, então todos os intervalos- T^2 simultâneos conteriam o zero.

Os intervalos simultâneos via Bonferroni de 95% também cobrem o zero, isto é,

$$IC_{95\%}(\xi_1) = (-20, 57; 1, 85) \quad \text{e} \quad IC_{95\%}(\xi_2) = (-2, 97; 29, 52).$$

Neste exemplo, o experimentador dividiu a amostra primeiro agitando a solução e depois colocando-a em duas garrafas para a análise química.

Esta foi uma boa estratégia, pois uma divisão simples da amostra em duas partes obtidas pela primeira metade do topo em uma garrafa e o resto na outra poderia resultar em sólidos suspensos na metade inferior devido à forma como os líquidos foram colocados nas garrafas.

Os dois laboratórios, então, não estariam necessariamente trabalhando com unidades experimentais similares, e as conclusões não pertenceriam à competência do laboratório, técnicas de medição, etc.

Sempre que um investigador puder controlar a designação de tratamentos às unidades experimentais, o emparelhamento apropriado de unidades experimentais ou a aleatorização da designação de tratamentos às unidades, permitirão uma análise estatística dos resultados mais apropriada.

Tratamentos em experimentos de medidas repetidas

Outra generalização da estatística t -Student emparelhada surge em situações nas quais q tratamentos são comparados com respeito a uma única variável resposta.

Cada unidade experimental recebe cada tratamento uma vez sobre sucessivos períodos de tempo. A i -ésima observação é

$$\mathbf{X}_i = \begin{bmatrix} X_{1i} \\ X_{2i} \\ \vdots \\ X_{qi} \end{bmatrix}, \quad i = 1, 2, \dots, n.$$

Aqui X_{ji} é a resposta do j -ésimo tratamento, $j = 1, 2, \dots, q$, sobre a i -ésima unidade experimental, $i = 1, 2, \dots, n$.

O nome medidas repetidas deriva do fato de que todos os tratamentos são administrados a cada unidade.

Para propósitos de comparação, podemos considerar os seguintes contrastes dos componentes do vetor de médias μ :

$$\begin{bmatrix} \mu_1 - \mu_2 \\ \mu_1 - \mu_3 \\ \vdots \\ \mu_1 - \mu_q \end{bmatrix} = \begin{bmatrix} 1 & -1 & 0 & \cdots & 0 \\ 1 & 0 & -1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & 0 & 0 & \cdots & -1 \end{bmatrix} \begin{bmatrix} \mu_1 \\ \mu_2 \\ \vdots \\ \mu_q \end{bmatrix} = \mathbf{C}_1 \mu$$

ou

$$\begin{bmatrix} \mu_2 - \mu_1 \\ \mu_3 - \mu_2 \\ \vdots \\ \mu_q - \mu_{q-1} \end{bmatrix} = \begin{bmatrix} -1 & 1 & 0 & \cdots & 0 & 0 \\ 0 & -1 & 1 & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \cdots & -1 & 1 \end{bmatrix} \begin{bmatrix} \mu_1 \\ \mu_2 \\ \vdots \\ \mu_q \end{bmatrix} = \mathbf{C}_2 \mu.$$

Ambas $\mathbf{C}_1$ e $\mathbf{C}_2$ são chamadas matrizes de contraste, pois suas $q - 1$ linhas são linearmente independentes e cada uma delas é um vetor de contraste.

A natureza do planejamento elimina muito da influência da variação unidade a unidade sobre as comparações de tratamento. O experimentador deve aleatorizar a ordem na qual os tratamentos são apresentados a cada unidade experimental.

Quando as médias dos diferentes tratamentos são iguais, $\mathbf{C}_1\boldsymbol{\mu} = \mathbf{C}_2\boldsymbol{\mu} = \mathbf{0}$.

Em resumo, a hipótese de que não há diferença entre tratamentos se torna $\mathbf{C}\boldsymbol{\mu} = \mathbf{0}$ para qualquer escolha adequada de $\mathbf{C}$.

Consequentemente, baseados nos contrastes observados $\mathbf{C}\mathbf{x}_i$ nas observações, teremos média amostral $\mathbf{C}\bar{\mathbf{x}}$ e matriz de covariâncias amostral $\mathbf{CSC}^\top$. Para testar

$$\begin{cases} H_0 : & \mathbf{C}\boldsymbol{\mu} = \mathbf{0} \\ H_1 : & \mathbf{C}\boldsymbol{\mu} \neq \mathbf{0} \end{cases},$$

usaremos a estatística³

$$T^2 = n(\mathbf{C}\bar{\mathbf{x}})^\top (\mathbf{CSC}^\top)^{-1} (\mathbf{C}\bar{\mathbf{x}}).$$

³ $\mathbf{C}$ não é quadrada! Logo, não simplificamos a expressão de T^2 .

Teste para a igualdade de tratamentos em experimentos de medidas repetidas

Sejam

$$\mathbf{X}_i \stackrel{ind}{\sim} \mathcal{N}_q(\boldsymbol{\mu}, \boldsymbol{\Sigma}), \quad i = 1, 2, \dots, n,$$

e $\mathbf{C}$ a matriz de contraste $(q - 1) \times q$. Para

$$\begin{cases} H_0 : & \mathbf{C}\boldsymbol{\mu} = \mathbf{0} \\ H_1 : & \mathbf{C}\boldsymbol{\mu} \neq \mathbf{0} \end{cases},$$

rejeitaremos H_0 , a um nível de significância α , se

$$T^2 = n(\mathbf{C}\bar{\mathbf{x}})^\top (\mathbf{CSC}^\top)^{-1} (\mathbf{C}\bar{\mathbf{x}}) > \frac{(n-1)(q-1)}{n-q+1} F_{q-1, n-q+1, 1-\alpha},$$

em que $\bar{\mathbf{x}} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i$ e $\mathbf{S} = \frac{1}{n-1} \sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{x}})(\mathbf{x}_i - \bar{\mathbf{x}})^\top$.

Região de $100(1 - \alpha)\%$ de confiança para $\mathbf{C}\boldsymbol{\mu}$:

$$n(\mathbf{C}\bar{\mathbf{x}} - \mathbf{C}\boldsymbol{\mu})^\top (\mathbf{C}\mathbf{S}\mathbf{C}^\top)^{-1} (\mathbf{C}\bar{\mathbf{x}} - \mathbf{C}\boldsymbol{\mu}) \leq \frac{(n-1)(q-1)}{n-q+1} F_{q-1, n-q+1, 1-\alpha}$$

Intervalos- T^2 simultâneos de $100(1 - \alpha)\%$ de confiança para contrastes isolados $\mathbf{c}^\top \boldsymbol{\mu}$:

$$\left(\mathbf{c}^\top \bar{\mathbf{x}} - \xi_{n, q-1, 1-\alpha} \sqrt{\frac{\mathbf{c}^\top \mathbf{S} \mathbf{c}}{n}}; \mathbf{c}^\top \bar{\mathbf{x}} + \xi_{n, q-1, 1-\alpha} \sqrt{\frac{\mathbf{c}^\top \mathbf{S} \mathbf{c}}{n}} \right),$$

em que

$$\xi_{n, q-1, 1-\alpha} = \sqrt{\frac{(n-1)(q-1)}{n-q+1} F_{q-1, n-q+1, 1-\alpha}}.$$

Os intervalos de Bonferroni simultâneos de $100(1 - \alpha)\%$ de confiança para contrastes isolados $\mathbf{c}^\top \boldsymbol{\mu}$:

$$\left(\mathbf{c}^\top \bar{\mathbf{x}} - t_{n-1, 1-\frac{\alpha}{2(q-1)}} \sqrt{\frac{\mathbf{c}^\top \mathbf{S} \mathbf{c}}{n}}; \mathbf{c}^\top \bar{\mathbf{x}} + t_{n-1, 1-\frac{\alpha}{2(q-1)}} \sqrt{\frac{\mathbf{c}^\top \mathbf{S} \mathbf{c}}{n}} \right).$$

Exemplo

Anestésicos melhorados são muitas vezes desenvolvidos primeiramente estudando seus efeitos em animais. Em um estudo, 19 cães receberam inicialmente a droga “pentobarbital”. Depois, cada cão recebeu dióxido de carbono (CO_2) em cada um de dois níveis de pressão (alta e baixa). Depois, halotano (H) foi adicionado e a administração de CO_2 foi repetida. A resposta, milissegundos entre batimentos cardíacos, foi medida para as quatro combinações de tratamento, a saber,

Tratamento 1: H ausente, CO_2 alta pressão;

Tratamento 2: H ausente, CO_2 baixa pressão;

Tratamento 3: H presente, CO_2 alta pressão; e

Tratamento 4: H presente, CO_2 baixa pressão.

Deseja-se testar os efeitos do anestésico quanto à pressão do CO_2 e quanto à presença de halotano.

[Ver arquivo exemplo_12.r](#)

Há três contrastes de tratamento que podem ser de interesse do experimentador.

Sejam μ_1 , μ_2 , μ_3 e μ_4 as correspondentes respostas médias para os Tratamentos 1, 2, 3 e 4. Então,

- $(\mu_3 + \mu_4) - (\mu_1 + \mu_2)$: contraste do Halotano representado pela diferença entre presença e ausência de H;
- $(\mu_1 + \mu_3) - (\mu_2 + \mu_4)$: contraste entre pressão alta e pressão baixa de CO_2 ; e
- $(\mu_1 + \mu_4) - (\mu_2 + \mu_3)$: contraste representando a interação entre pressão de CO_2 e H.

Rejeitamos H_0 porque $T^2 = 116,02$ é maior que o valor crítico $\xi_{19;3;0,95}^2 = 10,93$.

Como H_0 é rejeitada, faz sentido estudar os intervalos- T^2 simultâneos de 95% dos efeitos estudados.

Intervalo para o efeito do Halotano: (135,65 ; 282,98);

Intervalo para o efeito do CO_2 : (-114,73 ; -5,38); e

Intervalo para o efeito de interação $\text{H} \times \text{CO}_2$: (-78,73 ; 53,15).

Portanto, concluímos que a interação é desprezível, que a presença do Halotano aumenta o intervalo entre batimentos cardíacos e que na pressão baixa de CO_2 o intervalo entre batimentos cardíacos é maior.

Cuidados com estas interpretações simples dos resultados devem ser considerados. O efeito-H aparente pode ser devido à tendência do tempo. Ideal: a ordem temporal de todos os tratamentos deve ser determinada ao acaso.

O teste estudado é apropriado quando a matriz de covariâncias Σ não pode assumir qualquer estrutura especial. Se é razoável supor que ela tem uma estrutura particular, testes construídos para matrizes de covariâncias com estruturas especiais têm maior poder do que o teste apresentado aqui.

Análise de variância multivariada (MANOVA)

Parte 2: Comparação de vetores de médias para duas populações - amostras independentes

A estatística T^2 também é apropriada para comparar respostas de um conjunto de observações experimentais (**população 1**) com outro conjunto independente de observações experimentais (**população 2**). Isto pode ser feito sem explicitamente controlar a variabilidade unidade a unidade, como no caso das comparações emparelhadas.

Se possível, as unidades experimentais devem ser aleatoriamente designadas aos conjuntos de condições experimentais. A aleatorização irá, em alguma extensão, abrandar a variabilidade unidade a unidade na comparação subsequente de tratamentos. Apesar da perda em precisão relativa às comparações emparelhadas, as inferências no caso de duas populações são, ordinariamente, aplicáveis a uma coleção mais geral de unidades experimentais simplesmente porque a homogeneidade unitária não é mais exigida.

Considere uma amostra aleatória de tamanho n_1 da **população 1** e uma amostra aleatória de tamanho n_2 , independente da primeira amostra, da **população 2**.

Deseja-se fazer inferência sobre os vetores de médias das duas populações, por exemplo,

$$\begin{cases} H_0 : & \mu_1 - \mu_2 = \mathbf{0} \\ H_1 : & \mu_1 - \mu_2 \neq \mathbf{0} \end{cases}$$

Suposições básicas:

- (a) $(\mathbf{X}_{11}, \mathbf{X}_{12}, \dots, \mathbf{X}_{1n_1})$ é uma amostra aleatória de uma população p -variada com vetor de médias μ_1 e matriz de covariâncias Σ_1 .
- (b) $(\mathbf{X}_{21}, \mathbf{X}_{22}, \dots, \mathbf{X}_{2n_2})$ é uma amostra aleatória de uma população p -variada com vetor de médias μ_2 e matriz de covariâncias Σ_2 .
- (c) As duas amostras são independentes.

Suposições adicionais (tamanhos amostrais pequenos):

- (d) Ambas as populações são normais multivariadas.
- (e) $\Sigma_1 = \Sigma_2$.

A suposição de igualdade das covariâncias é bem mais forte do que a sua contrapartida univariada. Aqui estamos assumindo que vários pares de covariâncias são aproximadamente iguais.

Quando $\Sigma_1 = \Sigma_2 = \Sigma$, temos que $\mathbf{S}_1$ e $\mathbf{S}_2$, sob a suposição de normalidade, têm distribuições independentes de Wishart com $n_1 - 1$ e $n_2 - 1$ graus de liberdade, respectivamente, e parâmetro de escala Σ , em que $E((n_\ell - 1)\mathbf{S}_\ell) = (n_\ell - 1)\Sigma$, para $\ell = 1, 2$.

Assim, **podemos definir um estimador combinado para Σ dado por**

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

Para testar as hipóteses

$$\begin{cases} H_0 : \mu_1 - \mu_2 = \delta_0 \\ H_1 : \mu_1 - \mu_2 \neq \delta_0, \end{cases}$$

usaremos a estatística T^2 correspondente à diferença entre as médias amostrais, a saber,

$$T^2 = \left(\bar{\mathbf{X}}_1 - \bar{\mathbf{X}}_2 - \delta_0 \right)^\top \left[\left(\frac{1}{n_1} + \frac{1}{n_2} \right) \mathbf{S}_c \right]^{-1} \left(\bar{\mathbf{X}}_1 - \bar{\mathbf{X}}_2 - \delta_0 \right).$$

Além disso,

$$T^2 \sim \frac{(n_1 + n_2 - 2)p}{n_1 + n_2 - p - 1} \mathcal{F}_{p, n_1 + n_2 - p - 1}.$$

Rejeitamos a H_0 se $T_{obs}^2 > c^2$ em que

$$c^2 = \frac{(n_1 + n_2 - 2)p}{n_1 + n_2 - p - 1} F_{p, n_1 + n_2 - p - 1, 1 - \alpha}.$$

Intervalos- T^2 simultâneos

Faça

$$\xi_{n_1+n_2-1,p,1-\alpha} = \sqrt{\left[\frac{(n_1 + n_2 - 2)p}{n_1 + n_2 - p - 1} \right] F_{p,n_1+n_2-p-1,1-\alpha}}.$$

Então, os intervalos simultâneos de $100(1 - \alpha)\%$ de confiança para $\mathbf{a}^\top(\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)$ são dados por

$$\mathbf{a}^\top(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2) \pm \xi_{n_1+n_2-1,p,1-\alpha} \sqrt{\mathbf{a}^\top \left(\frac{1}{n_1} + \frac{1}{n_2} \right) \mathbf{S}_c \mathbf{a}}$$

Em particular, os intervalos de $100(1 - \alpha)\%$ de confiança para $\mu_{1j} - \mu_{2j}$ serão dados por

$$\bar{x}_{1j} - \bar{x}_{2j} \pm \xi_{n_1+n_2-1,p,1-\alpha} \sqrt{\left(\frac{1}{n_1} + \frac{1}{n_2} \right) s_{jj,c}}, \quad j = 1, 2, \dots, p.$$

A situação na qual $\Sigma_1 \neq \Sigma_2$

Quando isto ocorre, não somos capazes de obter uma medida de “distância” T^2 , cuja distribuição não dependa das matrizes de covariâncias desconhecidas.

Sugere-se que qualquer discrepância da ordem de $\sigma_{1,ii} = 4\sigma_{2,ii}$, ou vice-versa, é provavelmente séria. Uma transformação pode melhorar a situação quando as variâncias marginais são bem diferentes.

Porém, para n_1 e n_2 grandes, podemos evitar complexidades devido a uma desigualdade das covariâncias.

Proposição

Suponha que $n_1 - p$ e $n_2 - p$ sejam grandes. Então, um hiper-elipsóide aproximado de $100(1 - \alpha)\%$ de confiança para $\mu_1 - \mu_2$ é dado para todo $\mu_1 - \mu_2$ que satisfaça

$$[\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2 - (\mu_1 - \mu_2)]^\top \left[\frac{1}{n_1} \mathbf{S}_1 + \frac{1}{n_2} \mathbf{S}_2 \right]^{-1} [\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2 - (\mu_1 - \mu_2)] \leq \chi_{p,1-\alpha}^2.$$

Também, intervalos de $100(1 - \alpha)\%$ de confiança para as combinações $\mathbf{a}^\top (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)$ são dados por

$$\mathbf{a}^\top (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2) \pm \sqrt{\chi_{p,1-\alpha}^2} \sqrt{\mathbf{a}^\top \left[\frac{1}{n_1} \mathbf{S}_1 + \frac{1}{n_2} \mathbf{S}_2 \right] \mathbf{a}}.$$

Uma aproximação da distribuição de T^2 para populações normais, quando os tamanhos amostrais não são grandes

Pode-se testar $H_0 : \boldsymbol{\mu}_1 - \boldsymbol{\mu}_2 = \mathbf{0}$ quando as matrizes de covariâncias são desiguais, mesmo para tamanhos amostrais moderados, desde que as populações sejam normais multivariadas. Esta situação costuma ser chamada de problema de Behrens-Fisher. O resultado exige que $n_1 > p$ e $n_2 > p$.

A abordagem depende de uma aproximação da distribuição da estatística

$$T^2 = \left[\bar{\mathbf{X}}_1 - \bar{\mathbf{X}}_2 - (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2) \right]^\top \left[\frac{1}{n_1} \mathbf{S}_1 + \frac{1}{n_2} \mathbf{S}_2 \right]^{-1} \left[\bar{\mathbf{X}}_1 - \bar{\mathbf{X}}_2 - (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2) \right].$$

Observe que é a mesma estatística usada no resultado anterior. No entanto, em vez de usar a distribuição aproximada de qui-quadrado para obter os valores críticos para testar H_0 , a aproximação recomendada para amostras moderadas a pequenas é dada por

$$T^2 \stackrel{a}{\sim} \frac{\nu p}{\nu - p + 1} \mathcal{F}_{p, \nu - p + 1},$$

com

$$\nu = \frac{p + p^2}{\sum_{i=1}^2 \frac{1}{n_i} \left\{ \text{tr} \left[\left(\frac{1}{n_i} \mathbf{S}_i \left(\frac{1}{n_1} \mathbf{S}_1 + \frac{1}{n_2} \mathbf{S}_2 \right)^{-1} \right)^2 \right] + \left(\text{tr} \left[\frac{1}{n_i} \mathbf{S}_i \left(\frac{1}{n_1} \mathbf{S}_1 + \frac{1}{n_2} \mathbf{S}_2 \right)^{-1} \right] \right)^2 \right\}}.$$

em que $\min\{n_1, n_2\} \leq \nu \leq n_1 + n_2$.

Esta aproximação reduz-se à solução de Welch para o problema de Behrens-Fisher univariado ($p = 1$).

MANOVA a um fator

Dados

Amostra 1:	$\mathbf{X}_{11}, \mathbf{X}_{12}, \dots, \mathbf{X}_{1n_1}$	da pop. 1
Amostra 2:	$\mathbf{X}_{21}, \mathbf{X}_{22}, \dots, \mathbf{X}_{2n_2}$	da pop. 2
$\vdots$	$\vdots$	
Amostra g :	$\mathbf{X}_{g1}, \mathbf{X}_{g2}, \dots, \mathbf{X}_{gn_g}$	da pop. g

Suposições:

- $(\mathbf{X}_{k1}, \mathbf{X}_{k2}, \dots, \mathbf{X}_{kn_k})$ é uma amostra de tamanho n_k de uma população com vetor de médias $\boldsymbol{\mu}_k$, $k = 1, 2, \dots, g$;
- todas as populações têm matrizes de covariâncias comum $\boldsymbol{\Sigma}$;
- cada população é normal p -variada; e
- as amostras são independentes.

Primeiro, faremos uma rápida revisão do modelo ANOVA a um fator univariado.

Neste caso, tem-se g amostras aleatórias independentes $X_{k1}, X_{k2}, \dots, X_{kn_k}$ de tamanhos n_k de distribuições $\mathcal{N}(\mu_k, \sigma^2)$, $k = 1, 2, \dots, g$.

As hipóteses de interesse são

$$\begin{cases} H_0 : & \mu_1 = \mu_2 = \dots = \mu_g \\ H_1 : & \text{pelo menos uma média é diferente das demais.} \end{cases}$$

O modelo ANOVA é definido por

$$X_{ki} = \underbrace{\mu_k}_{\text{média da } k\text{-ésima amostra}} + \underbrace{\varepsilon_{ki}}_{\text{erro aleatório } i\text{-ésima observação da } k\text{-ésima amostra}},$$

para $i = 1, 2, \dots, n_k$, e $k = 1, 2, \dots, g$.

Em geral, adota-se a reparametrização

$$\mu_k = \underbrace{\mu}_{\text{média global}} + \underbrace{\tau_k}_{\text{efeito do } k\text{-ésimo tratamento}}$$

e, assim, as hipóteses equivalentes são

$$\begin{cases} H_0 : \tau_1 = \tau_2 = \cdots = \tau_g = 0 \\ H_1 : \text{ pelo menos um deles é não nulo.} \end{cases}$$

Assim, o modelo é

$$X_{ki} = \mu + \tau_k + \varepsilon_{ki}, \quad i = 1, 2, \dots, n_k, \quad k = 1, 2, \dots, g.$$

Suposição: $\varepsilon_{ki} \stackrel{\text{ind}}{\sim} \mathcal{N}(0, \sigma^2)$.

Para definir de modo único os parâmetros do modelo e suas estimativas de mínimos quadrados é usual impor a restrição $\sum_{k=1}^g n_k \tau_k = 0$.

Se as amostras são balanceadas, isto é, se $n_1 = n_2 = \dots = n_g$, então a restrição se simplifica para $\sum_{k=1}^g \tau_k = 0$.

O procedimento básico adotado é a decomposição da soma de quadrados totais:

$$\begin{aligned}
 \underbrace{\sum_{k=1}^g \sum_{i=1}^{n_k} (X_{ki} - \bar{X}_{..})^2}_{\text{SQ total}} &= \sum_{k=1}^g \sum_{i=1}^{n_k} (X_{ki} - \bar{X}_{k.} + \bar{X}_{k.} - \bar{X}_{..})^2 \\
 &= \underbrace{\sum_{k=1}^g n_k (\bar{X}_{k.} - \bar{X}_{..})^2}_{\text{SQ tratamentos}} \\
 &+ \underbrace{\sum_{k=1}^g \sum_{i=1}^{n_k} (X_{ki} - \bar{X}_{k.})^2}_{\text{SQ devida ao erro aleatório}}.
 \end{aligned}$$

Os estimadores de mínimos quadrados de μ e τ_k são dados por

$$\hat{\mu} = \bar{X}_{..} = \frac{1}{n} \sum_{k=1}^g \sum_{i=1}^{n_k} X_{ki}, \quad \text{com } n = \sum_{k=1}^g n_k,$$

$$\hat{\tau}_k = \bar{X}_{k.} - \bar{X}_{..}, \quad \text{com } \bar{X}_{k.} = \frac{1}{n_k} \sum_{i=1}^{n_k} X_{ki}.$$

Um estimador não tendencioso para σ^2 é dado por

$$\tilde{\sigma}^2 = \frac{SQ_{Erro}}{n - g} = \frac{\sum_{k=1}^g \sum_{i=1}^{n_k} (X_{ki} - \bar{X}_{k.})^2}{n - g}, \quad \text{com } n = \sum_{k=1}^g n_k.$$

Tabela 1: ANOVA a um fator com g níveis.

Fonte de Variação	g.l.	Soma de Quadrados	Quadrado Médio	Razão F
Tratamento	$g - 1$	$\sum_{k=1}^g n_k (\bar{x}_{k.} - \bar{x}_{..})^2$	$SQ_{Trat} / (g - 1)$	$\frac{QM_{Trat}}{QM_{Erro}}$
Erro	$n - g$	$\sum_{k=1}^g \sum_{i=1}^{n_k} (x_{ki} - \bar{x}_{k.})^2$	$SQ_{Erro} / (n - g)$	
Total	$n - 1$	$\sum_{k=1}^g \sum_{i=1}^{n_k} (x_{ki} - \bar{x}_{..})^2$		

em que $n = \sum_{k=1}^g n_k$.

Se H_0 é verdadeira, $F \sim \mathcal{F}_{g-1, n-g}$.

Logo, rejeitaremos H_0 ao nível de significância α se $F > F_{g-1, n-g, 1-\alpha}$.

Observe que isto é equivalente a dizer que o teste rejeita H_0 para valores grandes da razão SQ_{Trat}/SQ_{Erro} , equivalentemente, valores grandes de $1 + SQ_{Trat}/SQ_{Erro}$ ou, equivalentemente, valores pequenos da recíproca, dada por

$$\frac{SQ_{Erro}}{SQ_{Trat} + SQ_{Erro}}.$$

Adiante, veremos que uma estatística de teste no caso multivariado, com suas devidas adaptações, tem forma similar a esta última.

MANOVA a um fator (continuação)

Modelo:

$$\mathbf{X}_{ki} = \boldsymbol{\mu} + \boldsymbol{\tau}_k + \boldsymbol{\varepsilon}_{ki}, \quad k = 1, 2, \dots, g, \quad i = 1, 2, \dots, n_k.$$

Suposição:

$$\boldsymbol{\varepsilon}_{ki} \sim \mathcal{N}_p(\mathbf{0}, \boldsymbol{\Sigma}).$$

Restrição:

$$\sum_{k=1}^g n_k \boldsymbol{\tau}_k = \mathbf{0}$$

Soma de Quadrados Total:

$$SQ_{Tot} = \sum_{k=1}^g \sum_{i=1}^{n_k} (\mathbf{x}_{ki} - \bar{\mathbf{x}}_{..})(\mathbf{x}_{ki} - \bar{\mathbf{x}}_{..})^{\top}.$$

O vetor de médias amostral é dado por

$$\bar{\mathbf{x}}_{..} = \frac{1}{n} \sum_{k=1}^g \sum_{i=1}^{n_k} \mathbf{x}_{ki} \quad \text{com} \quad n = \sum_{k=1}^g n_k.$$

Agora, SQ_{Tot} é uma matriz $p \times p$.

Os elementos de sua diagonal correspondem às somas de quadrados totais para cada medida isolada.

Fora da diagonal temos as somas dos produtos cruzados dos desvios, caracterizando a estrutura de dependência entre as p medidas estudadas.

Decomposição de SQ_{Tot} :

$$\begin{aligned} SQ_{Tot} &= \sum_{k=1}^g \sum_{i=1}^{n_k} (\mathbf{x}_{ki} - \bar{\mathbf{x}}_{..})(\mathbf{x}_{ki} - \bar{\mathbf{x}}_{..})^\top \\ &= \underbrace{\sum_{k=1}^g n_k (\bar{\mathbf{x}}_{k.} - \bar{\mathbf{x}}_{..})(\bar{\mathbf{x}}_{k.} - \bar{\mathbf{x}}_{..})^\top}_{SQ_{Trat}} \\ &\quad + \underbrace{\sum_{k=1}^g \sum_{i=1}^{n_k} (\mathbf{x}_{ki} - \bar{\mathbf{x}}_{k.})(\mathbf{x}_{ki} - \bar{\mathbf{x}}_{k.})^\top}_{SQ_{Res}} \\ &= SQ_{Trat} + SQ_{Res} = \mathbf{B} + \mathbf{W}. \end{aligned}$$

No caso univariado, a decomposição, válidas as suposições, fornece estatísticas independentes com distribuições de qui-quadrado, parâmetro de escala σ^2 .

No caso multivariado, o resultado é similar. A decomposição fornece estatísticas independentes com distribuição Wishart, parâmetro de escala Σ .

Observe que quanto ao número de graus de liberdade, eles são exatamente os mesmos do caso univariado.

Hipóteses:

$$\begin{cases} H_0 : \tau_1 = \tau_2 = \cdots = \tau_g = \mathbf{0} \\ H_1 : \text{pelo menos um deles é não nulo.} \end{cases}$$

Tabela 2: Tabela MANOVA a um fator com g níveis.

Fonte de Variação	g.l.	Matriz de Soma de Quadrados
Tratamento	$g - 1$	$\mathbf{B} = \sum_{k=1}^g n_k (\bar{\mathbf{X}}_{k.} - \bar{\mathbf{X}}_{..})(\bar{\mathbf{X}}_{k.} - \bar{\mathbf{X}}_{..})^\top$
Erro	$n - g$	$\mathbf{W} = \sum_{k=1}^g \sum_{i=1}^{n_k} (\mathbf{X}_{ki} - \bar{\mathbf{X}}_{k.})(\mathbf{X}_{ki} - \bar{\mathbf{X}}_{k.})^\top$
Total	$n - 1$	$\mathbf{B} + \mathbf{W} = \sum_{k=1}^g \sum_{i=1}^{n_k} (\mathbf{X}_{ki} - \bar{\mathbf{X}}_{..})(\mathbf{X}_{ki} - \bar{\mathbf{X}}_{..})^\top$

em que $n = \sum_{k=1}^g n_k$.

As diagonais principais de cada uma destas matrizes de somas de quadrados e produtos cruzados fornecem as somas de quadrados da ANOVA para cada medida separadamente.

Uma possível estatística de teste envolve variâncias generalizadas. Seja

$$\Lambda^* = \frac{|\mathbf{W}|}{|\mathbf{B} + \mathbf{W}|}.$$

Rejeitaremos H_0 se Λ^* for um valor pequeno.

A estatística Λ^* foi originalmente proposta por Wilks e corresponde a uma forma equivalente do teste $\mathcal{F}$ da hipótese de ausência de efeito de tratamento do caso univariado.

A estatística lambda de Wilks tem a vantagem de ser conveniente e estar relacionada ao critério da razão de verossimilhança.

A distribuição exata de Λ^* pode ser derivada para os casos especiais apresentados a seguir.

Casos especiais:

1. $p = 1, g \geq 2, \left(\frac{n-g}{g-1} \right) \left(\frac{1-\Lambda^*}{\Lambda^*} \right) \sim \mathcal{F}_{g-1, n-g}.$
2. $p = 2, g \geq 2, \left(\frac{n-g-1}{g-1} \right) \left(\frac{1-\sqrt{\Lambda^*}}{\sqrt{\Lambda^*}} \right) \sim \mathcal{F}_{2(g-1), 2(n-g-1)}.$
3. $p \geq 1, g = 2, \left(\frac{n-p-1}{p} \right) \left(\frac{1-\Lambda^*}{\Lambda^*} \right) \sim \mathcal{F}_{p, n-p-1}.$
4. $p \geq 1, g = 3, \left(\frac{n-p-2}{p} \right) \left(\frac{1-\sqrt{\Lambda^*}}{\sqrt{\Lambda^*}} \right) \sim \mathcal{F}_{2p, 2(n-p-2)}.$

Para outros casos e tamanhos amostrais grandes, uma modificação de Λ^* devido à Bartlett pode ser usada para testar H_0 .

Bartlett mostrou que se H_0 é verdadeira e n é grande, então

$$- \left(n - 1 - \frac{(p+g)}{2} \right) \ln \Lambda^* \stackrel{a}{\sim} \chi^2_{p(g-1)}$$

Consequentemente, se n é grande, rejeitamos H_0 ao nível de significância α , se

$$- \left(n - 1 - \frac{(p+g)}{2} \right) \ln \Lambda^* > \chi^2_{p(g-1), 1-\alpha}.$$

Exemplo

O Departamento de Saúde e Serviços Sociais de Wisconsin reembolsa casas de repouso para idosos no estado por serviços fornecidos. O departamento desenvolve um conjunto de fórmulas para as taxas de cada facilidade, baseado em fatores tais como nível de cuidados, taxa média de salários, e taxa média de salário no estado.

As casas de repouso podem ser classificadas com relação à propriedade - privadas, sem fins lucrativos e públicas - e também com relação à certificação - especializadas em enfermagem, unidades de cuidados intermediários, ou uma combinação dos dois.

Um objetivo de um estudo foi investigar os efeitos de propriedade e certificação (ou ambos) sobre os custos. Quatro custos, calculados por paciente-dia, medidos em horas-paciente, foram usados:

X_1 : custo de mão de obra de enfermagem;

X_2 : custo de nutricionista;

X_3 : custo de trabalho de manutenção e funcionamento; e

X_4 : custo de limpeza e lavanderia.

Um total de $n = 516$ observações sobre cada um dos $p = 4$ custos foram inicialmente separados pelo tipo de propriedade.

Tabela 3: Estatísticas sumário para cada um dos $g = 3$ grupos.

Grupo k	n_k	Vetor de médias amostral
$k = 1$, privado	$n_1 = 271$	$\bar{\mathbf{x}}_1^T = (2,066; 0,480; 0,082; 0,360)$
$k = 2$, não-lucrativo	$n_2 = 138$	$\bar{\mathbf{x}}_2^T = (2,167; 0,596; 0,124; 0,418)$
$k = 3$, público	$n_3 = 107$	$\bar{\mathbf{x}}_3^T = (2,273; 0,521; 0,125; 0,383)$

$$\mathbf{S}_1 = \begin{bmatrix} 0,291 & -0,001 & 0,002 & 0,010 \\ -0,001 & 0,011 & 0,000 & 0,003 \\ 0,002 & 0,000 & 0,001 & 0,000 \\ 0,010 & 0,003 & 0,000 & 0,010 \end{bmatrix},$$

$$\mathbf{S}_2 = \begin{bmatrix} 0,561 & 0,011 & 0,001 & 0,037 \\ 0,011 & 0,025 & 0,004 & 0,007 \\ 0,001 & 0,004 & 0,005 & 0,002 \\ 0,037 & 0,007 & 0,002 & 0,019 \end{bmatrix} \text{ e}$$

$$\mathbf{S}_3 = \begin{bmatrix} 0,261 & 0,030 & 0,003 & 0,018 \\ 0,030 & 0,017 & 0,000 & 0,006 \\ 0,003 & 0,000 & 0,004 & 0,001 \\ 0,018 & 0,006 & 0,001 & 0,013 \end{bmatrix}.$$

Como os $\mathbf{S}_k$'s parecem similares, é razoável supor matriz de covariâncias iguais.

Calculando as matrizes da tabela MANOVA obtemos

$$\mathbf{W} = \begin{bmatrix} 182,962 & 4,408 & 1,695 & 9,581 \\ 4,408 & 8,200 & 0,633 & 2,428 \\ 1,695 & 0,633 & 1,484 & 0,394 \\ 9,581 & 2,428 & 0,394 & 6,538 \end{bmatrix},$$

com 513 graus de liberdade e

$$\mathbf{B} = \begin{bmatrix} 3,475 & 1,111 & 0,821 & 0,584 \\ 1,111 & 1,225 & 0,453 & 0,610 \\ 0,821 & 0,453 & 0,235 & 0,230 \\ 0,584 & 0,610 & 0,230 & 0,304 \end{bmatrix},$$

com 2 graus de liberdade.

Neste caso temos $\Lambda^* = \frac{|\mathbf{W}|}{|\mathbf{B} + \mathbf{W}|} \simeq 0,771$.

Como $p = 4$ e $g = 3$, observe que estamos sob o caso especial 4:

$$p \geq 1, \quad g = 3, \quad \left(\frac{n-p-2}{p} \right) \left(\frac{1 - \sqrt{\Lambda^*}}{\sqrt{\Lambda^*}} \right) \sim \mathcal{F}_{2p, 2(n-p-2)}.$$

Temos que $\left(\frac{n-p-2}{p} \right) \left(\frac{1 - \sqrt{\Lambda^*}}{\sqrt{\Lambda^*}} \right) \simeq 17,67$.

Ao nível de significância de 1% temos $F_{8; 1000; 0,99} \simeq \chi_{8; 0,99}^2 / 8 = 2,51$.

Logo, rejeitamos H_0 , a hipótese de que não há diferença no custo médio entre os diferentes tipos de propriedades.

Vale comparar este resultado com o procedimento sugerido por Bartlett.

Temos que $-\left(n-1-\frac{(p+g)}{2}\right) \ln \Lambda^* \simeq 132,76$.

Ao nível de significância de 1% tem-se $\chi^2_{8;0,99} = 20,09$, o que novamente leva à rejeição de H_0 .

Este resultado é consistente com o resultado baseado na estatística $\mathcal{F}$.

Intervalos simultâneos de confiança para efeitos de tratamento

Quando H_0 é rejeitada, os efeitos que levaram à rejeição são de interesse. Para comparações pareadas, a abordagem de Bonferroni pode ser usada para construir intervalos simultâneos de confiança para os componentes dos vetores diferença $\tau_k - \tau_\ell$.

Seja τ_{kj} o j -ésimo componente de τ_k .

Como τ_k é estimado por $\hat{\tau}_k = \bar{\mathbf{X}}_k - \bar{\mathbf{X}}_{..}$, temos que o j -ésimo componente de τ_k é estimado por $\bar{X}_{k.(j)} - \bar{X}_{..(j)}$.

Assim, $\tau_{kj} - \tau_{\ell j}$ é estimado por

$$\hat{\tau}_{kj} - \hat{\tau}_{\ell j} = \bar{X}_{k.(j)} - \bar{X}_{\ell.(j)}$$

que é uma diferença entre médias de duas amostras independentes.

O teste t -Student para duas amostras independentes é válido para um nível de significância apropriado.

$$\text{Var}(\hat{\tau}_{kj} - \hat{\tau}_{\ell j}) = \left(\frac{1}{n_k} + \frac{1}{n_\ell} \right) \sigma_{jj}$$

com σ_{jj} o j -ésimo elemento da diagonal de Σ .

Como sugerido, σ_{jj} é estimado pelo j -ésimo elemento da diagonal da matriz $\mathbf{W}$ dividido pelos respectivos graus de liberdade, a saber, $\hat{\sigma}_{jj} = \frac{1}{n-g} w_{jj}$.

Resta identificar o número de intervalos. Como são g tratamentos para p medidas, segue que ao todo teremos $m = p \binom{g}{2} = pg(g-1)/2$ intervalos.

Logo, com confiança pelo menos $1 - \alpha$ os intervalos simultâneos para $\tau_{kj} - \tau_{\ell j}$ são dados por

$$(\bar{X}_{k.(j)} - \bar{X}_{\ell.(j)}) \pm t_{n-g, 1-\alpha/(pg(g-1))} \sqrt{\left(\frac{1}{n_k} + \frac{1}{n_\ell} \right) \frac{w_{jj}}{n-g}}.$$

Exemplo

Vimos no exemplo anterior - referente aos dados sobre casas de repouso para idosos de Wisconsin - que há diferenças de custos conforme o tipo de propriedade do estabelecimentos: privado, não lucrativo e público. Podemos, então calcular os intervalos simultâneos para estimar as magnitudes das diferenças. Temos que

$$\hat{\tau}_1^T = (-0,070; -0,039; -0,020; -0,020)$$

$$\hat{\tau}_2^T = (0,031; 0,076; 0,022; 0,038)$$

$$\hat{\tau}_3^T = (0,137; 0,002; 0,023; 0,003)$$

e

$$W = \begin{bmatrix} 182,962 & 4,408 & 1,695 & 9,581 \\ 4,408 & 8,200 & 0,633 & 2,428 \\ 1,695 & 0,633 & 1,484 & 0,394 \\ 9,581 & 2,428 & 0,394 & 6,538 \end{bmatrix},$$

com 513 graus de liberdade.

Tabela 4: Intervalos simultâneos de pelo menos 95% de confiança.

Grupos	Comp.	Dif. médias	Limite Inf.	Limite Sup.
1,2	1	-0,10	-0,28	0,08
	2	-0,12	-0,15	-0,08
	3	-0,04	-0,06	-0,01
	4	-0,06	-0,09	-0,02
1,3	1	-0,21	-0,40	-0,01
	2	-0,04	-0,08	0,00
	3	-0,04	-0,06	-0,03
	4	-0,02	-0,06	0,01
2,3	1	-0,11	-0,33	0,12
	2	0,08	0,03	0,12
	3	0,00	-0,02	0,02
	4	0,04	-0,01	0,08

Foram destacados os intervalos que excluíram o zero.

Com relação ao custo mão de obra de enfermagem: as casas de repouso particulares parecem ter um custo mais baixo do que as públicas.

Com relação ao custo de nutricionista: as casas particulares parecem ter custo menor do que as sem fins lucrativos e, as sem fins lucrativos mostram um custo maior do que as públicas.

Com relação ao custo de manutenção e operação: as particulares mostram ter um custo menor tanto em comparação com as sem fins lucrativos, como com as públicas.

Com relação ao custo de limpeza e lavanderia: as particulares parecem ter custo menor do que as sem fins lucrativos.

Teste para a igualdade das matrizes de covariâncias

Uma das suposições feitas no modelo MANOVA foi a de que todas as populações têm variâncias iguais. Portanto, deve-se verificar se esta suposição é razoável ou não.

Um teste comumente usado é o teste M de Box.

Com g populações, tem-se

$$\begin{cases} H_0 : \mathbf{\Sigma}_1 = \mathbf{\Sigma}_2 = \cdots = \mathbf{\Sigma}_g = \mathbf{\Sigma} \\ H_1 : \text{pelo menos duas matrizes de covariâncias não são iguais.} \end{cases}$$

Supondo populações normais multivariadas, uma estatística de razão de verossimilhança é dada por

$$\Lambda = \prod_{k=1}^g \left(\frac{|\mathbf{s}_k|}{|\mathbf{s}_c|} \right)^{(n_k-1)/2} \quad \text{com} \quad \mathbf{s}_c = \frac{1}{n-g} \sum_{k=1}^g (n_k-1) \mathbf{s}_k.$$

O teste de Box é baseado na aproximação de qui-quadrado da distribuição amostral de $-2 \ln \Lambda$.

Faça $M = -2 \ln \Lambda$, então

$$M = (n - g) \ln |\mathbf{S}_c| - \sum_{k=1}^g [(n_k - 1) \ln |\mathbf{S}_k|].$$

Se H_0 é verdadeira, as matrizes de covariâncias devem ser bem parecidas e, conseqüentemente, também serão similares aos elementos em $\mathbf{S}_c$. Nesse caso, Λ estará próximo de 1 e, M , próximo de 0.

Teste M de Box para a igualdade de matrizes de covariância

Faça

$$u = \left[\sum_{k=1}^g \frac{1}{n_k - 1} - \frac{1}{\sum_{k=1}^g (n_k - 1)} \right] \left[\frac{2p^2 + 3p - 1}{6(p+1)(g-1)} \right]$$

em que p é o número de variáveis e g é o número de populações. Então,

$$C = (1 - u)M \stackrel{a}{\sim} \chi^2_{\nu}, \quad \text{com} \quad \nu = g \frac{1}{2} p(p+1) - \frac{1}{2} p(p+1) = \frac{1}{2} p(p+1)(g-1).$$

Logo, ao nível de significância α rejeitamos H_0 se

$$C > \chi^2_{\nu, 1-\alpha}.$$

A aproximação de χ^2 de Box funciona melhor se cada $n_k > 20$ e se p e g não são superiores a 5.

O teste M é sensível a algumas formas de não-normalidade. Porém, com amostras razoavelmente grandes, os testes MANOVA para efeitos de tratamento são mais robustos à não-normalidade, isto é, podemos prosseguir com a MANOVA usual, mesmo que o teste M rejeite H_0 .

Perspectivas e estratégia na análise de modelos multivariados

Quando se tem várias características, é importante controlar a probabilidade global de tomar uma decisão incorreta. Isto é particularmente importante quando testamos a igualdade entre dois ou mais tratamentos.

Um único teste multivariado com seu p -valor associado é preferível a realizar um grande número de testes univariados.

A saída do teste multivariado nos dirá se vale ou não a pena olhar de forma mais minuciosa variável por variável e analisar grupo por grupo.

Um único teste multivariado é recomendável em vez de, digamos, p testes univariados, pois como o próximo exemplo demonstrará, testes univariados ignoram informações importantes e podem levar a conclusões erradas.

Exemplo

Suponha que tenhamos coletado medidas sobre duas variáveis X_1 e X_2 para 10 unidades experimentais, selecionadas ao acaso de cada uma de duas populações (dois grupos). Os dados, hipotéticos são listados a seguir.

Grupo 1		Grupo 2	
X_1	X_2	X_1	X_2
5,0	3,0	4,6	4,9
4,5	3,2	4,9	5,9
6,0	3,5	4,0	4,1
6,0	4,6	3,8	5,4
6,2	5,6	6,2	6,1
6,9	5,2	5,0	7,0
6,8	6,0	5,3	4,7
5,3	5,5	7,1	6,6
6,6	7,3	5,8	7,8
7,3	6,5	6,8	8,0

É possível ver pelos diagramas de ponto de cada variável e o diagrama de dispersão das duas, que há boa sobreposição para os dois grupos em questão.

O teste de comparação dos dois vetores de médias, ao nível de significância de 10%, resulta em $T^2 = 17,29$ com valor crítico ($\alpha = 0,1$) dado por 12,94, tal que a hipótese nula, de igualdade entre os dois vetores de médias deve ser rejeitada. O valor- p deste teste é aproximadamente 0,0033.

Se realizarmos as ANOVAS separadamente para cada variável, teremos para a variável X_1 a razão $F = 2,46$ com valor- p superior a 0,10, implicando a não rejeição de que as médias são iguais para esta componente.

O mesmo sucede com a variável X_2 com $F = 2,68$ e, assim, também não rejeitamos a hipótese de que estas médias são iguais.

O teste T^2 leva em conta a correlação positiva entre as duas medidas para cada grupo - informação que é simplesmente ignorada pelos testes univariados.

Estes exemplos demonstram a eficácia de um teste multivariado relativo às suas contrapartidas univariadas. No exemplo sobre o tratamento de esgotos (efluentes) também percebemos exatamente essa aparente incoerência.

Três outros testes multivariados também são comuns

1. TRAÇO DE LAWLEY-HOTELLING: $\text{tr}(\mathbf{B}\mathbf{W}^{-1})$.
2. TRAÇO DE PILLAI: $\text{tr}([\mathbf{B} + \mathbf{W}]^{-1})$.
3. MAIOR RAIZ QUADRADA DE ROY: maior autovalor de $\mathbf{W}(\mathbf{B} + \mathbf{W})^{-1}$.

Todos estes testes são aproximadamente equivalentes quando se tem grandes amostras.

Quando, e somente quando, os testes multivariados apontam diferenças, deve-se olhar os componentes separadamente.

No **R**, pacote **stats**, há a função **manova** que calcula todas estas estatísticas e também as matrizes de somas de quadrados para um ou mais fatores.

Estratégia para a comparação multivariada de tratamentos

1. Identifique primeiro a presença pontos aberrantes (*outliers*).
2. Realize um teste multivariado.
3. Se o teste multivariado apontou diferenças, calcule os intervalos simultâneos de Bonferroni que interessam.

Regressão linear

Sejam r covariáveis, $z_1, z_2, \dots, z_r$, relacionadas a uma variável resposta y .

O **modelo de regressão linear múltiplo** (univariado) é dado pela equação:

$$\underbrace{y}_{\text{resposta}} = \underbrace{\beta_0 + \beta_1 z_1 + \dots + \beta_r z_r}_{\text{média; parte estrutural}} + \underbrace{\varepsilon}_{\text{erro; parte aleatória}}$$

O modelo é dito linear, pois a parte estrutural é linear nos parâmetros β_j , $j = 1, 2, \dots, r$.

Se dispomos de n observações independentes, então

$$y_i = \beta_0 + \beta_1 z_{i1} + \dots + \beta_r z_{ir} + \varepsilon_i, \quad i = 1, 2, \dots, n.$$

Suposições:

S1: $E(\varepsilon_i) = 0, \forall i = 1, 2, \dots, n$.

S2: $\text{Var}(\varepsilon_i) = \sigma^2, \forall i = 1, 2, \dots, n$ (homocedasticidade).

S3: $\text{Cov}(\varepsilon_i, \varepsilon_k) = 0, \forall i \neq k, i, k \in \{1, 2, \dots, n\}$.

Na notação matricial, tem-se

$$\underbrace{\mathbf{y}}_{n \times 1} = \underbrace{\mathbf{Z}}_{n \times (r+1)} \underbrace{\boldsymbol{\beta}}_{(r+1) \times 1} + \underbrace{\boldsymbol{\varepsilon}}_{n \times 1},$$

com

S1: $E(\boldsymbol{\varepsilon}) = \mathbf{0}$;

S2 e S3: $\text{Var}(\boldsymbol{\varepsilon}) = \sigma^2 \mathbf{I}_n$;

$$\mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix}; \quad \mathbf{Z} = \begin{bmatrix} 1 & z_{11} & z_{12} & \dots & z_{1r} \\ 1 & z_{21} & z_{22} & \dots & z_{2r} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & z_{n1} & z_{n2} & \dots & z_{nr} \end{bmatrix};$$

$$\boldsymbol{\beta} = \begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_r \end{bmatrix}; \quad \text{e} \quad \boldsymbol{\varepsilon} = \begin{bmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \vdots \\ \varepsilon_n \end{bmatrix}$$

Observe que ainda não fizemos nenhuma suposição acerca da distribuição dos erros.

Para efeito de obter os estimadores de mínimos quadrados, de fato, não é necessária nenhuma suposição sobre a distribuição da parte aleatória. Porém, para fins de inferência, será necessário.

Estimadores de Mínimos Quadrados

Suponha que a matriz $\mathbf{Z}$ seja de posto completo tal que suas colunas formam um conjunto de vetores linearmente independentes. Neste caso, a matriz $\mathbf{Z}^\top \mathbf{Z}$ é não-singular e o estimador de mínimos quadrados do vetor β é dado por

$$\hat{\beta} = (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{y}.$$

Os valores ajustados são, então, dados por

$$\hat{\mathbf{y}} = \mathbf{Z}\hat{\beta} = \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{y} = \mathbf{H}\mathbf{y}, \quad \text{em que} \quad \mathbf{H} = \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top,$$

e os resíduos

$$\hat{\varepsilon} = \mathbf{y} - \hat{\mathbf{y}} = [\mathbf{I} - \mathbf{H}]\mathbf{y} = \mathbf{P}\mathbf{y}, \quad \text{em que} \quad \mathbf{P} = \mathbf{I} - \mathbf{H},$$

satisfazendo as seguintes relações (somente quando houver a constante β_0 no modelo)

$$\mathbf{Z}^\top \hat{\varepsilon} = \mathbf{0} \quad \text{e} \quad \hat{\mathbf{y}}^\top \hat{\varepsilon} = 0.$$

Temos que $\mathbf{H}$ e $\mathbf{P}$ são matrizes idempotentes ($\mathbf{H} = \mathbf{H}\mathbf{H}$ e $\mathbf{H} = \mathbf{H}^\top$).

A soma de quadrados de resíduos é

$$SQ_{Res} = \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \hat{\boldsymbol{\varepsilon}}^\top \hat{\boldsymbol{\varepsilon}} = \mathbf{y}^\top \mathbf{P} \mathbf{y} = \mathbf{y}^\top \mathbf{y} - \mathbf{y}^\top \mathbf{Z} \hat{\boldsymbol{\beta}}.$$

Observe que

$$\sum_{i=1}^n y_i^2 = \mathbf{y}^\top \mathbf{y} = (\mathbf{y} - \hat{\mathbf{y}} + \hat{\mathbf{y}})^\top (\mathbf{y} - \hat{\mathbf{y}} + \hat{\mathbf{y}}) = \hat{\mathbf{y}}^\top \hat{\mathbf{y}} + \hat{\boldsymbol{\varepsilon}}^\top \hat{\boldsymbol{\varepsilon}}.$$

O coeficiente de determinação R^2

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

O R^2 fornece a proporção da variação total dos y_i 's que é “explicada” pelas covariáveis.

Por um lado, o R^2 seria igual 1 se a equação do modelo se ajustasse perfeitamente aos dados.

Por outro lado, o R^2 seria zero se $\hat{\beta}_0 = \bar{y}$ e os demais seriam todos nulos. Neste caso, as covariáveis não exerceriam nenhuma influência sobre a resposta.

O R^2 deve ser olhado com cuidado na verificação do modelo, pois valores altos de R^2 não necessariamente implicam que o modelo ajustado é bom. Além dessa medida é fundamental realizar uma análise dos resíduos.

Ademais, um R^2 não tão elevado, para um modelo ajustado cuja análise de resíduos foi boa, pode ser considerado.

Considere o modelo

$$\mathbf{y} = \mathbf{Z}\boldsymbol{\beta} + \boldsymbol{\varepsilon}.$$

$E(\mathbf{y}) = \mathbf{Z}\boldsymbol{\beta}$ é uma combinação linear das linhas da matriz $\mathbf{Z}$ com coeficientes $\beta_0, \beta_1, \dots, \beta_r$.

A medida que $\boldsymbol{\beta}$ varia, $\mathbf{Z}\boldsymbol{\beta}$ gera o “plano modelo” de todas as combinações lineares das colunas de $\mathbf{Z}$.

Geralmente, o vetor observado $\mathbf{y}$ não pertencerá ao plano modelo devido ao erro aleatório. Isto é, $\mathbf{y}$ não é uma combinação linear das colunas de $\mathbf{Z}$.

Uma vez que as observações se tornam disponíveis, a solução de mínimos quadrados é obtida a partir do vetor desvio dado por

$$\underbrace{\mathbf{y}}_{\text{vetor de observação}} - \underbrace{\mathbf{Z}\boldsymbol{\beta}}_{\text{vetor no plano modelo}}$$

O quadrado do módulo deste vetor é

$$S(\boldsymbol{\beta}) = (\mathbf{y} - \mathbf{Z}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{Z}\boldsymbol{\beta}).$$

$S(\hat{\beta})$ é tão pequeno quanto possível, quando $\hat{\beta}$ é selecionado tal que $\mathbf{Z}\hat{\beta}$ é o ponto no plano modelo mais próximo do vetor de observação.

De fato, $\mathbf{Z}\hat{\beta}$ é a projeção ortogonal do vetor de observação $\mathbf{y}$ sobre o plano modelo.

O vetor $\hat{\epsilon}$ é ortogonal ao plano modelo.

Como $\hat{\mathbf{y}} = \mathbf{Z}\hat{\beta} = \mathbf{Z}(\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \mathbf{y} = \mathbf{H}\mathbf{y}$, dizemos que $\mathbf{H}$ é a matriz de projeção ortogonal do vetor de observações sobre o plano modelo $\mathbf{Z}\beta$.

Similarmente, $\mathbf{P}$ é a matriz de projeção de $\mathbf{y}$ sobre o plano perpendicular ao modelo.

Inferências sobre os parâmetros

Sob o modelo $\mathbf{y} = \mathbf{Z}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$ com $E(\boldsymbol{\varepsilon}) = \mathbf{0}$ e $\text{Var}(\boldsymbol{\varepsilon}) = \sigma^2 \mathbf{I}$, o estimador de mínimos quadrados de $\boldsymbol{\beta}$ é $\hat{\boldsymbol{\beta}} = (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{y}$. Então,

$$\begin{aligned} E(\hat{\boldsymbol{\beta}}) &= \boldsymbol{\beta}, \\ \text{Var}(\hat{\boldsymbol{\beta}}) &= \sigma^2 (\mathbf{Z}^\top \mathbf{Z})^{-1}, \\ E(\hat{\boldsymbol{\varepsilon}}) &= \mathbf{0}, \text{ e} \\ \text{Var}(\hat{\boldsymbol{\varepsilon}}) &= \sigma^2 \mathbf{P}. \end{aligned}$$

Observe que apesar dos erros serem supostos independentes, os resíduos não o serão necessariamente, pois a matriz $\mathbf{P} = \mathbf{I} - \mathbf{H}$ pode não ser diagonal.

Temos que

$$E(\hat{\boldsymbol{\varepsilon}}^\top \hat{\boldsymbol{\varepsilon}}) = (n - r - 1)\sigma^2,$$

tal que $s^2 = \frac{\hat{\boldsymbol{\varepsilon}}^\top \hat{\boldsymbol{\varepsilon}}}{n - r - 1} = \frac{\mathbf{y}^\top \mathbf{P} \mathbf{y}}{n - r - 1}$ é um estimador não tendencioso de σ^2 .

Além disso, $\hat{\boldsymbol{\beta}}$ e $\hat{\boldsymbol{\varepsilon}}$ são não correlacionados.

Para ir adiante nas inferências se impõe a necessidade de se fazer suposições sobre a distribuição dos erros.

A suposição adicional aqui é $\varepsilon \sim \mathcal{N}_n(\mathbf{0}, \sigma^2 \mathbf{I})$.

Com esta suposição adicional, é possível provar que $\hat{\beta}$ é o estimador de máxima verossimilhança de β . Além disso,

$$\hat{\beta} \sim \mathcal{N}_{r+1}(\beta, \sigma^2(\mathbf{Z}^\top \mathbf{Z})^{-1})$$

e

$$\frac{(n - r - 1)s^2}{\sigma^2} = \frac{\hat{\varepsilon}^\top \hat{\varepsilon}}{\sigma^2} \sim \chi_{n-r-1}^2.$$

Podemos então construir uma região de confiança conjunta para β usando

$$\frac{(\beta - \hat{\beta})^\top \mathbf{Z}^\top \mathbf{Z}(\beta - \hat{\beta})}{s^2} \sim (r + 1)F_{r+1, n-r-1}.$$

Intervalos simultâneos de $100(1 - \alpha)\%$ para os β_i 's

$$\text{IC}_{100(1-\alpha)\%}(\beta_i) = \hat{\beta}_i \pm \sqrt{\widehat{\text{Var}}(\hat{\beta}_i)} \sqrt{(r+1)F_{r+1, n-r-1, 1-\alpha}}, \quad i = 0, 1, \dots, r,$$

em que $\widehat{\text{Var}}(\hat{\beta}_i)$ é o i -ésimo elemento da diagonal de $s^2(\mathbf{Z}^\top \mathbf{Z})^{-1}$.

é comum ignorar a simultaneidade dos intervalos se trabalhando com os intervalos separados que são mais estreitos.

$$\text{IC}_{100(1-\alpha)\%}(\beta_i) = \hat{\beta}_i \pm \sqrt{\widehat{\text{Var}}(\hat{\beta}_i)} t_{n-r-1, 1-\alpha/2}.$$

Neste caso, deve-se atentar para o fato de que aqui a confiança não é simultânea para os $r + 1$ intervalos.

Uma abordagem que é interessante de ser usada aqui é a abordagem de Bonferroni.

Inferências a partir da função de regressão estimada

Estimação da média de uma observação nova $E(y_0|\mathbf{z}_0)$, com $\mathbf{z}_0^\top = (1, z_{01}, \dots, z_{0r})$.

A média é dada por $E(y_0|\mathbf{z}_0) = \beta_0 + \beta_1 z_{01} + \dots + \beta_r z_{0r} = \mathbf{z}_0^\top \boldsymbol{\beta}$.

Logo, $E(\widehat{y_0}|\mathbf{z}_0) = \mathbf{z}_0^\top \widehat{\boldsymbol{\beta}}$.

Temos que

$$\mathbf{z}_0^\top \widehat{\boldsymbol{\beta}} \sim \mathcal{N}(\mathbf{z}_0^\top \boldsymbol{\beta}, \sigma^2 \mathbf{z}_0^\top (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{z}_0).$$

O intervalo de $100(1-\alpha)\%$ de confiança para $\mathbf{z}_0^\top \boldsymbol{\beta}$ é dado por

$$IC_{100(1-\alpha)\%}(\mathbf{z}_0^\top \boldsymbol{\beta}) = \mathbf{z}_0^\top \widehat{\boldsymbol{\beta}} \pm t_{n-r-1, 1-\alpha/2} \sqrt{s^2 \mathbf{z}_0^\top (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{z}_0}.$$

Previsão de uma nova observação, quando $\mathbf{z} = \mathbf{z}_0$

O intervalo de previsão 100(1- α)% de confiança para y_0 é dado por

$$IC_{100(1-\alpha)\%}(y_0) = \mathbf{z}_0^\top \hat{\boldsymbol{\beta}} \pm t_{n-r-1, 1-\alpha/2} \sqrt{s^2(1 + \mathbf{z}_0^\top (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{z}_0)}.$$

Alguns outros tópicos que devem ser revisados pelo aluno:

- Verificação do ajuste e outros aspectos da regressão.
- Análise de resíduos.
- Pontos de alavanca e de influência.
- Seleção de covariáveis.
- Comparação de modelos.

Na notação matricial, tem-se

$$\underbrace{\mathbf{Y}}_{n \times p} = \underbrace{\mathbf{Z}}_{n \times (r+1)} \underbrace{\boldsymbol{\beta}}_{(r+1) \times p} + \underbrace{\boldsymbol{\epsilon}}_{n \times p}$$

com

$$\mathbf{Y} = \begin{bmatrix} y_{11} & y_{12} & \cdots & y_{1p} \\ y_{21} & y_{22} & \cdots & y_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ y_{n1} & y_{n2} & \cdots & y_{np} \end{bmatrix} = [\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_p],$$

$$\mathbf{Z} = \begin{bmatrix} 1 & z_{11} & z_{12} & \cdots & z_{1r} \\ 1 & z_{21} & z_{22} & \cdots & z_{2r} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & z_{n1} & z_{n2} & \cdots & z_{nr} \end{bmatrix} = [\mathbf{1}, \mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_r],$$

$$\boldsymbol{\beta} = \begin{bmatrix} \beta_{01} & \beta_{02} & \cdots & \beta_{0p} \\ \beta_{11} & \beta_{12} & \cdots & \beta_{1p} \\ \vdots & \vdots & \ddots & \vdots \\ \beta_{r1} & \beta_{r2} & \cdots & \beta_{rp} \end{bmatrix} = [\boldsymbol{\beta}_1, \boldsymbol{\beta}_2, \dots, \boldsymbol{\beta}_p], \quad \mathbf{e}$$

$$\boldsymbol{\epsilon} = \begin{bmatrix} \epsilon_{11} & \epsilon_{12} & \cdots & \epsilon_{1p} \\ \epsilon_{21} & \epsilon_{22} & \cdots & \epsilon_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ \epsilon_{n1} & \epsilon_{n2} & \cdots & \epsilon_{np} \end{bmatrix} = [\boldsymbol{\epsilon}_1, \boldsymbol{\epsilon}_2, \dots, \boldsymbol{\epsilon}_p].$$

Suposições do modelo $\underbrace{\mathbf{Y}}_{n \times p} = \underbrace{\mathbf{Z}}_{n \times (r+1)} \underbrace{\boldsymbol{\beta}}_{(r+1) \times p} + \underbrace{\boldsymbol{\varepsilon}}_{n \times p}$:

- $E(\boldsymbol{\varepsilon}_j) = \mathbf{0}$, e
- $\text{Cov}(\boldsymbol{\varepsilon}_j, \boldsymbol{\varepsilon}_k) = \sigma_{jk} \mathbf{I}_n$, para $j, k = 1, 2, \dots, p$.

As p medidas sobre a i -ésima observação têm matriz de covariância dada por $\boldsymbol{\Sigma}$, mas medidas provenientes observações diferentes são não correlacionadas.

Temos que $\boldsymbol{\beta}$ e $\boldsymbol{\Sigma}$ são desconhecidos.

Observe que a j -ésima coluna da matriz resposta segue o modelo linear univariado múltiplo dado por

$$\mathbf{y}_j = \mathbf{Z}\boldsymbol{\beta}_j + \boldsymbol{\varepsilon}_j, \quad \text{para } j = 1, 2, \dots, p.$$

com $\text{Var}(\boldsymbol{\varepsilon}_j) = \sigma_{jj} \mathbf{I}_n$.

Estimação por Mínimos Quadrados:

$$\mathbf{B} = (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{Y}.$$

Matriz de somas de quadrados e produtos cruzados:

$$(\mathbf{Y} - \mathbf{ZB})^\top (\mathbf{Y} - \mathbf{ZB}).$$

Valores ajustados:

$$\hat{\mathbf{Y}} = \mathbf{ZB} = \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{Y}$$

Resíduos:

$$\hat{\boldsymbol{\varepsilon}} = \mathbf{Y} - \hat{\mathbf{Y}} = \mathbf{Y} - \mathbf{ZB} = [\mathbf{I} - \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top] \mathbf{Y}.$$

Propriedades:

$$\mathbf{Z}^\top \hat{\boldsymbol{\varepsilon}} = \underbrace{\mathbf{0}}_{(r+1) \times p} \quad \text{e} \quad \hat{\mathbf{Y}}^\top \hat{\boldsymbol{\varepsilon}} = \underbrace{\mathbf{0}}_{p \times p};$$

$$\mathbf{Y}^\top \mathbf{Y} = \hat{\mathbf{Y}}^\top \hat{\mathbf{Y}} + \hat{\boldsymbol{\varepsilon}}^\top \hat{\boldsymbol{\varepsilon}}; \quad \text{e}$$

$$\hat{\boldsymbol{\varepsilon}}^\top \hat{\boldsymbol{\varepsilon}} = \mathbf{Y}^\top \mathbf{Y} - \mathbf{B}^\top (\mathbf{Z}^\top \mathbf{Z}) \mathbf{B}.$$

Para o estimador de mínimos quadrados $\mathbf{B}$ e com a matriz $\mathbf{Z}$ de posto completo, tem-se

$$\begin{aligned} E(\mathbf{B}) &= \beta, \\ \text{Cov}(\hat{\beta}_j, \hat{\beta}_k) &= \sigma_{jk}(\mathbf{Z}^\top \mathbf{Z})^{-1}, \\ E(\hat{\varepsilon}) &= \underbrace{\mathbf{0}}_{\text{matriz nula}}, \text{ e} \\ E\left(\frac{1}{n-r-1} \hat{\varepsilon}^\top \hat{\varepsilon}\right) &= \Sigma. \end{aligned}$$

$\mathbf{B}$ e $\hat{\varepsilon}$ são não correlacionados.

Estimador não tendencioso de Σ :

$$\mathbf{S} = \frac{\hat{\varepsilon}^\top \hat{\varepsilon}}{n-r-1}.$$

Inferências a partir da função de regressão estimada

Suponha que o modelo $\mathbf{Y} = \mathbf{Z}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$, com erros normais tenha sido ajustado. Se o ajuste for considerado bom, ele poderá ser usado para fins de previsão.

Um problema é prever a média correspondente a um vetor de covariáveis $\mathbf{z}_0$.

Inferências sobre a média podem ser feitas se usando os mesmos resultados estudados no caso univariado.

$$\mathbf{B}^\top \mathbf{z}_0 \sim \mathcal{N}_p(\boldsymbol{\beta}^\top \mathbf{z}_0, \mathbf{z}_0^\top (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{z}_0 \boldsymbol{\Sigma})$$

e

$$n\hat{\boldsymbol{\Sigma}} \sim \mathcal{W}_p(n - r - 1, \boldsymbol{\Sigma}), \quad \text{independentemente.}$$

Intervalos- T^2 simultâneos de $100(1-\alpha)\%$ de confiança para $E(y_k|\mathbf{z}_0) = \mathbf{z}_0^\top \boldsymbol{\beta}_k$:

$$\mathbf{z}_0^\top \hat{\boldsymbol{\beta}}_k \pm \sqrt{\frac{p(n-r-1)}{n-r-p} F_{p, n-r-p, 1-\alpha}} \sqrt{\mathbf{z}_0^\top (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{z}_0 \left(\frac{n}{n-r-1} \hat{\sigma}_{kk} \right)},$$

para $k = 1, 2, \dots, p$.

O outro problema está voltado para a previsão de uma nova resposta $\mathbf{y}_0$ dado o vetor de covariáveis $\mathbf{z}_0$.

Agora,

$$\mathbf{y}_0 \sim \mathcal{N}_p(\boldsymbol{\beta}^\top \mathbf{z}_0, (1 + \mathbf{z}_0^\top (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{z}_0) \boldsymbol{\Sigma}).$$

E, assim, intervalos simultâneos de previsão de $100(1-\alpha)\%$ para a k -ésima componente de $\mathbf{y}_0$, $y_{0,k}$, $k = 1, 2, \dots, p$, são dados por

$$\mathbf{z}_0^\top \hat{\boldsymbol{\beta}}_k \pm \sqrt{\frac{p(n-r-1)}{n-r-p} F_{p, n-r-p, 1-\alpha}} \sqrt{(1 + \mathbf{z}_0^\top (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{z}_0) \left(\frac{n}{n-r-1} \hat{\sigma}_{kk} \right)},$$

O conceito de regressão linear

O modelo de regressão linear clássico diz respeito à associação entre uma variável dependente (ou resposta) y e uma coleção de covariáveis

$z_1, z_2, \dots, z_r$.

O modelo que consideramos trata y como uma variável aleatória cuja média depende dos valores fixados das covariáveis.

A média é uma função linear dos coeficientes da regressão $\beta_0, \beta_1, \dots, \beta_r$.

O modelo de regressão linear também surge em uma configuração diferente.

Suponha que todas as variáveis $y, z_1, z_2, \dots, z_r$ são variáveis aleatórias e têm uma distribuição conjunta com média μ de dimensão $(r + 1) \times 1$ e variância Σ de dimensão $(r + 1) \times (r + 1)$.

Particionando μ e Σ de forma apropriada temos

$$\mu = \begin{bmatrix} \mu_y \\ \mu_z \end{bmatrix} \quad \text{e} \quad \Sigma = \begin{bmatrix} \sigma_{yy} & \Sigma_{yz} \\ \Sigma_{zy} & \Sigma_{zz} \end{bmatrix},$$

com

$$\Sigma_{yz} = [\sigma_{yz_1}, \sigma_{yz_2}, \dots, \sigma_{yz_r}].$$

Σ_{zz} é suposta ser não singular. Se isso não acontecer, significa que pelo menos uma das covariáveis é combinação linear das demais, de modo que esta covariável é redundante e pode ser eliminada do problema.

Preditor Linear:

$$\hat{\beta}_0 + \hat{\beta}_1 z_1 + \dots + \hat{\beta}_r z_r = \hat{\beta}_0 + \hat{\beta}^\top \mathbf{z}, \quad \text{com} \quad \hat{\beta}^\top = (\hat{\beta}_1, \hat{\beta}_2, \dots, \hat{\beta}_r).$$

Para um preditor dessa forma, o erro de previsão é dado por:

$$y - \hat{\beta}_0 - \hat{\beta}^\top \mathbf{z}.$$

Como este erro é aleatório, costuma-se escolher $\hat{\beta}_0$ e $\hat{\beta}$ de modo a minimizar o erro quadrático médio

$$\text{EQM} = E \left([y - \hat{\beta}_0 - \hat{\beta}^\top \mathbf{z}]^2 \right).$$

O EQM depende da distribuição conjunta de y e $\mathbf{z}$ somente através de μ e Σ .
é possível expressar o preditor linear ótimo em função dessas quantidades.

Resultado

O preditor linear $\beta_0 + \beta^\top \mathbf{z}$ com coeficientes $\beta = \Sigma_{zz}^{-1} \Sigma_{zy}$, $\beta_0 = \mu_y - \beta^\top \mu_z$ tem EQM mínimo entre todos os preditores lineares de y . Além disso,

$$E \left([y - \hat{\beta}_0 - \hat{\beta}^\top \mathbf{z}]^2 \right) = \sigma_{yy} - \Sigma_{yz} \Sigma_{zz}^{-1} \Sigma_{zy}.$$

Também $\beta_0 + \beta^\top \mathbf{z}$ é o preditor linear de maior correlação com y dada por

$$\sqrt{\frac{\Sigma_{yz} \Sigma_{zz}^{-1} \Sigma_{zy}}{\sigma_{yy}}}.$$

Esse coeficiente de correlação é chamado coeficiente de correlação múltiplo da população e será denotado por $\rho_{y(z)}$.

$$\rho_{y(z)} = \sqrt{\frac{\mathbf{\Sigma}_{yz} \mathbf{\Sigma}_{zz}^{-1} \mathbf{\Sigma}_{zy}}{\sigma_{yy}}}.$$

Observe que este coeficiente de correlação múltiplo, diferente dos outros varia entre 0 e 1.

O quadrado do coeficiente de correlação múltiplo populacional é chamado **coeficiente de determinação** populacional.

O coeficiente de determinação populacional tem uma interpretação importante. O EQM em usar $\beta_0 + \beta^\top \mathbf{z}$ para prever y é

$$\sigma_{yy} - \mathbf{\Sigma}_{yz} \mathbf{\Sigma}_{zz}^{-1} \mathbf{\Sigma}_{zy} = \sigma_{yy} - \sigma_{yy} \rho_{y(z)}^2 = \sigma_{yy} (1 - \rho_{y(z)}^2).$$

Se $\rho_{y(z)} = 0$, $\mathbf{z}$ não acrescenta nenhum poder de previsão. No outro extremo, y pode ser perfeitamente explicada por $\mathbf{z}$.

Previsão de várias variáveis

A extensão destes resultados para a previsão de várias respostas $y_1, y_2, \dots, y_p$ é quase imediata.

Suponha

$$\begin{pmatrix} \mathbf{y}_{p \times 1} \\ \mathbf{z}_{r \times 1} \end{pmatrix} \sim \mathcal{N}_{p+r}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$$

com

$$\boldsymbol{\mu} = \begin{pmatrix} \boldsymbol{\mu}_y \\ \boldsymbol{\mu}_z \end{pmatrix} \quad \text{e} \quad \boldsymbol{\Sigma} = \begin{bmatrix} \boldsymbol{\Sigma}_{yy} & \boldsymbol{\Sigma}_{yz} \\ \boldsymbol{\Sigma}_{zy} & \boldsymbol{\Sigma}_{zz} \end{bmatrix}.$$

Vimos que

$$E(\mathbf{y} | z_1, z_2, \dots, z_r) = \boldsymbol{\mu}_y + \boldsymbol{\Sigma}_{yz} \boldsymbol{\Sigma}_{zz}^{-1} (\mathbf{z} - \boldsymbol{\mu}_z).$$

Este valor esperado condicional, considerado como uma função de $z_1, z_2, \dots, z_r$, é chamado vetor de regressão multivariada de $\mathbf{y}$ em $\mathbf{z}$.

Ele se compõe de p regressões múltiplas univariadas.

Por exemplo, o primeiro componente do vetor de média condicional é:

$$\mu_{y_1} + \Sigma_{y_1 z} \Sigma_{zz}^{-1} (\mathbf{z} - \mu_z) = E(y_1 | z_1, z_2, \dots, z_r)$$

que minimiza o erro de previsão nesse componente.

A matriz $\beta_{p \times r} = \Sigma_{yz} \Sigma_{zz}^{-1}$ é chamada matriz de coeficientes de regressão.

O vetor de erro de previsão é:

$$\mathbf{v} = \mathbf{y} - \mu_y - \Sigma_{yz} \Sigma_{zz}^{-1} (\mathbf{z} - \mu_z) \quad \text{e} \quad E(\mathbf{v} \mathbf{v}^T) = \Sigma_{yy.z} = \Sigma_{yy} - \Sigma_{yz} \Sigma_{zz}^{-1} \Sigma_{zy}.$$

Como μ e Σ costumam ser desconhecidos, eles devem ser estimados.

Coeficiente de correlação parcial

Considere $\begin{pmatrix} \mathbf{y}_{p \times 1} \\ \mathbf{z}_{r \times 1} \end{pmatrix} \sim \mathcal{N}_{p+r}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$.

O coeficiente de correlação parcial entre y_j e y_k , eliminando-se $\mathbf{z}$, é definido por

$$\rho_{y_j y_k \cdot \mathbf{z}} = \frac{\sigma_{y_j y_k \cdot \mathbf{z}}}{\sqrt{\sigma_{y_j y_j \cdot \mathbf{z}} \sigma_{y_k y_k \cdot \mathbf{z}}}}$$

com $\sigma_{y_j y_k \cdot \mathbf{z}}$ a entrada (j, k) da matriz $\boldsymbol{\Sigma}_{yy \cdot \mathbf{z}} = \boldsymbol{\Sigma}_{yy} - \boldsymbol{\Sigma}_{yz} \boldsymbol{\Sigma}_{zz}^{-1} \boldsymbol{\Sigma}_{zy}$.

O coeficiente de correlação amostral correspondente é dado por

$$r_{y_j y_k \cdot} = \frac{s_{y_j y_k \cdot \mathbf{z}}}{\sqrt{s_{y_j y_j \cdot \mathbf{z}} s_{y_k y_k \cdot \mathbf{z}}}}$$

com $s_{y_j y_k \cdot \mathbf{z}}$ a entrada (j, k) da matriz $\mathbf{S}_{yy \cdot \mathbf{z}} = \mathbf{S}_{yy} - \mathbf{S}_{yz} \mathbf{S}_{zz}^{-1} \mathbf{S}_{zy}$.

Análise em componentes principais

O propósito da **análise de componentes principais** (ACP) é substituir as variáveis originais por um número menor de variáveis que são funções 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.

Em geral, espera-se que os primeiras componentes chamados **componentes principais** (em número menor do que o de variáveis originais) compreendam a maior parte da variação total no conjunto de dados original tal que a dimensionalidade efetiva dos dados pode ser reduzida.

A **análise fatorial** (AF) tem um propósito similar, mas é baseada em um modelo estatístico próprio que especifica um dado número de variáveis subjacentes chamadas fatores. A AF é uma técnica voltada para a “explicação” da estrutura de covariâncias das variáveis, em vez de olhar apenas as variâncias.

Devido às similaridades das duas técnicas os dois métodos são algumas vezes confundidos. Porém, existem diferenças fundamentais que serão apresentadas ao longo deste curso.

Em muitas aplicações, os componentes principais ou fatores aparecem como o objetivo final da análise e, os pesquisadores tentam, então, interpretá-los de forma significativa. Porém, o principal benefício da ACP está na possibilidade de reduzir a dimensão do problema de forma a simplificar as análises subsequentes.

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 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) 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 .

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.

Por exemplo, construir o gráfico de dispersão dos escores dos dois primeiros componentes para cada observação é uma forma de tentar identificar conglomerados (agrupamentos).

A ACP transforma um conjunto de variáveis correlacionadas em um conjunto de variáveis não correlacionadas. Se as variáveis originais são aproximadamente não correlacionadas entre si, então não existe razão para realizarmos uma ACP.

Componentes principais populacionais

Algebricamente, CP 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 ao rotacionar o sistema original. Os novos eixos representam as direções com variabilidade máxima e fornecem descrição mais simples e parcimoniosa da estrutura de covariância.

Os CP dependem somente da estrutura de covariância Σ (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.

Seja $\mathbf{X} = (X_1, X_2, \dots, X_p)^\top$ um vetor aleatório p -dimensional, com média $\boldsymbol{\mu}$ e matriz de covariâncias $\boldsymbol{\Sigma}$. Desejamos encontrar um novo conjunto de variáveis, $\mathbf{Y} = (Y_1, Y_2, \dots, Y_p)^\top$, 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 tomado 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 $\mathbf{a}_j = (a_{j1}, a_{j2}, \dots, a_{jp})^\top$ 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 ao escolher $\mathbf{a}_1$ 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 ao escolher $\mathbf{a}_2$ 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 correlacionadas com os anteriores e tenham variâncias decrescentes.

Definição

Se $\mathbf{A}$ é uma matriz $p \times p$ e $\mathbf{v}$ é um vetor não nulo de tamanho p tal que

$$\mathbf{A}\mathbf{v} = \lambda\mathbf{v}, \quad \text{com } \lambda \geq 0,$$

então dizemos que λ é um autovalor correspondente a $\mathbf{v}$. Da teoria de decomposição espectral de $\mathbf{A}$, temos que $\mathbf{A} = \mathbf{P}\mathbf{\Lambda}\mathbf{P}^\top$ com $\mathbf{P} = (\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_p)$ e $\mathbf{\Lambda} = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_p)$.

Proposição 1

Seja Σ 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 Σ 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:

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

e

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

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

Prova

Vimos que

$$\max_{\mathbf{a} \neq \mathbf{0}} \frac{\mathbf{a}^\top \boldsymbol{\Sigma} \mathbf{a}}{\mathbf{a}^\top \mathbf{a}} = \lambda_1.$$

O máximo é atingido quando $\mathbf{a} = \mathbf{v}_1$.

Mas $\mathbf{v}_1^\top \mathbf{v}_1 = 1$ e, assim,

$$\max_{\mathbf{a} \neq \mathbf{0}} \frac{\mathbf{a}^\top \boldsymbol{\Sigma} \mathbf{a}}{\mathbf{a}^\top \mathbf{a}} = \lambda_1 = \frac{\mathbf{v}_1^\top \boldsymbol{\Sigma} \mathbf{v}_1}{\mathbf{v}_1^\top \mathbf{v}_1} = \mathbf{v}_1^\top \boldsymbol{\Sigma} \mathbf{v}_1 = \text{Var}(Y_1).$$

Similarmente, vimos que

$$\max_{\mathbf{a} \perp \mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k} \frac{\mathbf{a}^\top \boldsymbol{\Sigma} \mathbf{a}}{\mathbf{a}^\top \mathbf{a}} = \lambda_{k+1} \quad \text{para } k = 1, 2, \dots, p-1.$$

Para a escolha $\mathbf{a} = \mathbf{v}_{k+1}$ com $\mathbf{v}_{k+1}^\top \mathbf{v}_j = 0$, para $j = 1, 2, \dots, k$ e $k = 1, 2, \dots, p-1$,

$$\text{Var}(Y_{k+1}) = \lambda_{k+1} \quad \text{e} \quad \text{Cov}(Y_{k+1}, Y_j) = 0, \quad j = 1, 2, \dots, k.$$

Proposição 2

Seja $\mathbf{X}$ um 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,

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

Prova

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

Logo,

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

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 $\mathbf{v}_j = (v_{j1}, v_{j2}, \dots, v_{jp})^\top$ que definem os componentes 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 Σ do vetor aleatório $\mathbf{X}$, então

$$\rho(Y_j, X_k) = \frac{v_{jk} \sqrt{\lambda_j}}{\sqrt{\sigma_{kk}}}, \quad \text{para } j, k = 1, 2, \dots, p,$$

são os coeficientes de correlação entre Y_j e X_k .

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

Prova

Faça $\mathbf{e}_k = (0, \dots, 0, \underbrace{1}_{k\text{-ésima entrada}}, 0, \dots, 0)^\top$ tal que $X_k = \mathbf{e}_k^\top \mathbf{X}$ e

$$\begin{aligned} \text{Cov}(X_k, Y_j) &= \text{Cov}(\mathbf{e}_k^\top \mathbf{X}, \mathbf{v}_j^\top \mathbf{X}) = \mathbf{e}_k^\top \Sigma \mathbf{v}_j = \mathbf{e}_k^\top \mathbf{v}_j \lambda_j = v_{jk} \lambda_j \Rightarrow \\ \rho(Y_j, X_k) &= \frac{\text{Cov}(Y_j, X_k)}{\sqrt{\text{Var}(Y_j)} \sqrt{\text{Var}(X_k)}} = \frac{v_{jk} \lambda_j}{\sqrt{\lambda_j} \sqrt{\sigma_{kk}}} = \frac{v_{jk} \sqrt{\lambda_j}}{\sqrt{\sigma_{kk}}}. \end{aligned}$$

Exemplo

Suponha um vetor aleatório de dimensão 3 cuja matriz de covariâncias é dada por

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

Os pares de autovalores e autovetores associados a Σ são:

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

Portanto, os componentes principais são:

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

Ver arquivo exemplo_13.r

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 5: 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

ACP considerando uma normal multivariada

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

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

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 $(1/\lambda_1, \mathbf{v}_1), (1/\lambda_2, \mathbf{v}_2), \dots, (1/\lambda_p, \mathbf{v}_p)$. Note que $\boldsymbol{\Sigma}^{-1} = \mathbf{P}\boldsymbol{\Lambda}^{-1}\mathbf{P}$.

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 λ_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 $\mathbf{v}_j$.

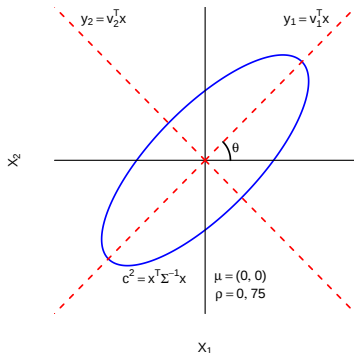


Figura 7: Elipse de densidade constante e os componentes principais para um vetor aleatório normal bivariado com $\mu = \mathbf{0}$ e $\rho = 0,75$. (Ver arquivo exemplo_14.r)

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

Componentes principais obtidos das variáveis padronizadas

Os componentes também podem ser obtidos via variáveis padronizadas. Na notação matricial

$$\mathbf{Z} = \left[\boldsymbol{\Psi}^{1/2} \right]^{-1} (\mathbf{X} - \boldsymbol{\mu}) \quad \text{com} \quad \boldsymbol{\Psi}^{1/2} = \text{diag}(\sqrt{\sigma_{11}}, \sqrt{\sigma_{22}}, \dots, \sqrt{\sigma_{pp}}).$$

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 δ_j/p , para $j = 1, 2, \dots, p$, em que δ_j é 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

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

Os pares de autovalores e autovetores de Σ são dados por

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

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

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

Ver arquivo exemplo_15.r

Os respectivos componentes principais são, então,

- relativos a Σ :

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

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

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 Σ .

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 Σ dados por 0,040 e 0,999, respectivamente.

Este exemplo demonstra que componentes principais obtidas de Σ 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.

Discussão

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.

ACP para matrizes de covariância com estruturas especiais

Existem certos padrões das matrizes de covariâncias e correlações cujos componentes principais podem ser expressos de forma simples. Por exemplo, suponha que

$$\Sigma = \text{diag}(\sigma_{11}, \sigma_{22}, \dots, \sigma_{pp}).$$

Fazendo $\mathbf{e}_j = (0, \dots, 0, \underbrace{1}_{j\text{-ésima posição}}, 0, \dots, 0)^\top$ é fácil ver que $\Sigma \mathbf{e}_j = \sigma_{jj} \mathbf{e}_j$

e concluímos que os pares $(\sigma_{jj}, \mathbf{e}_j)$, $j = 1, 2, \dots, p$ são pares de autovalores e autovetores de Σ . Como a combinação linear $\mathbf{e}_j^\top \mathbf{X} = X_j$, o conjunto de componentes principais é exatamente o conjunto de variáveis originais não correlacionadas, exceto que agora apresentadas em ordem decrescente de variância.

Para uma matriz de covariâncias diagonal, nada se ganha em usar ACP. De outro ponto de vista, se $\mathbf{X}$ tem distribuição $\mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$, os contornos de densidade constante são elipsóides cujos eixos já estão na direção da variação máxima. Consequentemente, não há necessidade de rotacionar o sistema de coordenadas originais.

Padronização não altera substancialmente a situação para $\boldsymbol{\Sigma}$ diagonal. Neste caso $\mathbf{R} = \mathbf{I}_p$. Claramente, $\mathbf{R}\mathbf{e}_j = 1 \times \mathbf{e}_j$, tal que o autovalor 1 é de multiplicidade p e

$$\mathbf{e}_j = (0, \dots, 0, \underbrace{1}_{j\text{-ésima posição}}, 0, \dots, 0)^\top, \quad j = 1, 2, \dots, p,$$

é uma escolha conveniente para os autovetores. Consequentemente, os componentes principais obtidos por $\mathbf{R}$ são também as variáveis originais, agora padronizadas, $Z_1, Z_2, \dots, Z_p$. No caso de autovalores iguais e sob normalidade, os contornos de densidade constante serão esferóides.

Outra matriz de covariância padrão, que em geral descreve correspondência entre certas variáveis biológicas tem a seguinte forma geral

$$\Sigma = \begin{bmatrix} \sigma^2 & \rho\sigma^2 & \cdots & \rho\sigma^2 \\ \rho\sigma^2 & \sigma^2 & \cdots & \rho\sigma^2 \\ \vdots & \vdots & \ddots & \vdots \\ \rho\sigma^2 & \rho\sigma^2 & \cdots & \sigma^2 \end{bmatrix}.$$

A matriz de correlações correspondente é

$$\mathcal{R} = \begin{bmatrix} 1 & \rho & \cdots & \rho \\ \rho & 1 & \cdots & \rho \\ \vdots & \vdots & \ddots & \vdots \\ \rho & \rho & \cdots & 1 \end{bmatrix}$$

Observe que com este padrão as variáveis $X_1, X_2, \dots, X_p$ são igualmente correlacionadas.

Não é difícil mostrar que os p autovalores correspondentes à matriz de correlações podem ser divididos em dois grupos.

Quando $\rho > 0$, o maior autovalor é $\lambda_1 = 1 + (p - 1)\rho$ com autovetor associado

$$\mathbf{v}_1 = \left(\frac{1}{\sqrt{p}}, \frac{1}{\sqrt{p}}, \dots, \frac{1}{\sqrt{p}} \right)^\top.$$

Os $p - 1$ autovalores restantes são $\lambda_2 = \lambda_3 = \dots = \lambda_p = 1 - \rho$ e uma escolha apropriada para os autovetores é

$$\mathbf{v}_2 = \left(\frac{1}{\sqrt{1 \times 2}}, \frac{-1}{\sqrt{1 \times 2}}, 0, \dots, 0 \right)^\top$$

$$\mathbf{v}_3 = \left(\frac{1}{\sqrt{2 \times 3}}, \frac{1}{\sqrt{2 \times 3}}, \frac{-2}{\sqrt{2 \times 3}}, 0, \dots, 0 \right)^\top$$

$$\vdots \quad \vdots \quad \quad \quad \vdots$$

$$\mathbf{v}_j = \left(\frac{1}{\sqrt{(j-1) \times j}}, \dots, \frac{1}{\sqrt{(j-1) \times j}}, \frac{-(j-1)}{\sqrt{(j-1) \times j}}, 0, \dots, 0 \right)^\top$$

$$\vdots \quad \vdots \quad \quad \quad \vdots$$

$$\mathbf{v}_p = \left(\frac{1}{\sqrt{(p-1) \times p}}, \dots, \frac{1}{\sqrt{(p-1) \times p}}, \frac{-(p-1)}{\sqrt{(p-1) \times p}} \right)^\top.$$

O primeiro componente principal é

$$Y_1 = \mathbf{v}_1^\top \mathbf{Z} = \frac{1}{\sqrt{p}} \sum_{i=1}^p Z_i$$

que é proporcional à soma das p variáveis padronizadas. Ele pode ser olhado como um “índice” com pesos iguais para todas as variáveis. A proporção da variação total explicada por este componente é dada por

$$\frac{\lambda_1}{p} = \frac{1 + (p-1)\rho}{p} = \rho + \frac{1-\rho}{p}.$$

Se ρ estiver próximo de 1 ou p for grande temos que $\frac{\lambda_1}{p} \simeq \rho$. Por exemplo, se $\rho = 0,80$ e $p = 5$, este componente explicará cerca de 84% da variação total do dados.

Quando ρ está perto de 1, os $p-1$ autovalores restantes contribuem relativamente muito pouco e podem ser desprezados.

Se o vetor aleatório $(Z_1, Z_2, \dots, Z_p)$ tem distribuição normal multivariada com correlações iguais duas a duas, então o elipsóide de densidade constante terá uma forma de charuto.

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

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.

O valor do j -ésimo CP para esta observação é $\hat{y}_j = \hat{\mathbf{v}}_j^\top \mathbf{x}$.

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

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

Denotaremos

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

Exemplo

Um censo forneceu informações sobre 5 variáveis socioeconômicas em 61 setores censitários na área de Madison, Wisconsin.

As variáveis observadas foram as seguintes:

- 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?

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

Tabela 6: Médias amostrais das variáveis.

Ver exemplo_16.r

	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

Tabela 7: Matriz de covariâncias 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

Tabela 8: Matriz de correlações amostrais das variáveis.

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

Os **desvios padrões** são dados por:

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

	$\hat{v}_1$	$\hat{v}_2$	$\hat{v}_3$	$\hat{v}_4$	$\hat{v}_5$
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

Tabela 9: Rotação das variáveis.

	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

Tabela 10: Importância das componentes.

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.

	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

Tabela 11: Correlações entre os componentes e as variáveis

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 (continuação)

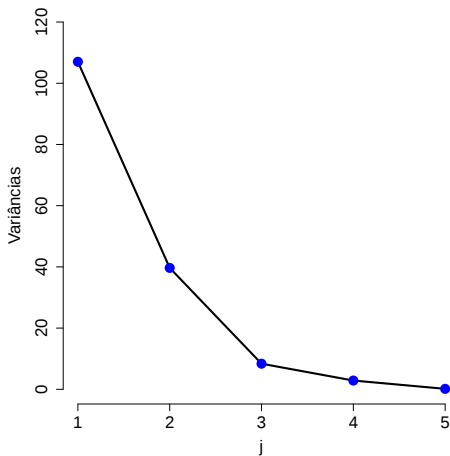


Figura 8: Autovalores ordenados de forma não crescente.

Componentes principais usando as variáveis padronizadas

Os componentes principais podem ser obtidos via matriz de correlações amostrais $\mathbf{R}$ 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

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 $(\text{preço de fechamento da semana corrente} - \text{preço de fechamento da semana anterior}) / (\text{preço de fechamento da semana anterior})$, ajustada para *stock splits* e dividendos.

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 se mover junto em resposta à conjuntura econômica.

Variável	JP	Citi	Fargo	Royal	Exxon
$\bar{x}$	0,00106	0,00066	0,00163	0,00405	0,00404

Tabela 12: Médias amostrais das variáveis.

Ver exemplo_17.r

	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

Tabela 13: Matriz de correlações amostrais das variáveis.

	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

Tabela 14: Importância das componentes.

	$\hat{v}_1$	$\hat{v}_2$	$\hat{v}_3$	$\hat{v}_4$	$\hat{v}_5$
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

Tabela 15: Rotação das variáveis.

- 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, constatamos 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 e as variáveis é dada a seguir.

	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

Tabela 16: Correlações entre os componentes e as variáveis.

A seguir mostramos o *scree plot*.

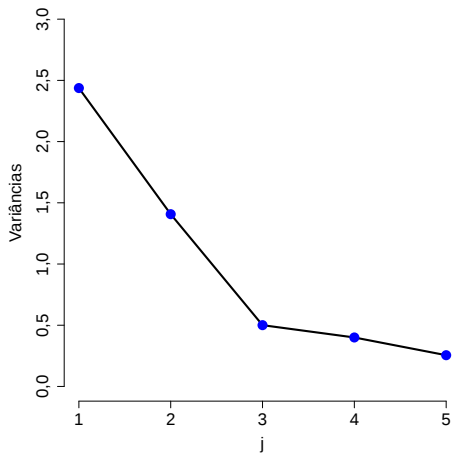


Figura 9: Autovalores ordenados de forma não crescente.

Gráficos dos componentes principais

Gráficos dos CP podem revelar observações suspeitas bem como desvios da suposição de normalidade. Como os CP são combinações lineares das variáveis originais, é razoável esperar que eles se comportem 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 são aproximadamente normais, quando eles são usados como dados de entrada em análises posteriores.

Os últimos CP 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.

Assim, as magnitudes dos últimos CP 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 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. Eles ajudam a identificar observações suspeitas.

Exemplo (continuação)

Usando os dados sobre ações, podemos construir os gráficos de dispersão dos dois primeiros CP e respectivos *qq-plots* para normalidade, e para os dois últimos CP para identificar a presença de observações suspeitas.

Ver exemplo_18.r

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 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 $\lambda_1 > \lambda_2 > \dots > \lambda_p > 0$ 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 λ_j , $j = 1, 2, \dots, p$, os autovalores de $\boldsymbol{\Sigma}$. Faça $\boldsymbol{\lambda} = (\lambda_1, \lambda_2, \dots, \lambda_p)^\top$. 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

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

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

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

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

$$\widehat{\lambda}_j \stackrel{ind,a}{\sim} \mathcal{N}\left(\lambda_j, \frac{2\lambda_j^2}{n}\right).$$

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

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

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 λ_j 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 $\frac{1}{n} \hat{\mathbf{V}}_j$ em que os elementos de $\hat{\mathbf{V}}_j$ são obtidos a partir de $\mathbf{V}_j$ substituindo suas entradas por suas estimativas $\hat{\lambda}_j$ e $\hat{\mathbf{v}}_j$.

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

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

Assim, os IC simultâneos são dados por

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

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

Teste para estrutura de correlações iguais

$$H_0 : \mathbf{R} = \mathbf{R}_0 = \begin{bmatrix} 1 & \rho & \cdots & \rho \\ \rho & 1 & \cdots & \rho \\ \vdots & \vdots & \ddots & \vdots \\ \rho & \rho & \cdots & 1 \end{bmatrix} \quad \text{versus} \quad H_1 : \mathbf{R} \neq \mathbf{R}_0.$$

É possível construir um teste destas hipóteses usando a estatística da razão de verossimilhança. Porém, apresentaremos um teste equivalente, devido à Lawley, que é construído a partir dos elementos fora da diagonal de $\mathbf{R}$.

Faça

$$\bar{r}_k = \frac{1}{p-1} \sum_{i=1, i \neq k}^p r_{ik}, \quad k = 1, 2, \dots, p,$$

$$\bar{r} = \frac{2}{p(p-1)} \sum \sum_{i < k} r_{ik}, \quad \text{e}$$

$$\hat{\gamma} = \frac{(p-1)^2 [1 - (1 - \bar{r})^2]}{p - (p-2)(1 - \bar{r})^2}.$$

É evidente que $\bar{r}_k$ é a média dos elementos fora da diagonal na k -ésima coluna de $\mathbf{R}$ e $\bar{r}$ é a média global dos elementos fora da diagonal de $\mathbf{R}$.

Testes aproximados para grandes amostras rejeitam H_0 ao nível de significância α se

$$T > \chi^2_{(p+1)(p-2)/2, 1-\alpha} \quad \text{com}$$

$$T = \frac{n-1}{(1-\bar{r})^2} \left[\sum \sum_{i < k} (r_{ik} - \bar{r})^2 - \hat{\gamma} \sum_{k=1}^p (\bar{r}_k - \bar{r})^2 \right].$$

Exemplo

Desejamos fazer um teste para a estrutura de equicorrelação. Para isto, temos

$$\mathbf{R} = \begin{bmatrix} 1 & 0,7501 & 0,6329 & 0,6363 \\ 0,7501 & 1 & 0,6925 & 0,7386 \\ 0,6329 & 0,6925 & 1 & 0,6625 \\ 0,6363 & 0,7386 & 0,6625 & 1 \end{bmatrix},$$

$\bar{r}_1 = 0,6731$, $\bar{r}_2 = 0,7271$, $\bar{r}_3 = 0,66626$, $\bar{r}_4 = 0,6791$, $\bar{r} = 0,6855$ e $\hat{\gamma} = 2,1329$.

Fazendo os cálculos, obtemos $T \simeq 11,4$.

Temos também $\chi^2_{5; 0,95} \simeq 11,07$.

Logo, rejeitamos H_0 ao nível de significância de 5%.

Supondo população normal multivariada, existe um teste assintótico para verificar se todas as variáveis são independentes ($\mathcal{R} = \mathbf{I}_p$).

Análise fatorial

O propósito fundamental da **análise fatorial (AF)** é descrever, se possível, as relações de covariâncias entre muitas variáveis em função de umas poucas, porém não observáveis, quantidades aleatórias subjacentes que são chamadas **fatores**.

O modelo de AF é motivado pelo seguinte argumento.

Suponha que as variáveis podem ser agrupadas por suas correlações. Isto é, suponha que todas as variáveis dentro de um grupo particular sejam altamente correlacionadas entre si, mas apresentam correlações relativamente pequenas com variáveis de outros grupos.

Então, é razoável propor um único constructo (fator) subjacente para cada grupo, que resume as correlações observadas no grupo.

AF pode ser considerada uma extensão da ACP. Ambas podem ser vistas como voltadas para aproximar a matriz de covariâncias Σ . Porém, a aproximação baseada no modelo AF é mais elaborada.

A questão primária na AF é se os dados são consistentes com uma estrutura prescrita.

Modelo fatorial ortogonal

Seja $\mathbf{X}$ vetor aleatório observável $p \times 1$, com $\boldsymbol{\mu} = E(\mathbf{X})$ e $\boldsymbol{\Sigma} = \text{Var}(\mathbf{X})$.

O modelo fatorial (MF) postula que $\mathbf{X}$ é linearmente dependente de $m < p$ variáveis aleatórias não observáveis $F_1, F_2, \dots, F_m$ chamadas fatores comuns, e p fontes de variação adicionais $\varepsilon_1, \varepsilon_2, \dots, \varepsilon_p$ chamadas erros ou, algumas vezes, fatores específicos.

$$X_j - \mu_j = \ell_{j1}F_1 + \ell_{j2}F_2 + \dots + \ell_{jm}F_m + \varepsilon_j, \quad \text{para } j = 1, 2, \dots, p,$$

ou, em notação matricial,

$$\underbrace{\mathbf{X} - \boldsymbol{\mu}}_{p \times 1} = \underbrace{\mathbf{L}}_{p \times m} \underbrace{\mathbf{F}}_{m \times 1} + \underbrace{\boldsymbol{\varepsilon}}_{p \times 1},$$

em que

$$\boldsymbol{\varepsilon} = \begin{bmatrix} \varepsilon_1 & \varepsilon_2 & \dots & \varepsilon_p \end{bmatrix}^T \quad (\text{vetor aleatório}),$$

$$\mathbf{F} = \begin{bmatrix} F_1 & F_2 & \dots & F_m \end{bmatrix}^T \quad (\text{vetor aleatório}),$$

$$\mathbf{L} = \begin{bmatrix} \ell_1 & \ell_2 & \dots & \ell_m \end{bmatrix} \quad \text{e}$$

$$\ell_k = \begin{bmatrix} \ell_{1k} & \ell_{2k} & \dots & \ell_{pk} \end{bmatrix}^T, \quad k = 1, 2, \dots, m.$$

Os coeficientes ℓ_{jk} são chamados cargas da j -ésima variável sobre o k -ésimo fator ($j = 1, 2, \dots, p$ e $k = 1, 2, \dots, m$).

O fator específico ε_j está associado à j -ésima variável (X_j).

Os p desvios $X_j - \mu_j$ são expressos em função de $p + m$ quantidades aleatórias: $F_1, F_2, \dots, F_m$ e $\varepsilon_1, \varepsilon_2, \dots, \varepsilon_p$, não observáveis.

Isto distingue o modelo AF do modelo de regressão linear múltipla multivariado no qual as variáveis explicativas são observáveis.

Com tantas quantidades não observáveis, uma verificação do modelo fatorial a partir das observações sobre $X_1, X_2, \dots, X_p$ é desanimadora.

Porém, com algumas suposições adicionais sobre os vetores aleatórios $\mathbf{F}$ e $\boldsymbol{\varepsilon}$, o modelo AF implicará em certas relações de covariâncias que podem ser verificadas.

Suposições adicionais:

1. $E(\mathbf{F}) = \mathbf{0}$ e $\text{Var}(\mathbf{F}) = E(\mathbf{F}\mathbf{F}^\top) = \mathbf{I}_m$;
2. $E(\varepsilon) = \mathbf{0}$ e $\text{Var}(\varepsilon) = E(\varepsilon\varepsilon^\top) = \Psi = \text{diag}\{\psi_1, \psi_2, \dots, \psi_p\}$; e
3. $\mathbf{F}$ e ε são independentes tal que $\text{Cov}(\mathbf{F}, \varepsilon) = \mathbf{0}_{m \times p}$.

Com estas suposições o modelo apresentado é chamado **modelo fatorial ortogonal (MFO)**.

O MFO implica a seguinte estrutura de covariâncias para $\mathbf{X}$:

$$\begin{aligned}(\mathbf{X} - \mu)(\mathbf{X} - \mu)^\top &= (\mathbf{L}\mathbf{F} + \varepsilon)(\mathbf{L}\mathbf{F} + \varepsilon)^\top \\ &= \mathbf{L}\mathbf{F}\mathbf{F}^\top\mathbf{L}^\top + \mathbf{L}\mathbf{F}\varepsilon^\top + \varepsilon\mathbf{F}^\top\mathbf{L}^\top + \varepsilon\varepsilon^\top.\end{aligned}$$

tal que

$$\begin{aligned}\Sigma &= \text{Var}(\mathbf{X}) = E\left((\mathbf{X} - \mu)(\mathbf{X} - \mu)^\top\right) \\ &= \mathbf{L}E(\mathbf{F}\mathbf{F}^\top)\mathbf{L}^\top + \mathbf{L}E(\mathbf{F}\varepsilon^\top) + E(\varepsilon\mathbf{F}^\top)\mathbf{L}^\top + E(\varepsilon\varepsilon^\top) \\ &= \mathbf{L}\mathbf{L}^\top + \Psi.\end{aligned}$$

Também, é fácil verificar que,

$$\text{Cov}(\mathbf{X}, \mathbf{F}) = \text{E} \left((\mathbf{X} - \boldsymbol{\mu}) \mathbf{F}^\top \right) = \mathbf{L}.$$

Logo, o MFO implica em $\boldsymbol{\Sigma} = \mathbf{L}\mathbf{L}^\top + \boldsymbol{\Psi}$ e $\text{Cov}(\mathbf{X}, \mathbf{F}) = \mathbf{L}$ tal que

$$\text{Cov}(X_j, f_k) = \ell_{jk} \text{ e } \text{Cor}(X_j, f_k) = \frac{\ell_{jk}}{\sqrt{\sigma_{jj}}}, j = 1, 2, \dots, p \text{ e } k = 1, 2, \dots, m.$$

O modelo $\mathbf{X} - \boldsymbol{\mu} = \mathbf{L}\mathbf{F} + \boldsymbol{\varepsilon}$ é linear nos fatores comuns. Se as p respostas são, de fato, relacionadas aos fatores subjacentes, mas a relação não é linear, tal como

$$X_1 - \mu_1 = \ell_{11}F_1F_3 + \varepsilon_1, \quad , X_2 - \mu_2 = \ell_{22}F_2F_3 + \varepsilon_3, \quad \text{etc.},$$

então a estrutura de covariâncias dada por $\mathbf{L}\mathbf{L}^\top + \boldsymbol{\Psi}$ pode não ser adequada.

A suposição de linearidade é fundamental na formulação do modelo fatorial usual.

- A porção da variância da j -ésima variável contribuída pelos m fatores é chamada j -ésima **comunalidade**.
- A porção da variância devida ao fator específico ε_j é chamada **variância específica**.

Denotaremos a j -ésima comunalidade por h_j^2 e a variância específica por ψ_j . Assim,

$$\underbrace{\sigma_{jj}}_{\text{Var}(X_j)} = \underbrace{\ell_{j1}^2 + \ell_{j2}^2 + \cdots + \ell_{jm}^2}_{\text{comunalidade}} + \underbrace{\psi_j}_{\text{var. específica}},$$

$$h_j^2 = \sum_{k=1}^m \ell_{jk}^2 \quad \text{e} \quad \sigma_{jj} = h_j^2 + \psi_j.$$

A j -ésima comunalidade é a soma dos quadrados das cargas da j -ésima variável sobre os m fatores, $j = 1, 2, \dots, p$.

Exemplo

Aqui, faremos uma verificação da relação $\Sigma = \mathbf{LL}^\top + \Psi$ para $m = 2$.

$$\Sigma = \begin{bmatrix} 19 & 30 & 2 & 12 \\ 30 & 57 & 5 & 23 \\ 2 & 5 & 38 & 47 \\ 12 & 23 & 47 & 68 \end{bmatrix} = \begin{bmatrix} 4 & 1 \\ 7 & 2 \\ -1 & 6 \\ 1 & 8 \end{bmatrix} \begin{bmatrix} 4 & 7 & -1 & 1 \\ 1 & 2 & 6 & 8 \end{bmatrix} + \begin{bmatrix} 2 & 0 & 0 & 0 \\ 0 & 4 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 3 \end{bmatrix},$$

ou, equivalentemente, $\Sigma = \mathbf{LL}^\top + \Psi$ com

$$\mathbf{L}^\top = \begin{bmatrix} 4 & 7 & -1 & 1 \\ 1 & 2 & 6 & 8 \end{bmatrix} \quad \text{e} \quad \Psi = \text{diag}(2, 4, 1, 3).$$

Variável	Comunalidade	Var. específica	Variância
X_1	16+1=17	2	19
X_2	49+4=53	4	57
X_3	1+36=37	1	38
X_4	1+64=65	3	68

Ver arquivo exemplo_19.r

O MFO assume que os $p(p+1)/2$ elementos distintos da matriz de covariâncias de $\mathbf{X}$, $\mathbf{\Sigma}$, podem ser representados pelas pm cargas das p variáveis sobre os m fatores e as p variâncias específicas.

Por exemplo, se $p = 10$ e $m = 2$ o MFO simplifica a estrutura de covariâncias com 55 parâmetros desconhecidos para a estimação de 30.

Se $m = p$, qualquer matriz de covariâncias $\mathbf{\Sigma}$ pode ser decomposta em $\mathbf{LL}^T$, assim $\mathbf{\Psi}$ pode ser a matriz nula.

Porém, é quando m é pequeno relativamente a p que a AF terá utilidade.

Neste caso, o MFO fornece uma explicação simples das covariâncias de $\mathbf{X}$ com menos parâmetros dos que os $p(p+1)/2$ em $\mathbf{\Sigma}$.

Mas, muitas das matrizes Σ não podem ser decompostas na forma $LL^T + \Psi$ quando m é bem menor do que p .

Exemplo

Suponha $p = 3$, $m = 1$ e $\Sigma = \begin{bmatrix} 1 & 0,9 & 0,7 \\ 0,9 & 1 & 0,4 \\ 0,7 & 0,4 & 1 \end{bmatrix}$.

Usando o MFO a estrutura imposta para Σ é $\underbrace{L}_{3 \times 1} L^T + \Psi$ ou

$$\begin{array}{ll} \ell_{11}^2 + \psi_1 = 1 & \ell_{11}\ell_{21} = 0,9 \\ \ell_{21}^2 + \psi_2 = 1 & \ell_{11}\ell_{31} = 0,7 \\ \ell_{31}^2 + \psi_3 = 1 & \ell_{21}\ell_{31} = 0,4 \end{array}$$

Resolvendo o sistema encontramos uma solução única, porém inconsistente com as interpretações estatísticas dadas aos parâmetros. Por exemplo, conclui-se que $\text{Cov}(X_1, f_1) = \text{Cor}(X_1, f_1) = \ell_{11} = \pm 1,255!!!$

Observe que $\text{Var}(X_1) = 1$ e $\text{Var}(F_1) = 1$.

Ver arquivo exemplo_20.r

Além disso, também obtem-se $\psi_1 = -0,575!$

Assim, para este exemplo, com $m = 1$ é possível obter uma solução numérica única para a equação $\Sigma = \mathbf{L}\mathbf{L}^\top + \Psi$. No entanto, a solução não é apropriada.

Quando $m > 1$, há sempre uma ambiguidade inerente associada ao modelo fatorial. Para ver isto, seja $\mathbf{Q}$ uma matriz ortogonal $m \times m$ tal que $\mathbf{Q}^\top \mathbf{Q} = \mathbf{Q}\mathbf{Q}^\top = \mathbf{I}_m$. Então,

$$\mathbf{X} - \mu = \mathbf{L}\mathbf{F} + \varepsilon = \mathbf{L}\mathbf{Q}\mathbf{Q}^\top \mathbf{F} + \varepsilon = \mathbf{L}^* \mathbf{F}^* + \varepsilon.$$

com $\mathbf{L}^* = \mathbf{L}\mathbf{Q}$ e $\mathbf{F}^* = \mathbf{Q}^\top \mathbf{F}$.

Como $E(\mathbf{F}^*) = \mathbf{Q}^\top E(\mathbf{F}) = \mathbf{0}$ e $\text{Var}(\mathbf{F}^*) = \mathbf{Q}^\top \text{Var}(\mathbf{F}) \mathbf{Q} = \mathbf{Q}^\top \mathbf{I}_m \mathbf{Q} = \mathbf{I}_m$ é impossível, com base nas observações sobre $\mathbf{X}$, distinguir as cargas em $\mathbf{L}$ das cargas em $\mathbf{L}^*$. Isto é, os fatores $\mathbf{F}$ e $\mathbf{F}^*$, têm as mesmas propriedades estatísticas e, apesar das cargas em $\mathbf{L}^*$ serem, em geral, diferentes das cargas em $\mathbf{L}$, ambas geram a mesma matriz de covariâncias Σ .

$$\Sigma = \mathbf{L}\mathbf{L}^\top + \Psi = (\mathbf{L}\mathbf{Q})(\mathbf{L}\mathbf{Q})^\top + \Psi = [\mathbf{L}^*][\mathbf{L}^*]^\top + \Psi.$$

Esta ambiguidade fornece a justificativa para o uso da “rotação de fatores”, pois as transformações lineares correspondentes às matrizes ortogonais correspondem às rotações (e reflexões) no sistema de coordenadas $\mathbf{X}$.

A análise do MFO prossegue se impondo condições que permitem estimar univocamente $\mathbf{L}$ e Ψ .

A matriz de cargas $\mathbf{L}$ é, então, rotacionada (pré-multiplicada por uma matriz ortogonal). A rotação é determinada por algum critério de facilitação da interpretação dos fatores.

Obtidas as cargas e as variâncias específicas, os fatores são identificados, e valores estimados para os fatores em si, chamados **escores dos fatores**, são calculados.

Métodos de estimação

Dadas as observações $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$, p -variadas, a AF busca responder a seguinte questão.

O MFO, com um número pequeno de fatores, representa adequadamente os dados?

Em essência, lidamos com esse problema tentando verificar a estrutura de covariâncias dos dados.

$\mathbf{S}$ é um estimador de $\mathbf{\Sigma}$. Se os elementos fora da diagonal de $\mathbf{S}$ são pequenos (fora da diagonal de $\mathbf{R}$ estão próximos de 0), as variáveis são não correlacionadas e uma AF não será útil.

Nestas circunstâncias, os fatores específicos representarão o papel principal, enquanto que o objetivo central da AF é determinar uns poucos fatores comuns importantes.

Se $\mathbf{\Sigma}$ parece se desviar de uma matriz diagonal, então um MFO pode ser investigado, e o problema inicial será estimar as cargas ℓ_{jk} e as variâncias específicas ψ_j , $j = 1, 2, \dots, p$ e $k = 1, 2, \dots, m$.

Consideraremos dois métodos de estimação de parâmetros: o método dos **componentes principais** e o método da **máxima verossimilhança**.

A solução de cada método pode ser rotacionada de forma a facilitar a interpretação dos fatores comuns.

é sempre prudente tentar mais de um método de solução, se o MFO é adequado para o problema em mãos, as soluções deveriam ser consistentes umas com as outras.

Os métodos de estimação e rotação requerem cálculos iterativos. Vários programas estatísticos dispõem de rotinas com esse fim. No **R** está disponível a função **factanal**.

Método das componentes principais

Considere a decomposição espectral de Σ com $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p \geq 0$ tal que $\Sigma = \sum_{j=1}^p \lambda_j \mathbf{v}_j \mathbf{v}_j^\top$ e observe que podemos escrever $\Sigma = \mathbf{B}\mathbf{B}^\top$ com

$$\mathbf{B} = [\sqrt{\lambda_1} \mathbf{v}_1 \quad \sqrt{\lambda_2} \mathbf{v}_2 \quad \dots \quad \sqrt{\lambda_p} \mathbf{v}_p].$$

Isto ajusta a estrutura de covariâncias prescrita para o MFO quando $m = p$ e $\psi_j = 0 \ \forall j$. As cargas da j -ésima coluna são dadas por $\sqrt{\lambda_j} \mathbf{v}_j$.

$$\Sigma = \mathbf{L}\mathbf{L}^\top + \underbrace{\Psi}_{=0} = \mathbf{B}\mathbf{B}^\top.$$

Exceto pelo fator de escala $\sqrt{\lambda_j}$, as cargas do j -ésimo fator para as p variáveis são os coeficientes do j -ésimo componente principal da população.

Apesar da representação aqui considerada ser exata, ela não é particularmente útil, pois usa tantos fatores quantas são as variáveis e não permite qualquer variação nos fatores específicos dados por ε .

É desejável, no entanto, modelos que expliquem a estrutura de covariâncias em função de apenas uns poucos fatores comuns.

Uma abordagem possível é, quando os últimos $p - m$ autovalores são bem pequenos, desprezar a contribuição de $\lambda_j \mathbf{v}_j \mathbf{v}_j^\top$, $j = m + 1, m + 2, \dots, p$ para Σ .

Desprezando esta contribuição, obtemos

$$\Sigma \simeq \sum_{j=1}^m \lambda_j \mathbf{v}_j \mathbf{v}_j^\top = \underbrace{\mathbf{L}}_{p \times m} \mathbf{L}^\top \quad \text{com}$$

$$\mathbf{L} = [\ell_1 \quad \ell_2 \quad \cdots \quad \ell_m] = [\sqrt{\lambda_1} \mathbf{v}_1 \quad \sqrt{\lambda_2} \mathbf{v}_2 \quad \cdots \quad \sqrt{\lambda_m} \mathbf{v}_m].$$

Esta aproximação assume que os fatores específicos ε são de menor importância e também podem ser ignorados na decomposição da matriz Σ .

Se os fatores específicos estão incluídos no modelo, suas variâncias podem ser tomadas como os elementos da diagonal da matriz

$$\Sigma - \mathbf{L} \mathbf{L}^\top.$$

Incluindo os fatores específicos, temos que

$$\begin{aligned}\Sigma &\simeq \mathbf{L}\mathbf{L}^\top + \Psi \\ &= \begin{bmatrix} \sqrt{\lambda_1}\mathbf{v}_1 & \sqrt{\lambda_2}\mathbf{v}_2 & \cdots & \sqrt{\lambda_m}\mathbf{v}_m \end{bmatrix} \begin{bmatrix} \sqrt{\lambda_1}\mathbf{v}_1^\top \\ \sqrt{\lambda_2}\mathbf{v}_2^\top \\ \vdots \\ \sqrt{\lambda_m}\mathbf{v}_m^\top \end{bmatrix} + \text{diag}(\psi_1, \psi_2, \dots, \psi_p)\end{aligned}$$

em que $\psi_j = \sigma_{jj} - \sum_{k=1}^m \ell_{jk}^2$, para $j = 1, 2, \dots, p$.

Na aplicação desta abordagem aos dados $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ observados é usual primeiro centrar as observações trabalhando com $\mathbf{x}_i - \bar{\mathbf{x}}$, $i = 1, 2, \dots, n$. Observe que a matriz de covariâncias amostral $\mathbf{S}$ não se altera com esta transformação de localização.

Nos casos em que as unidades das variáveis são não comensuráveis é desejável trabalhar com as variáveis padronizadas decompondo, então, a matriz de correlações.

A padronização evita o problema de se ter uma variável com variância grande influenciando indevidamente a determinação das cargas sobre os fatores.

O uso da aproximação sugerida por $\mathbf{\Sigma} \simeq \mathbf{LL}^\top + \mathbf{\Psi}$ acima é conhecido como solução via componentes principais (CP).

A AF via CP de $\mathbf{S}$ é obtida em função da decomposição espectral da matriz de covariâncias amostral $\mathbf{S}$. Sejam $(\hat{\lambda}_j, \hat{\mathbf{v}}_j)$ os pares de autovalores e autovetores de $\mathbf{S}$ com $\hat{\lambda}_1 \geq \hat{\lambda}_2 \geq \dots \geq \hat{\lambda}_p \geq 0$.

Seja $m < p$ o número de fatores comuns. Então,

$$\begin{aligned}\hat{\mathbf{L}} &= \begin{bmatrix} \sqrt{\hat{\lambda}_1} \hat{\mathbf{v}}_1 & \sqrt{\hat{\lambda}_2} \hat{\mathbf{v}}_2 & \dots & \sqrt{\hat{\lambda}_m} \hat{\mathbf{v}}_m \end{bmatrix} \quad \text{e} \\ \hat{\mathbf{\Psi}} &= \text{diag}(\hat{\psi}_1, \hat{\psi}_2, \dots, \hat{\psi}_p),\end{aligned}$$

em que $\hat{\psi}_j = s_{jj} - \sum_{k=1}^m \hat{\ell}_{jk}^2$, isto é, $\hat{\mathbf{\Psi}}$ matriz diagonal cujos elementos diagonais são os elementos diagonais da matriz $\mathbf{S} - \hat{\mathbf{L}}\hat{\mathbf{L}}^\top$.

Comunalidades: $\hat{h}_j^2 = \sum_{k=1}^m \hat{\ell}_{jk}^2$.

Observações:

1. A solução da AF via CP com base na matriz de correlações amostral $\mathbf{R}$ é obtida da mesma forma. Neste caso, denotaremos os autovalores por $\hat{\delta}_j$ e os autovetores por $\hat{\mathbf{w}}_j$, para $j = 1, 2, \dots, p$.
2. Na AF via CP as cargas dos fatores não se alteram se aumentarmos o número de fatores.

Pela definição dos $\hat{\psi}_j$, os elementos da diagonal de $\mathbf{S}$ são iguais aos elementos da diagonal de $\hat{\mathbf{L}}\hat{\mathbf{L}}^\top + \hat{\Psi}$. Porém, os elementos fora da diagonal de $\mathbf{S}$ não serão, em geral, reproduzidos por $\hat{\mathbf{L}}\hat{\mathbf{L}}^\top + \hat{\Psi}$.

Como, então, selecionar m ?

Se m não é determinado a priori por considerações tais como teóricas, sua escolha pode ser baseada sobre os autovalores estimados, da mesma forma que em ACP.

Considere a matriz residual dada por $\mathbf{S} - \widehat{\mathbf{L}}\widehat{\mathbf{L}}^\top - \widehat{\boldsymbol{\Psi}}$, cujos elementos da diagonal são nulos. Se os outros elementos também forem próximos de zero, podemos, subjetivamente, considerar o MFO a m fatores apropriado.

Analiticamente temos:

Soma dos quadrados das entradas de

$$\mathbf{S} - \widehat{\mathbf{L}}\widehat{\mathbf{L}}^\top - \widehat{\boldsymbol{\Psi}} \leq \widehat{\lambda}_{m+1}^2 + \widehat{\lambda}_{m+2}^2 + \cdots + \widehat{\lambda}_p^2.$$

Consequentemente, um valor pequeno da soma dos quadrados dos autovalores estimados desprezados implica em um valor pequeno da soma de quadrados dos erros de aproximação.

Idealmente, as contribuições dos primeiros poucos fatores às variâncias amostrais das variáveis deve ser grande.

A contribuição para s_{jj} do primeiro fator é $\hat{\ell}_{j1}^2$. A contribuição total para a variância amostral total $\text{tr}(\mathbf{S}) = \sum_{j=1}^p s_{jj}$ do primeiro fator comum é

$$\sum_{j=1}^p \hat{\ell}_{j1}^2 = \left(\sqrt{\hat{\lambda}_1} \hat{\mathbf{v}}_1 \right)^\top \left(\sqrt{\hat{\lambda}_1} \hat{\mathbf{v}}_1 \right) = \hat{\lambda}_1, \text{ pois } \hat{\mathbf{v}}_1 \text{ tem norma } 1.$$

Resumindo, a proporção da variação amostral total devida ao k -ésimo fator ($k = 1, 2, \dots, m$) é dada por

$$\frac{\hat{\lambda}_k}{\text{tr}(\mathbf{S})} \quad \text{se AF via } \mathbf{S}$$

$$\frac{\hat{\delta}_k}{p} \quad \text{se AF via } \mathbf{R}.$$

Este critério costuma ser usado como um esquema heurístico para determinar o número de fatores comuns apropriado.

Note que $\hat{\lambda}_k = \sqrt{\hat{\lambda}_k} \mathbf{v}_k^\top \mathbf{v}_k \sqrt{\hat{\lambda}_k} = \boldsymbol{\ell}_k^\top \boldsymbol{\ell}_k = \sum_{j=1}^p \hat{\ell}_{jk}^2$ (soma de quadrados da k -ésima coluna de $\mathbf{L}$ - correspondente ao k -ésimo fator f_k).

O número de fatores comuns retido no modelo é crescente até que uma proporção adequada da variação amostral total tenha sido explicada.

Outra convenção muito usada nos pacotes estatísticos é fazer m igual ao número de autovalores maiores do que 1, se estamos com os dados padronizados ou igual ao número de autovalores **positivos** (“bem maiores que zero”) de **S**, se estamos trabalhando com os dados originais.

Essas “regras do polegar” não devem ser aplicadas indiscriminadamente. Por exemplo, $m = p$, se a regra para **S** for obedecida, já que esperamos que todos os autovalores de **S** sejam positivos para tamanhos amostrais grandes.

A melhor abordagem é reter poucos em vez de muitos fatores, assumindo que eles fornecem uma interpretação satisfatória dos dados e produzem um ajuste adequado para **S** ou **R**.

Exemplo

Em um estudo sobre a preferências de consumo, uma amostra de consumidores foi solicitada a classificar vários atributos de um novo produto. As respostas, em uma escala diferencial semântica de 1 a 7 pontos, foram tabuladas e a matriz de correlações dos atributos foi calculada. Os atributos classificados foram sabor (X_1), preço (X_2), aroma (X_3), “bom para lanche” (X_4) e “fornece muita energia” (X_5).

$$R = \begin{bmatrix} 1 & 0,02 & 0,96 & 0,42 & 0,01 \\ 0,02 & 1 & 0,13 & 0,71 & 0,85 \\ 0,96 & 0,13 & 1 & 0,50 & 0,11 \\ 0,42 & 0,71 & 0,50 & 1 & 0,79 \\ 0,01 & 0,85 & 0,11 & 0,79 & 1 \end{bmatrix}$$

A partir das entradas de R , verificamos que as variáveis 1 e 3 (gosto e aroma) e as variáveis 2 e 5 (preço e energia) formam grupos. A variável 4 (bom para lanche) está mais próxima do grupo (2,5) do que grupo (1,3).

Ver arquivo exemplo_21.r

Dados estes resultados e o pequeno número de variáveis, podemos esperar que as relações lineares aparentes entre as variáveis podem ser explicadas em função de, no máximo, 2 ou 3 fatores comuns.

Calculando a decomposição espectral de $\mathbf{R}$ obtemos $\hat{\delta}_1 \simeq 2,853$ e $\hat{\delta}_2 \simeq 1,806$ tal que ambos explicam cerca 93,2% da variação amostral total. Logo, a escolha $m = 2$ parece apropriada.

Na tabela a seguir são apresentadas as estimativas das cargas dos fatores sobre as cinco variáveis, das communalidades e das variâncias.

Variável	$\hat{\ell}_1$	$\hat{\ell}_2$	$\hat{h}_j^2$	$\hat{\psi}_j$
Z_1	-0,560	0,816	0,979	0,021
Z_2	-0,777	-0,524	0,879	0,121
Z_3	-0,645	0,748	0,976	0,024
Z_4	-0,939	-0,105	0,893	0,107
Z_5	-0,798	-0,543	0,932	0,068

Lembre-se que $\hat{\ell}_k = \sqrt{\hat{\delta}_k} \mathbf{w}_k \Rightarrow \hat{\delta}_k = \hat{\ell}_k^\top \hat{\ell}_k, k = 1, 2.$

$$\widehat{LL}^T + \widehat{\Psi} = \begin{bmatrix} 1,000 & 0,007 & 0,972 & 0,440 & 0,004 \\ 0,007 & 1,000 & 0,110 & 0,785 & 0,905 \\ 0,972 & 0,110 & 1,000 & 0,528 & 0,109 \\ 0,440 & 0,785 & 0,528 & 1,000 & 0,807 \\ 0,004 & 0,905 & 0,109 & 0,807 & 1,000 \end{bmatrix}$$

reproduz **R** adequadamente. A matriz residual correspondente a solução $m = 2$ é:

$$\mathbf{R} - \widehat{LL}^T - \widehat{\Psi} = \begin{bmatrix} 0,000 & 0,013 & -0,012 & -0,020 & 0,006 \\ 0,013 & 0,000 & 0,020 & -0,075 & -0,055 \\ -0,012 & 0,020 & 0,000 & -0,028 & 0,001 \\ -0,020 & -0,075 & -0,028 & 0,000 & -0,017 \\ 0,006 & -0,055 & 0,001 & -0,017 & 0,000 \end{bmatrix}$$

As communalidades indicam que os dois fatores levam em conta a maior parte da variação de cada variável.

Interpretação: Uma rotação dos fatores frequentemente revela uma estrutura simples e auxilia na interpretação dos fatores. Voltaremos a este exemplo depois de tratarmos este tema.

Exemplo

Considere novamente os dados sobre preços de $p = 5$ ações durante 103 semanas consecutivas. Realizamos uma ACP sobre estes dados padronizados e verificamos que os dois primeiros componentes correspondiam a cerca de 77% da variação total. Fazendo $m = 1$ e $m = 2$, podemos facilmente obter a solução do MFO via CP.

Especificamente, as cargas dos fatores estimadas são os coeficientes dos componentes principais (autovetores de $\mathbf{R}$), corrigidas pela raiz quadrada dos correspondentes autovalores.

A tabela a seguir apresenta as estimativas das cargas dos fatores, das comunalidades, das variâncias específicas e a proporção total da variação explicada por cada fator para as soluções do modelo MFO com $m = 1$ e $m = 2$.

[Ver arquivo exemplo_22.r](#)

Variável	$m = 1$		$m = 2$		
	$\hat{\ell}_1$	$\hat{\psi}$	$\hat{\ell}_1$	$\hat{\ell}_2$	$\hat{\psi}$
JP	-0,732	0,464	-0,732	0,437	0,273
Citi	-0,831	0,309	-0,831	0,280	0,230
Fargo	-0,726	0,473	-0,726	0,374	0,333
Royal	-0,605	0,634	-0,605	-0,694	0,153
Exxonl	-0,563	0,683	-0,563	-0,719	0,166
Prop. Acum.	0,487		0,487	0,769	

A matriz residual correspondente a solução $m = 2$ é:

$$R - \hat{L}\hat{L}^T - \hat{\Psi} = \begin{bmatrix} 0,000 & -0,099 & -0,185 & -0,025 & 0,056 \\ -0,099 & 0,000 & -0,134 & 0,014 & -0,054 \\ -0,185 & -0,134 & 0,000 & 0,003 & 0,006 \\ -0,025 & 0,014 & 0,003 & 0,000 & -0,156 \\ 0,056 & -0,054 & 0,006 & -0,156 & 0,000 \end{bmatrix}$$

A proporção da variação total explicada por dois fatores é bem maior do que a explicada quando se considera o modelo a um fator. No entanto, o modelo a dois fatores parece produzir correlações amostrais maiores. Este é particularmente o caso de r_{13} .

Parece bastante claro que o primeiro fator representa as condições econômicas gerais e pode ser chamado de fator de mercado. Todas as cargas das ações são altas sobre este fator, e as cargas são praticamente iguais.

O segundo fator contrasta as ações do setor bancário com as ações de petróleo. Assim, este fator parece diferenciar ações em diferentes ramos de atividade e pode ser chamado de fator de atividade.

Resumindo, as taxas de retorno parecem ser determinadas pelas condições gerais de mercado e atividades que são exclusivas para diferentes ramos de atividade, bem como um fator residual específico relacionado a cada empresa.

Método da máxima verossimilhança

$$\mathbf{F} \sim \mathcal{N}_m(\mathbf{0}, \mathbf{I}_m), \quad \varepsilon \sim \mathcal{N}_p(\mathbf{0}, \Psi) \quad \text{e} \quad \mathbf{X} \sim \mathcal{N}_p(\mu, \Sigma).$$

Temos,

$$\begin{aligned} L(\mu, \Sigma) &= (2\pi)^{-np/2} |\Sigma|^{-n/2} \\ &\times \exp \left\{ -\frac{1}{2} \text{tr} \left(\Sigma^{-1} \left[\sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{x}})(\mathbf{x}_i - \bar{\mathbf{x}})^\top \right] \right) \right\} \\ &\times \exp \left\{ -\frac{1}{2} \text{tr} \left(\Sigma^{-1} \left[n(\bar{\mathbf{x}} - \mu)(\bar{\mathbf{x}} - \mu)^\top \right] \right) \right\} \\ &= (2\pi)^{-(n-1)p/2} |\Sigma|^{-(n-1)/2} \\ &\times \exp \left\{ -\frac{1}{2} \text{tr} \left(\Sigma^{-1} \left[\sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{x}})(\mathbf{x}_i - \bar{\mathbf{x}})^\top \right] \right) \right\} \\ &\times (2\pi)^{-p/2} |\Sigma|^{-1/2} \exp \left\{ -\frac{n}{2} (\bar{\mathbf{x}} - \mu) \Sigma^{-1} (\bar{\mathbf{x}} - \mu) \right\}. \end{aligned}$$

que depende de $\mathbf{L}$ e Ψ em função de $\Sigma = \mathbf{L}\mathbf{L}^\top + \Psi$. Este modelo ainda não está bem definido devido à multiplicidade de escolhas para $\mathbf{L}$ que são possíveis pelas transformações ortogonais da solução.

É desejável tornar $\mathbf{L}$ bem definida se impondo uma condição de unicidade que seja computacionalmente conveniente:

$$\mathbf{L}^T \boldsymbol{\Psi} \mathbf{L} = \boldsymbol{\Delta} \quad \text{com} \quad \boldsymbol{\Delta} \quad \text{uma matriz diagonal.}$$

As estimativas de máxima verossimilhança de $\mathbf{L}$ e $\boldsymbol{\Psi}$ dadas por $\tilde{\mathbf{L}}$ e $\tilde{\boldsymbol{\Psi}}$ são obtidas por maximização numérica. Os pacotes estatísticos em geral dispõem de rotinas para calcular tais estimativas.

Proposição

Seja $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$ uma amostra aleatória de $\mathcal{N}_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$, $\boldsymbol{\Sigma} = \mathbf{L}\mathbf{L}^\top + \boldsymbol{\Psi}$ para o MFO com m fatores comuns. Seja $(\tilde{\mathbf{L}}, \tilde{\boldsymbol{\Psi}})$ a estimativa de máxima verossimilhança de $(\mathbf{L}, \boldsymbol{\Psi})$ sob a condição $\mathbf{L}^\top \boldsymbol{\Psi} \mathbf{L} = \boldsymbol{\Delta}$ matriz diagonal. Então, a estimativa de máxima verossimilhança das communalidades são dadas por (soma de quadrados da j -ésima linha de $\mathbf{L}$)

$$\tilde{h}_j^2 = \tilde{\ell}_{j1}^2 + \tilde{\ell}_{j2}^2 + \dots + \tilde{\ell}_{jm}^2 \quad \text{para } j = 1, 2, \dots, p.$$

Assim, a proporção da variação total devida ao k -ésimo fator é dada por

$$\frac{\tilde{\lambda}_k}{\text{tr}(\mathbf{S})} = \frac{\tilde{\boldsymbol{\ell}}_k^\top \tilde{\boldsymbol{\ell}}_k}{\text{tr}(\mathbf{S})} = \frac{\sum_{j=1}^p \tilde{\ell}_{jk}^2}{\text{tr}(\mathbf{S})}, \quad k = 1, 2, \dots, m \quad (\text{Invariância do EMV}).$$

Se usarmos as variáveis padronizadas $\mathbf{Z} = \mathbf{V}^{-1/2}(\mathbf{X} - \boldsymbol{\mu})$, temos

$$\begin{aligned}\mathbf{R} &= \mathbf{V}^{-1/2} \boldsymbol{\Sigma} \mathbf{V}^{-1/2} \\ &= \mathbf{V}^{-1/2} (\mathbf{L} \mathbf{L}^\top + \boldsymbol{\Psi}) \mathbf{V}^{-1/2} \\ &= (\mathbf{V}^{-1/2} \mathbf{L}) (\mathbf{V}^{-1/2} \mathbf{L})^\top + \mathbf{V}^{-1/2} \boldsymbol{\Psi} \mathbf{V}^{-1/2} \\ &= \mathbf{L}_Z \mathbf{L}_Z^\top + \boldsymbol{\Psi}_Z.\end{aligned}$$

Dado $\tilde{\mathbf{L}}_Z$ e $\tilde{\boldsymbol{\Psi}}_Z$ via $\mathbf{R}$, obtemos

$$\begin{aligned}\tilde{\mathbf{L}} &= \tilde{\mathbf{V}}^{1/2} \tilde{\mathbf{L}}_Z \\ \tilde{\boldsymbol{\Psi}} &= \tilde{\mathbf{V}}^{1/2} \tilde{\boldsymbol{\Psi}}_Z \tilde{\mathbf{V}}^{1/2} \\ \tilde{\ell}_{jk} &= \tilde{\ell}_{\mathbf{Z},jk} \sqrt{\tilde{\sigma}_{jk}} \\ \tilde{\psi}_j &= \tilde{\psi}_{\mathbf{Z},j} \tilde{\sigma}_{jj},\end{aligned}$$

com $\tilde{\sigma}_{jj}$ o EMV de σ_{jj} , $j = 1, 2, \dots, p$ e $k = 1, 2, \dots, m$.

Logo, na análise fatorial via método da máxima verossimilhança há equivalência entre decompor $\mathbf{S}$ ou $\mathbf{R}$. O mesmo não se aplica a análise fatorial via componentes principais.

Exemplo

Para os dados do mercado de ações, obteremos a solução do MFO via máxima verossimilhança.

Variável	$m = 1$		$m = 2$		
	$\tilde{\ell}_1$	$\tilde{\psi}$	$\tilde{\ell}_1$	$\tilde{\ell}_2$	$\tilde{\psi}$
JP	0,715	0,488	0,121	0,754	0,417
Citi	0,875	0,234	0,328	0,786	0,275
Fargo	0,663	0,560	0,188	0,650	0,542
Royal	0,346	0,880	0,997	-0,007	0,005
Exxonl	0,279	0,922	0,685	0,026	0,530
Prop. Acum.	0,383		0,324	0,646	

Ver arquivo exemplo_23.r

A matriz residual correspondente a solução $m = 2$ é:

$$R - \tilde{L}\tilde{L}^T - \tilde{\Psi} =$$

$$\begin{bmatrix} 0,00000011 & 0,00000750 & -0,00256422 & -0,00033256 & 0,05198222 \\ 0,00000750 & 0,00000003 & 0,00160887 & 0,00021162 & -0,03307885 \\ -0,00256422 & 0,00160887 & 0,00000005 & -0,00000952 & 0,00055472 \\ -0,00033256 & 0,00021162 & -0,00000952 & -0,00000156 & 0,00012189 \\ 0,05198222 & -0,03307885 & 0,00055472 & 0,00012189 & 0,00000027 \end{bmatrix}$$

Os elementos da matriz residual via máxima verossimilhança são bem menores do que os obtidos via método de componentes principais.

Portanto, a abordagem da máxima verossimilhança é mais apropriada.

A porcentagem da variação acumulada explicada pelos fatores é maior para os fatores via componentes principais do que para os fatores via máxima verossimilhança. Isto não é uma surpresa, pois o método componentes principais justamente tem como critério escolher o fator com maior proporção da variação total. As cargas obtidas via componentes principais estão relacionadas aos componentes principais que têm, por planejamento, uma propriedade de otimização da variância.

Verificando a solução via máxima verossimilhança, observamos que todas as cargas do fator 1 são positivas. Chamamos este fator de fator de mercado. A interpretação do segundo fator não é tão clara quanto a interpretação sugerida na solução via componentes principais. As ações dos bancos têm cargas grandes e positivas, enquanto que as ações de petróleo têm cargas desprezíveis.

A partir desta perspectiva, o segundo fator diferencia as ações de bancos das ações de petróleo e pode ser chamado fator de ramo de atividade. Alternativamente, o segundo fator pode ser simplesmente chamado “fator de banco”.

Os padrões das cargas iniciais dos fatores para a solução de máxima verossimilhança estão restritos pela condição de unicidade ($\tilde{\mathbf{L}}^T \tilde{\Psi} \tilde{\mathbf{L}}$ é diagonal). Portanto, padrões úteis não serão revelados, em geral, até que os fatores sejam rotacionados.

Teste para o número de fatores comuns para grandes amostras

A suposição de normalidade leva diretamente a um teste da adequação do modelo a m fatores.

Suponha que m fatores comuns foram considerados no ajuste.

Neste caso $\Sigma = \mathbf{L}\mathbf{L}^\top + \Psi$. Assim, testar a adequação de m fatores é equivalente a testar

$$\left\{ \begin{array}{ll} H_0 : & \Sigma = \underbrace{\mathbf{L}}_{p \times m} \mathbf{L}^\top + \underbrace{\Psi}_{\text{matriz diagonal}} \\ H_1 : & \Sigma \text{ é qualquer outra matriz positiva definida.} \end{array} \right.$$

Quando Σ não tem forma especial, o valor máximo da função de verossimilhança é atingido quando

$$\tilde{\mu} = \bar{\mathbf{x}} \quad \text{e} \quad \tilde{\Sigma} = \mathbf{S}_n = \frac{n-1}{n} \mathbf{S},$$

em que

$$L(\mu, \Sigma) \propto |\mathbf{S}_n|^{-n/2} \exp \left\{ -\frac{np}{2} \right\}.$$

Sob H_0 , o valor máximo da função de verossimilhança é atingido quando

$$\tilde{\boldsymbol{\mu}} = \bar{\mathbf{x}} \quad \text{e} \quad \tilde{\boldsymbol{\Sigma}} = \tilde{\mathbf{L}}\tilde{\mathbf{L}}^{\top} + \tilde{\boldsymbol{\Psi}},$$

em que

$$L(\boldsymbol{\mu}, \boldsymbol{\Sigma}) = L(\boldsymbol{\mu}, \mathbf{L}, \boldsymbol{\Psi}) \propto |\tilde{\mathbf{L}}\tilde{\mathbf{L}}^{\top} + \tilde{\boldsymbol{\Psi}}|^{-n/2} \exp \left\{ -\frac{n}{2} \text{tr} \left(\left[\tilde{\mathbf{L}}\tilde{\mathbf{L}}^{\top} + \tilde{\boldsymbol{\Psi}} \right]^{-1} \mathbf{S}_n \right) \right\}.$$

Segue, após diversas simplificações que

$$-2 \ln \boldsymbol{\Lambda} = n \ln \left(\frac{|\tilde{\mathbf{L}}\tilde{\mathbf{L}}^{\top} + \tilde{\boldsymbol{\Psi}}|}{|\mathbf{S}_n|} \right) \quad \text{com} \quad \frac{(p-m)^2 - p - m}{2},$$

graus de liberdade.

Note que sob H_0 ,

$$\tilde{\boldsymbol{\Sigma}} = \mathbf{S}_n = \tilde{\mathbf{L}}\tilde{\mathbf{L}}^{\top} + \tilde{\boldsymbol{\Psi}} \Rightarrow \text{tr} \left(\left[\tilde{\mathbf{L}}\tilde{\mathbf{L}}^{\top} + \tilde{\boldsymbol{\Psi}} \right]^{-1} \mathbf{S}_n \right) = \text{tr} (\mathbf{I}_p) = p.$$

Bartlett sugeriu a seguinte correção

$$(n - 1 - (2p + 4m + 5)/6) \ln \left(\frac{|\tilde{\mathbf{L}}\tilde{\mathbf{L}}^\top + \tilde{\Psi}|}{|\mathbf{S}_n|} \right) \stackrel{a}{\sim} \chi_\nu^2,$$

em que $\nu = \frac{(p - m)^2 - p - m}{2}$ para n e $n - p$ grandes.

Como $\nu > 0$, segue que $m < \frac{1}{2}(2p + 1 - \sqrt{8p + 1})$ de modo a poder aplicar este teste.

Observe que na implementação deste teste para verificar se m é adequado para o número de fatores comuns, estamos comparando as variâncias generalizadas $|\tilde{\mathbf{L}}\tilde{\mathbf{L}}^\top + \tilde{\Psi}|$ e $|\mathbf{S}_n|$. Se n é grande e m é relativamente pequeno comparado a p , H_0 será comumente rejeitada, levando à retenção de mais fatores comuns. Porém, $\tilde{\Sigma} = \tilde{\mathbf{L}}\tilde{\mathbf{L}}^\top + \tilde{\Psi}$ pode ser bem parecida com $\mathbf{S}_n$ tal que adicionar mais fatores não fornecerá ganhos, mesmo que os novos fatores sejam “significantes”. Algum julgamento deve ser exercido na escolha de m (por exemplo a proporção acumulada da variabilidade total).

Exemplo (continuação)

Para os dados do mercado de ações, faremos o teste de adequação de $m = 2$ da solução do MFO via máxima verossimilhança.

Primeiro, note para $n = 103$ e $p = 5$ devemos ter m menor que

$$\frac{1}{2}(2p + 1 - \sqrt{8p + 1}) = 2,298.$$

Logo, podemos testar a adequação para $m = 1$ e $m = 2$.

Para $m = 1$, a estatística de qui-quadrado é 62,22 com 5 grau de liberdade. O valor p é 10^{-12} . Logo, rejeitamos H_0 .

Para $m = 2$, a estatística de qui-quadrado é 2,00 com 1 grau de liberdade. O valor p é 0,16. Logo, não rejeitamos H_0 .

Nota: os resultados apresentados acima diferem pouca coisa dos resultados obtidos via função `factanal` do R. Isto se deve aos cálculos numéricos e suas possíveis aproximações.

Rotação de fatores

Vimos que todas as cargas são obtidas das cargas iniciais por uma transformação ortogonal com a mesma habilidade em reproduzir a matriz de covariâncias (correlações).

Da álgebra matricial, sabemos que uma transformação ortogonal corresponde a uma rotação (ou reflexão) dos eixos coordenados em relação à origem. Por esta razão, a técnica de realizar uma transformação ortogonal das cargas para obter as cargas rotacionadas é chamada **rotação de fatores**.

Se $\tilde{\mathbf{L}}_{p \times m}$ é a matriz das cargas estimadas, então $\tilde{\mathbf{L}}^* = \tilde{\mathbf{L}}\mathbf{Q}$, com $\mathbf{Q}\mathbf{Q}^\top = \mathbf{Q}^\top\mathbf{Q} = \mathbf{I}_m$, é uma matriz de cargas rotacionadas.

Além disso, a matriz de covariâncias (correlações) estimada se mantém inalterada, pois

$$\tilde{\mathbf{L}}\tilde{\mathbf{L}}^\top + \tilde{\Psi} = \tilde{\mathbf{L}}\mathbf{Q}\mathbf{Q}^\top\tilde{\mathbf{L}}^\top + \tilde{\Psi} = \tilde{\mathbf{L}}^*\tilde{\mathbf{L}}^{*\top} + \tilde{\Psi}.$$

A matriz residual também não se altera, pois

$$\mathbf{S}_n - \tilde{\mathbf{L}}\tilde{\mathbf{L}}^\top - \tilde{\Psi} = \mathbf{S}_n - \tilde{\mathbf{L}}^*\tilde{\mathbf{L}}^{*\top} - \tilde{\Psi}$$

e as estimativas das comunalidades e variâncias específicas também não se alteram.

Do ponto de vista matemático, tanto faz usar $\tilde{\mathbf{L}}$ como usar $\tilde{\mathbf{L}}^*$.

Como as cargas das variáveis sobre os fatores comuns podem não ser facilmente interpretáveis, é usual rotacioná-las até que uma “estrutura mais simples” seja alcançada.

O ideal seria poder obter um padrão de cargas tais que a carga de cada variável seja bem grande em um único fator e tenha valores pequenos a moderados nos demais fatores.

Nem sempre será possível obter tal “estrutura simples”.

Concentraremos-nos em métodos gráficos e analíticos para determinar a rotação ortogonal visando uma estrutura mais simples das cargas.

Quando $m = 2$ ou os fatores comuns são considerados 2 de cada vez, uma transformação para uma estrutura simples pode ser determinada graficamente.

Fatores comuns não correlacionados são olhados como vetores unitários no sistema de eixos coordenados ortogonais representando as variáveis originais.

Um gráfico de dispersão dos pares de cargas dos fatores produz p pontos, $(\tilde{\ell}_{j1}, \tilde{\ell}_{j2})$, $j = 1, 2, \dots, p$, correspondentes a cada variável.

Os eixos coordenados podem ser visualmente rotacionados por um ângulo ϕ e as novas cargas rotacionadas $\tilde{\ell}_{jq}^*$, $j = 1, 2, \dots, p$, $q = 1, 2$, são determinadas pelas relações ($\tilde{\mathbf{L}}^* = \tilde{\mathbf{L}}\mathbf{Q}$)

$$\begin{aligned}\mathbf{Q} &= \begin{bmatrix} \cos \phi & \text{sen} \phi \\ -\text{sen} \phi & \cos \phi \end{bmatrix} && \text{sentido horário} \\ \mathbf{Q} &= \begin{bmatrix} \cos \phi & -\text{sen} \phi \\ \text{sen} \phi & \cos \phi \end{bmatrix} && \text{sentido anti-horário.}\end{aligned}$$

Estas relações raramente são implementadas nas análises gráficas visuais a duas dimensões.

Nesta situação, conglomerados de variáveis são frequentemente percebidos no diagrama de dispersão das cargas e estes conglomerados nos permitem identificar os fatores comuns sem precisar calcular o ângulo ϕ da rotação.

Por outro lado, para $m > 2$, orientações não são facilmente visualizadas, e as magnitudes das cargas rotacionadas devem ser investigadas para encontrar uma interpretação com algum significado dos dados originais. A escolha de uma matriz ortogonal $\mathbf{Q}$ que satisfaz uma certa medida analítica de estrutura simples será brevemente considerada.

O critério **varimax** é uma medida analítica para a rotação de fatores.

Defina $\check{\ell}_{jk}^* = \frac{\tilde{\ell}_{jk}^*}{\tilde{h}_j}$ como os **coeficientes rotacionados** corrigidos pelo inverso da raiz quadrada das comunalidades estimadas.

O procedimento varimax, também conhecido como **varimax normal**, seleciona a transformação ortogonal **Q** que maximiza

$$\eta(\check{\mathbf{L}}^*) = \frac{1}{p} \sum_{k=1}^m \left[\sum_{j=1}^p \check{\ell}_{jk}^{*4} - \left(\sum_{j=1}^p \check{\ell}_{jk}^{*2} \right)^2 / p \right] \text{ com } \check{\mathbf{L}}^* = \begin{bmatrix} \check{\ell}_1^* & \check{\ell}_2^* & \cdots & \check{\ell}_m^* \end{bmatrix}.$$

Os coeficientes rotacionados $\check{\ell}_{jk}^*$ têm o efeito de dar as variáveis com comunalidades relativamente pequenas mais peso na determinação da estrutura simples.

Depois de determinar **Q**, as cargas $\check{\ell}_{jk}^*$ são multiplicadas por $\tilde{h}_j$ tal que as comunalidades estimadas sejam preservadas, isto é, $\tilde{\ell}_{jk}^* = \check{\ell}_{jk}^* \tilde{h}_j$.

Uma interpretação simples para $\eta(\check{\mathbf{L}}^*)$ é que ela é proporcional a soma das variâncias dos quadrados das cargas corrigidas pelo inverso da raiz quadrada das comunalidades para o k -ésimo fator.

Efetivamente, maximizar $\eta(\check{\mathbf{L}}^*)$ corresponde a espalhar os quadrados das cargas sobre cada fator tanto quanto for possível. Portanto, espera-se encontrar grupos de variáveis com cargas grandes e grupos com cargas pequenas para cada fator após a rotação.

Algoritmos computacionais para a maximização de $\eta(\check{\mathbf{L}}^*)$ costumam estar disponíveis. No **R** podemos utilizar a função ou o argumento **varimax**.

Exemplo (continuação)

Rotação das cargas para os dados sobre preferência dos consumidores.

As cargas originais foram obtidas pelo método de componentes principais.

Variável	$\hat{\ell}_1$	$\hat{\ell}_2$
Z_1	-0,560	0,816
Z_2	-0,777	-0,524
Z_3	-0,645	0,748
Z_4	-0,939	-0,105
Z_5	-0,798	-0,543

Ver arquivo exemplo_24.r

Após inspeção visual é fácil identificar novos eixos ortogonais que produzem uma estrutura simples.

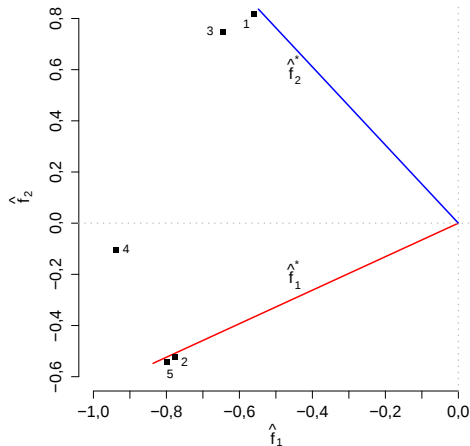


Figura 10: Rotação de fatores via varimax.

Variável	$\hat{\ell}_1$	$\hat{\ell}_2$	$\hat{\ell}_1^*$	$\hat{\ell}_2^*$
Z_1	-0,560	0,816	-0,021	0,989
Z_2	-0,777	-0,524	-0,937	-0,013
Z_3	-0,645	0,748	-0,130	0,979
Z_4	-0,939	-0,105	-0,843	0,427
Z_5	-0,798	-0,543	-0,965	-0,017

É claro que as variáveis 2, 4 e 5 definem o fator 1 (cargas altas neste fator e desprezíveis no fator 2), enquanto que as variáveis 1 e 3 definem o fator 2. Podemos chamar o fator 1 de fator **nutricional** e o fator 2 de fator de **paladar**.

A matriz de rotação varimax é

$$\mathbf{Q} = \begin{bmatrix} 0,837 & -0,548 \\ 0,548 & 0,837 \end{bmatrix}.$$

Exemplo (continuação)

Rotação varimax - mercado de ações - para obter as cargas rotacionadas dos fatores obtidos via método de máxima verossimilhança.

	$\tilde{\ell}_1$	$\tilde{\ell}_2$	$\tilde{\ell}_1^*$	$\tilde{\ell}_2^*$
JP	0,121	0,754	0,763	0,029
Citi	0,328	0,786	0,819	0,232
Fargo	0,188	0,650	0,668	0,108
Royal	0,997	-0,007	0,113	0,991
Exxon	0,685	0,026	0,108	0,677

As cargas rotacionadas indicam que as ações de bancos apresentam cargas altas no fator 1, enquanto que as ações de petróleo apresentam cargas altas no fator 2.

Os dois fatores rotacionados diferenciam os ramos de atividade. Fica difícil rotular estes fatores de forma inteligente. O fator 1 representa aquelas forças econômicas únicas que acarretam na movimento comum das ações dos bancos. O fator 2 parece representar condições econômicas que afetam ações de petróleo.

Ver arquivo exemplo_25.r

Como já foi observado, um fator geral (isto é, um fator para o qual todas as variáveis apresentam cargas altas) tende a ser “destruído” após a rotação. Por esta razão, em casos nos quais um fator geral é evidente, uma rotação ortogonal é algumas vezes realizada, mantendo o fator geral fixado.

Existem outros métodos de rotação ortogonal.

- Método **quartimax**: minimiza o número de fatores necessários para explicar uma variável.
- Método **equimax**: é um compromisso entre os métodos varimax e quartimax.

O método varimax é o mais popular.

Rotações oblíquas

Rotações ortogonais são apropriadas para um modelo fatorial no qual os fatores comuns são supostos ser independentes. Muitas investigações sociais também consideram rotações oblíquas dos fatores.

Se olharmos os m fatores comuns como eixos coordenados, o ponto com as m coordenadas $(\tilde{\ell}_{j1}, \tilde{\ell}_{j2}, \dots, \tilde{\ell}_{jm})$ representa a posição da j -ésima variável no espaço dos fatores. Supondo que as variáveis estão agrupadas em conglomerados disjuntos, uma rotação ortogonal para uma estrutura simples corresponde à uma rotação rígida dos eixos coordenados tal que os eixos, após a rotação, passam tão perto quanto possível destes conglomerados.

Uma rotação oblíqua para uma estrutura simples corresponde a uma rotação não rígida do sistema de coordenadas tal que os eixos rotacionados (não mais ortogonais) passam (quase) pelos conglomerados. Uma rotação oblíqua busca expressar cada variável em função de um número mínimo de fatores, preferencialmente um único fator.

Para rotações oblíquas, o método mais popular é o **promax** que tem a vantagem de ser rápido e conceitualmente simples. O método busca ajustar uma matriz alvo que tenha uma estrutura simples.

Escores dos fatores

Na análise fatorial, o interesse está comumente centrado nos parâmetros do modelo fatorial. Porém, os valores estimados dos fatores comuns, chamados **escores dos fatores**, podem também ser necessários para as análises.

Estas quantidades são em geral usadas para propósitos de diagnósticos, bem como entradas para análises subsequentes. Os escores dos fatores não são estimativas de parâmetros desconhecidos no sentido usual da estimação.

Ao contrário, eles são “estimativas” dos valores para os vetores fatoriais aleatórios não observáveis $\mathbf{F}_i$, $i = 1, 2, \dots, m$. Isto é, os escores $\tilde{\mathbf{F}}_i$ são as estimativas dos valores de $\mathbf{f}_i$ atingidos por $\mathbf{F}_i$, $i = 1, 2, \dots, n$.

Este problema de estimação é complicado, pois as quantidades não observáveis $\mathbf{F}_i$ e ε_i superam em número as observações $\mathbf{x}_i$, $i = 1, 2, \dots, n$.

Para lidar esta dificuldade, algumas abordagens, um tanto heurísticas, mas fundamentadas, para o problema de estimação dos valores dos fatores foram desenvolvidas. Apresentaremos duas delas. Ambas têm dois elementos em comum.

- (1) tratam as estimativas das cargas dos fatores $\tilde{\ell}_{jk}$ e das variâncias específicas $\tilde{\psi}_j$ como se fossem os valores verdadeiros.
- (2) envolvem transformações lineares dos dados originais, centradas ou padronizadas. Comumente, as cargas rotacionadas estimadas são usadas para calcular os escores dos fatores. As fórmulas computacionais são substituídas por cargas não rotacionadas e, portanto, não iremos diferenciar entre elas.

Método dos mínimos quadrados ponderados

Modelo: $\mathbf{X} - \boldsymbol{\mu} = \mathbf{LF} + \boldsymbol{\varepsilon}$.

A soma de quadrados dos erros ponderados pela recíproca de suas variâncias é dada por

$$S^2(\mathbf{F}) = \sum_{j=1}^p \frac{\varepsilon_j^2}{\psi_j} = \boldsymbol{\varepsilon}^\top \boldsymbol{\Psi}^{-1} \boldsymbol{\varepsilon} = (\mathbf{x} - \boldsymbol{\mu} - \mathbf{LF}) \boldsymbol{\Psi}^{-1} (\mathbf{x} - \boldsymbol{\mu} - \mathbf{LF})$$

Bartlett propôs escolher as estimativas $\tilde{\mathbf{F}}$ de $\mathbf{F}$ que minimizam $S^2(\mathbf{F})$.

Escore dos fatores obtidos por MQP a partir das estimativas MV:

$$\tilde{\mathbf{F}}_i = (\tilde{\mathbf{L}}^\top \tilde{\boldsymbol{\Psi}}^{-1} \tilde{\mathbf{L}})^{-1} \tilde{\mathbf{L}}^\top \tilde{\boldsymbol{\Psi}}^{-1} (\mathbf{x}_i - \tilde{\boldsymbol{\mu}}) = \tilde{\boldsymbol{\Delta}}^{-1} \tilde{\mathbf{L}}^\top \tilde{\boldsymbol{\Psi}}^{-1} (\mathbf{x}_i - \tilde{\boldsymbol{\mu}}),$$

$i = 1, 2, \dots, n$ ou, se a matriz de correlação é decomposta,

$$\tilde{\mathbf{F}}_i = (\tilde{\mathbf{L}}_z^\top \tilde{\boldsymbol{\Psi}}_z^{-1} \tilde{\mathbf{L}}_z)^{-1} \tilde{\mathbf{L}}_z^\top \tilde{\boldsymbol{\Psi}}_z^{-1} \mathbf{z}_i = \tilde{\boldsymbol{\Delta}}_z^{-1} \tilde{\mathbf{L}}_z^\top \tilde{\boldsymbol{\Psi}}_z^{-1} \mathbf{z}_i$$

$i = 1, 2, \dots, n$ com $\mathbf{z}_i = \mathbf{V}^{-1/2}(\mathbf{x}_i - \bar{\mathbf{x}})$ e $\tilde{\mathcal{R}} = \tilde{\mathbf{L}}_z \tilde{\mathbf{L}}_z^\top + \tilde{\boldsymbol{\Psi}}_z$.

Observações:

1. Se cargas rotacionadas $\tilde{\mathbf{L}}^* = \tilde{\mathbf{L}}\mathbf{Q}$ são usadas no lugar das cargas originais, os escores dos fatores correspondentes $\tilde{\mathbf{F}}_i^*$ serão dados por $\tilde{\mathbf{F}}_i^* = \mathbf{Q}^\top \tilde{\mathbf{F}}_i$, $k = 1, 2, \dots, n$.
2. Se as cargas dos fatores são estimadas via método dos componentes principais, costuma-se gerar os escores dos fatores se usando um procedimento de mínimos quadrados ordinários (não ponderados). Implicitamente, isto equivale a assumir que as variâncias específicas são iguais ou aproximadamente iguais. Os escores dos fatores, são, então,

$$\tilde{\mathbf{F}}_i = (\hat{\mathbf{L}}^\top \hat{\mathbf{L}})^{-1} \hat{\mathbf{L}}^\top (\mathbf{x}_i - \bar{\mathbf{x}}) \quad \text{para } \Sigma, \quad \text{e}$$

$$\tilde{\mathbf{F}}_i = (\hat{\mathbf{L}}_Z^\top \hat{\mathbf{L}}_Z)^{-1} \hat{\mathbf{L}}_Z^\top \mathbf{z}_i, \quad \text{para } \mathcal{R}.$$

Métodos de regressão

Suponha

$$(\mathbf{X} - \boldsymbol{\mu} | \mathbf{F}) \sim \mathcal{N}_p(\mathbf{0}, \mathbf{L}\mathbf{L}^\top + \boldsymbol{\Psi}) \quad \text{e} \quad \mathbf{F} \sim \mathcal{N}_m(\mathbf{0}, \mathbf{I}_m) \quad \text{tal que}$$

$$\begin{pmatrix} \mathbf{X} - \boldsymbol{\mu} \\ \mathbf{F} \end{pmatrix} \sim \mathcal{N}_{p+m}(\mathbf{0}, \boldsymbol{\Sigma}^*) \quad \text{com}$$

$$\boldsymbol{\Sigma}^* = \begin{bmatrix} \mathbf{L}\mathbf{L}^\top + \boldsymbol{\Psi} & \mathbf{L} \\ \mathbf{L}^\top & \mathbf{I}_m \end{bmatrix}.$$

Portanto,

$$\begin{aligned} \mathbb{E}(\mathbf{F} | \mathbf{x}) &= \mathbf{L}^\top (\mathbf{L}\mathbf{L}^\top + \boldsymbol{\Psi})^{-1} (\mathbf{x} - \boldsymbol{\mu}) \\ \text{Var}(\mathbf{F} | \mathbf{x}) &= \mathbf{I}_m - \mathbf{L}^\top (\mathbf{L}\mathbf{L}^\top + \boldsymbol{\Psi})^{-1} \mathbf{L} \end{aligned}$$

As quantidades em $\mathbf{L}^\top (\mathbf{L}\mathbf{L}^\top + \boldsymbol{\Psi})^{-1}$ correspondem aos coeficientes da regressão multivariada dos fatores sobre as variáveis. Estimativas destes coeficientes produzem os escores.

Escore via método de regressão:

$$\begin{aligned}\tilde{F}_i &= \tilde{\mathbf{L}}^\top \mathbf{S}^{-1}(\mathbf{x}_i - \bar{\mathbf{x}}), \quad i = 1, 2, \dots, n \quad \text{ou} \\ \tilde{F}_i &= \tilde{\mathbf{L}}_Z^\top \mathbf{R}^{-1} \mathbf{z}_i, \quad i = 1, 2, \dots, n.\end{aligned}$$

Uma medida de concordância dos escores gerados pelos dois métodos é fornecida pelo coeficiente de correlação amostral entre os escores de um mesmo fator.

Não há nenhuma indicação de que um dos métodos seja superior ao outro.

Na função **factanal** do **R** está disponível o argumento `scores=c("none", "regression", "Bartlett")`.

Exemplo (continuação)

Estimar os escores dos dois fatores - mercado de ações - pelos métodos de Bartlett e da regressão, usando o modelo ajustado por (máxima verossimilhança) para os dados sobre ações e calcule as correlações entre os escores de um mesmo fator.

A correlação estimada entre os fatores pelos dois métodos é aproximadamente

Regressão	Bartlet	
	$\tilde{F}_1$	$\tilde{F}_2$
$\tilde{F}_1$	1	
$\tilde{F}_2$	0	1

Ver arquivo exemplo_26.r

Perspectivas e estratégias para a análise fatorial

Há muitas decisões que devem ser tomadas em qualquer análise fatorial.

A mais importante é a escolha de m , o número de fatores comuns.

Apesar de um teste para a adequação do modelo para um dado m estar disponível (sob grandes amostras), ele é adequado somente para dados aproximadamente normais. Além disso, o teste provavelmente irá rejeitar o modelo para m pequeno, se o número de variáveis e observações for grande. No entanto, esta é a situação em análise fatorial que fornecerá uma aproximação útil.

Na maioria das vezes, a escolha de m é baseada em uma combinação de:

1. proporção da variância amostral explicada;
2. conhecimento da situação específica em questão; e
3. “razoabilidade” dos resultados (por exemplo, interpretação dos fatores).

A escolha do método de solução e tipo de rotação é uma decisão menos crucial. De fato, as análises fatoriais mais satisfatórias são aquelas nas quais as rotações são experimentadas com mais do que um método e todos os resultados substancialmente confirmam a mesma estrutura de fatores.

Não existe uma estratégia única para a análise fatorial.

Sugere-se a seguinte rotina.

1. Realize uma ACP. Este método é particularmente apropriado para um primeiro olhar sobre os dados. (Não é necessário que **S** ou **R** sejam não singulares.)
 - 1.1 Procure por observações suspeitas construindo os gráficos de dispersão dos escores dos fatores. Também, calcule escores padronizados para cada observação e distâncias quadradas como foi descrito na seção de verificação de normalidade.
 - 1.2 Tente uma rotação varimax.
2. Realize uma análise fatorial via máxima verossimilhança, incluindo uma rotação varimax.

3. Compare as soluções obtidas das duas análises fatoriais.
 - 3.1 Os grupos de cargas são parecidos?
 - 3.2 Construa o gráfico de dispersão dos escores do fatores obtidos por componentes principais versus os obtidos por máxima verossimilhança.
4. Repita os três primeiros passos para outros números m de fatores comuns. Fatores extras necessariamente contribuem na compreensão e interpretação dos dados?
5. Para grandes conjuntos de dados, divida-os ao meio e realize uma análise fatorial em cada metade. Compare os dois resultados entre si e com o obtido a partir do conjunto completo de dados para verificar a estabilidade da solução. (Os dados podem ser divididos se colocando a primeira metade dos casos em um grupo e, a segunda, em outro. Isto poderá revelar mudanças ao longo do tempo, se os dados forem temporais.)

A análise fatorial tem um apelo intuitivo forte para as ciências sociais e comportamentais. Nestas áreas, é natural olhar dados multivariados sobre processos animais, sociais, como manifestações de “traços” não observáveis.

A análise fatorial fornece um modo de explicar a variabilidade observada no comportamento em função destes traços.

Ainda, quando tudo é dito e feito, a análise fatorial permanece muito subjetiva. Os exemplos trabalhados, em comum com a maior parte das publicações, consistem de situações nas quais o **modelo fatorial ortogonal** fornece explicações razoáveis em função de uns poucos fatores interpretáveis.

Na prática, a grande maioria das análises fatoriais não produzem tais resultados claros.

O critério para julgar a qualidade de qualquer análise fatorial não foi bem quantificado.

Exemplo (continuação)

Dividir o conjunto de dados - mercado de ações - em dois subconjuntos, o primeiro contendo as 52 primeiras linhas e, o segundo, contendo as últimas 51 linhas. Rodar o modelo fatorial a dois fatores usando máxima verossimilhança e rotação varimax. Verificar a estabilidade da solução encontrada.

	Completo		Parte 1		Parte 2	
	$\tilde{\ell}_1$	$\tilde{\ell}_2$	$\tilde{\ell}_1$	$\tilde{\ell}_2$	$\tilde{\ell}_1$	$\tilde{\ell}_2$
JP	0,763	0,029	0,799	-0,006	0,741	0,061
Citi	0,819	0,232	0,813	0,258	0,805	0,231
Fargo	0,668	0,108	0,469	0,183	0,918	0,047
Royal	0,113	0,991	0,119	0,990	0,110	0,991
Exxon	0,108	0,677	0,180	0,687	0,081	0,684

Ver arquivo [exemplo_27.r](#)

Discriminação e classificação

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 como:

Discriminação: Descrever gráfica 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

Suponha que se dispõe 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 π_1 e π_2 as duas populações.

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

Se $\mathbf{x}$ vem da população π_1 dizemos que a distribuição caracterizada pela função (de densidade) de probabilidade conjunta de $\mathbf{X}$ é dada por $f_1(\mathbf{x})$, caso contrário, função (de densidade) de probabilidade 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 π_1 e, caso contrário, será classificada em π_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 π_2 em π_1 e vice-versa.

Um bom procedimento de classificação deve resultar em uma 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 π_1 em π_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 Ω 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 π_1 e, $R_2 = \Omega \setminus R_1$ o conjunto de valores para o qual classificamos o objeto como sendo de π_2 .

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

A probabilidade condicional $\Pr(2|1)$ de classificar um objeto de π_1 em π_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 π_2 em π_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:

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

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

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 π_1	Classificado em π_2
π_1	0	$C(2 1)$
π_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

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

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

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

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

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:

(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 se reduz à comparação de densidades tal que se $f_1(\mathbf{x}_0) \geq f_2(\mathbf{x}_0)$, $\mathbf{x}_0$ é classificado em π_1 . Caso, contrário, $\mathbf{x}_0$ é classificado em π_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 π_1 . Caso contrário, classificamos $\mathbf{x}_0$ em π_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 μ_1 , μ_2 e Σ desconhecidos.

Proposição

A regra do CECI mínimo é dada por

$$R_1: \left\{ \mathbf{x} \mid (\mu_1 - \mu_2)^\top \Sigma^{-1} \mathbf{x} - \frac{1}{2} (\mu_1 - \mu_2)^\top \Sigma^{-1} (\mu_1 + \mu_2) \geq \ln \left(\frac{C(1|2)p_2}{C(2|1)p_1} \right) \right\};$$

R_2 : Caso contrário.

Observação: Nas aplicações μ_1 , μ_2 e Σ são desconhecidos, por essa razão, a regra dada pela proposição deve ser modificada.

O estimador de Σ 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:

$$R_1: \left\{ \mathbf{x} \mid (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} \mathbf{x} - \frac{1}{2}(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1}(\bar{\mathbf{x}}_1 + \bar{\mathbf{x}}_2) \geq \ln \left(\frac{C(1|2)p_2}{C(2|1)p_1} \right) \right\};$$

R_2 : Caso contrário.

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 π_1 e π_2 e, então, designar $\mathbf{x}_0$ a π_1 ou a π_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 se comportar 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 se comparando os valores da estatística com o valor de

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

Exemplo

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. Além das três medidas, o biólogo identificou o gênero de cada lagarto *m*-macho, *f*-fêmea. O ojetivo é construir uma regra de classificação de gênero a partir das três medidas usando os dados disponíveis.

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

Tabela 17: Probabilidade a priori dos grupos.

Grupo 1	Grupo 2
0,48	0,52

Como nada foi dito na função do **R**, o mesmo adota prioris iguais às proporções amostrais.

Tabela 18: 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

Ver arquivo [exemplo_28.r](#)

Tabela 19: Coeficientes de discriminação linear.

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

Tabela 20: 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

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

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

usando um argumento completamente diferente. A ideia de Fisher foi transformar as observações multivariadas $\mathbf{x}$ em observações univariadas y tal que os y provenientes de π_1 e de π_2 sejam tão separados quanto possível.

Fisher sugeriu tomar combinações lineares de $\mathbf{x}$ para criar os y , 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 π_1 e os valores $y_{21}, y_{22}, \dots, y_{2n_2}$ para as observações de π_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

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

$$\begin{aligned}\hat{y}_0 &= (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1} \mathbf{x}_0 \\ \hat{m} &= \frac{1}{2}(\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)^\top \mathbf{S}_c^{-1}(\bar{\mathbf{x}}_1 + \bar{\mathbf{x}}_2) = \frac{1}{2}(\bar{y}_1 + \bar{y}_2).\end{aligned}$$

Aloque $\mathbf{x}_0$ a π_1 se $\hat{y}_0 \geq \hat{m}$.

Caso contrário, aloque $\mathbf{x}_0$ a π_2 .

Classificação é uma boa ideia?

Para duas populações, a separação máxima relativa que pode ser obtida se considerando 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 π_1 e π_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 π_1 e π_2 sejam normais multivariadas com uma matriz de covariâncias comum Σ . 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.

$$\frac{f_1(\mathbf{x})}{f_2(\mathbf{x})} = \left(\frac{|\Sigma_2|}{|\Sigma_1|} \right)^{1/2} \exp \left\{ -\frac{1}{2} \text{SQ} \right\},$$

em que

$$\text{SQ} = (\mathbf{x} - \mu_1)^\top \Sigma_1^{-1} (\mathbf{x} - \mu_1) - (\mathbf{x} - \mu_2)^\top \Sigma_2^{-1} (\mathbf{x} - \mu_2).$$

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

$$R_1: -\frac{1}{2}\mathbf{x}^\top(\boldsymbol{\Sigma}_1^{-1} - \boldsymbol{\Sigma}_2^{-1})\mathbf{x} + (\boldsymbol{\mu}_1^\top \boldsymbol{\Sigma}_1^{-1} - \boldsymbol{\mu}_2^\top \boldsymbol{\Sigma}_2^{-1})\mathbf{x} - k \geq \ln \left(\frac{C(1|2) p_2}{C(2|1) p_1} \right);$$

R_2 : caso contrário.

com

$$k = \frac{1}{2} \ln \left(\frac{|\boldsymbol{\Sigma}_2|}{|\boldsymbol{\Sigma}_1|} \right) + \frac{1}{2} (\boldsymbol{\mu}_1^\top \boldsymbol{\Sigma}_1^{-1} \boldsymbol{\mu}_1 - \boldsymbol{\mu}_2^\top \boldsymbol{\Sigma}_2^{-1} \boldsymbol{\mu}_2).$$

As regiões de classificação são quadráticas em $\mathbf{x}$. Quando $\boldsymbol{\Sigma}_1 = \boldsymbol{\Sigma}_2$, o termo quadrático $\mathbf{x}^\top(\boldsymbol{\Sigma}_1^{-1} - \boldsymbol{\Sigma}_2^{-1})\mathbf{x}$ se anula, e as regiões resultantes são aquelas obtidas anteriormente no caso de variâncias iguais.

Proposição

Sob normalidade multivariada com covariâncias desiguais, aloque $\mathbf{x}_0$ a π_1 se

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

Caso contrário, aloque $\mathbf{x}_0$ a π_2 .

Na prática, a regra de classificação acima é implementada se substituindo 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 π_1 se

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

Caso contrário, aloque $\mathbf{x}_0$ a π_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 π_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 ou 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

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,

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

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

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

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

Se $\delta^2 = 2,56$, teremos $\text{TEO} = \Phi(-0,8) = 0,2119$.

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 se calculando 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	
	π_1	π_2
π_1	n_{1c}	$n_{1M} = n_1 - n_{1c}$
π_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 se considerando 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 π_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 π_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 π_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 π_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 π_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 π_j em π_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 π_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 π_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

As regiões de classificação que minimizam o custo esperado de classificação incorreta são definidas por

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

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

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

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

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 π_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 π_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

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

População	Classificação em		
	π_1	π_2	π_3
π_1	$C(1 1) = 0$	$C(2 1) = 10$	$C(3 1) = 50$
π_2	$C(1 2) = 500$	$C(2 2) = 0$	$C(3 2) = 200$
π_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 π_k , $k = 1, 2, 3$ para a qual

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

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

Como o menor valor ocorre para $k = 2$, alocamos $\mathbf{x}_0$ a π_2 .

Se os custos de classificação incorreta fossem todos iguais, designaríamos $\mathbf{x}_0$ a π_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
3	0,7000

Nesse caso, $\mathbf{x}_0$ seria alocada a π_3 .

Classificação com populações normais

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

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

- Aloque $\mathbf{x}_0$ a π_k se

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

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

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

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

Regra da probabilidade total de classificação incorreta mínima - populações normais, covariâncias desiguais

- Aloque $\mathbf{x}_0$ a π_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 μ_k e Σ_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 π_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 π_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 π_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:

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

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

Podemos selecionar $\mathbf{a}$ 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 se criando 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 se comportar tanto pobremente como satisfatoriamente, dependendo das correlações entre as variáveis contínuas e qualitativas. “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. 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 se comportam tão bem quanto à regressão logística e as funções discriminantes discutidas aqui.

Seleção de variáveis

Em algumas aplicações da análise discriminante, os dados estão disponíveis para um grande número de variáveis. Alguns autores 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. 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. 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.