

ECONOMETRIA

Ralph dos Santos Silva

<http://www.ralphsilva.org/econometria.html>

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

Programa da disciplina

- Introdução à econometria.
- Modelo de regressão linear.
- Estimador de mínimos quadrados.
- Teste de hipóteses e seleção de modelos.
- Forma funcional e mudança de estrutura.
- Modelos de regressão não-linear, semi-paramétrico e não-paramétrico.
- Endogeneidade e estimação por variável instrumental.
- Modelo de regressão generalizado e heterocedasticidade.
- Sistema de equações.
- Modelos para dados de painel.
- Métodos de estimação em econometria.
- Método generalizado dos momentos.
- Estimação por máxima verossimilhança.
- Estimação e inferência utilizando métodos de simulação.
- Estimação e inferência bayesiana.

Referências

- Bauwens, L., Lubrano, M. e Richard, J-F. (2000). *Bayesian Inference in Dynamic Econometric Models*.
- Geweke, J. (2005). *Contemporary Bayesian Econometrics and Statistics*.
- Greene, W. H. (2011). *Econometric Analysis*, 7th ed.
- Gujarati, D. (2002). *Basic Econometrics*, 4th ed.
- Hayashi, F. (2000). *Econometrics*.
- Hendry, D. F. e Nielsen, B. (2007). *Econometric Modelling: A Likelihood Approach*.
- Koop, G. (2003). *Bayesian Econometrics*.
- Lancaster, T. (2004). *Introduction to Modern Bayesian Econometrics*.
- Stock, J. H. e Watson, M. W. (2006). *Introduction to Econometrics*.
- Wooldridge, J. M. (2006). *Introductory to Econometrics: A Modern Approach*.

Introdução

Econometria é a área que foca em aplicações de estatística matemática e as ferramentas da inferência estatística na mensuração empírica de relações postuladas pela teoria econômica.

Exemplo - Função de Consumo de Keynes

$$C = f(X) \quad \left\{ \begin{array}{l} C : \text{consumo} \\ X: \text{renda} \end{array} \right.$$

Temos

$$0 < \frac{dC}{dX} < 1,$$

a propensão marginal a consumir.

A propensão média a consumir (PMC), C/X , decai quando a renda sobe

$$\frac{d(C/X)}{dX} = \left(\frac{dC}{dX} - \frac{C}{X} \right) \frac{1}{X} < 0 \quad \Rightarrow \quad \frac{dC}{dX} < \frac{C}{X}.$$

O modelo de função consumo mais simples é o linear (com restrições)

$$C = \alpha + X\beta, \quad \alpha > 0 \quad \text{e} \quad 0 < \beta < 1.$$

- Podemos estudar alguns conjuntos de dados.
- Será verdade que $\alpha > 0$ e $0 < \beta < 1$?
- A relação é estável através do tempo?
- Os parâmetros mudam se considerarmos instantes de tempo diferentes? (Uma mudança na propensão média a economizar.)
- Existem diferenças sistemáticas nas relações para diferentes países? Como poderíamos explicar estas diferenças?
- Existem outros fatores que poderiam melhorar a habilidade do modelo em explicar a relação entre consumo e renda?

Exemplo

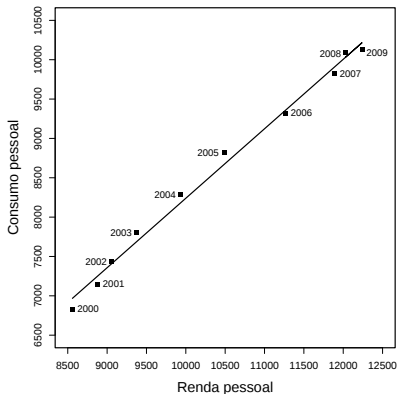


Figura 1: Consumo e renda agregados dos EUA entre 2000 e 2009.

O ajuste por mínimos quadrados resulta em $\hat{\alpha} = -587,19$ e $\hat{\beta} = 0,8827$.

Ver arquivo exemplo_01.r

Para os dados reais, temos uma estimativa negativa para a constante do modelo.

- A teoria está errada?
- Isto se deve a variabilidade dos dados?
- Será que o estimador de mínimos quadrados é adequado nesta situação?
- Como ficaria a estimação por máxima verossimilhança?
- Como ficaria a estimação bayesiana? E a distribuição a priori?

Elasticidade

Em economia, o termo elasticidade se refere a como uma variável econômica muda em função de outra variável econômica.

A elasticidade é medida em termos percentuais ao invés de absoluto.

Isto significa que medimos uma mudança na variável como uma porcentagem da quantidade original da variável.

A mudança em porcentagem da variável x é definida por

$$\text{Mudança percentual em } x = \frac{\text{Mudança em } x}{\text{Valor original de } x} = \frac{\Delta x}{x} \quad (\text{de } x \text{ para } x + \Delta x).$$

A elasticidade de y com respeito a x é dada por

$$\epsilon = \frac{\text{Mudança percentual em } y}{\text{Mudança percentual em } x} = \frac{\Delta y/y}{\Delta x/x}.$$

Em termos infinitesimais, temos que

$$\epsilon = \frac{\partial y/y}{\partial x/x} = \frac{x}{y} \frac{\partial y}{\partial x}.$$

Modelo de regressão linear múltipla

Notação

$$\begin{aligned}y &= f(x_1, x_2, \dots, x_p) + \varepsilon \\ &= x_1\beta_1 + x_2\beta_2 + \dots + x_p\beta_p + \varepsilon,\end{aligned}$$

sendo que

- y é a variável dependente;
- $(x_1, x_2, \dots, x_p)^\top$ é o vetor de variáveis independentes (regressoras, exploratórias ou explicativas);
- $(\beta_1, \beta_2, \dots, \beta_p)^\top$ é o vetor de coeficientes; e
- ε é um distúrbio aleatório.

Para o modelo de regressão acima, a **elasticidade de y com respeito à k -ésima variável x_k** é dada por

$$\epsilon = \frac{x_k}{y} \frac{\partial y}{\partial x_k} = \frac{x_k}{y} \beta_k.$$

Equação da demanda:

$$\text{quantidade} = \beta_1 + \text{preço} \times \beta_2 + \text{renda} \times \beta_3 + \varepsilon.$$

Equação da demanda inversa:

$$\text{preço} = \gamma_1 + \text{quantidade} \times \gamma_2 + \text{renda} \times \beta_3 + \xi.$$

Estas são relações válidas de um mercado.

Ganhos e educação:

$$\text{ganhos} = \beta_1 + \text{educação} \times \beta_2 + \text{idade} \times \beta_3 + \varepsilon,$$

$$\text{ganhos} = \beta_1 + \text{educação} \times \beta_2 + \text{idade} \times \beta_3 + \text{idade}^2 \times \beta_4 + \varepsilon.$$

Hipóteses do modelo de regressão linear

HP.1: Linearidade: $y_i = x_{i1}\beta_1 + x_{i2}\beta_2 + \cdots + x_{ip}\beta_p + \varepsilon_i$;

HP.2: Posto completo: não existe nenhuma relação exata entre as variáveis independentes do modelo;

HP.3: Exogeneidade das variáveis independentes:

$$E(\varepsilon_i | x_{i1}, x_{i2}, \dots, x_{ip}) = 0 \Rightarrow \text{Cov}(\mathbf{x}_i, \varepsilon_i) = 0$$

com $\mathbf{x}_i = (x_{i1}, x_{i2}, \dots, x_{ip})^\top$.

HP.4: Homocedasticidade e autocorrelação nula: cada distúrbio, ε_i , tem a mesma variância, σ^2 , e é não correlacionado com qualquer outro distúrbio ε_j ;

HP.5: Geração dos dados: os dados em $(x_{i1}, x_{i2}, \dots, x_{ip})$ podem ser qualquer mistura de constantes e variáveis aleatórias; e

HP.6: Distribuição normal: os distúrbios são normalmente distribuídos.

Notação

Temos que

$$\begin{aligned}y_i &= x_{i1}\beta_1 + x_{i2}\beta_2 + \cdots + x_{ip}\beta_p + \varepsilon_i \\ &= \mathbf{x}_i^\top \boldsymbol{\beta} + \varepsilon_i, \quad \text{para } i = 1, 2, \dots, n;\end{aligned}$$

- n é o número de observações;
- $\mathbf{x}_i = (x_{i1}, x_{i2}, \dots, x_{ip})^\top$ é o vetor da i -ésima observação das variáveis independentes;
- $\mathbf{x}_k = (x_{k1}, x_{k2}, \dots, x_{kn})^\top$ é o vetor da k -ésima variável independente;
- $\mathbf{y} = (y_1, y_2, \dots, y_n)^\top$ é o vetor da variável dependente;
- $\boldsymbol{\varepsilon} = (\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n)^\top$ é o vetor dos distúrbios;
- $\boldsymbol{\beta} = (\beta_1, \beta_2, \dots, \beta_p)^\top$ é o vetor de coeficientes; e
- $\mathbf{X} = (\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p)$ é a matriz de variáveis independentes.

Assim, o modelo de regressão linear múltipla pode ser reescrito por

$$\begin{aligned}\mathbf{y} &= \mathbf{x}_1\beta_1 + \mathbf{x}_2\beta_2 + \cdots + \mathbf{x}_p\beta_p + \boldsymbol{\varepsilon} \quad \text{ou} \\ \mathbf{y} &= \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}.\end{aligned}$$

Exemplos

- $y = \gamma x^\beta e^\varepsilon \Rightarrow \ln y = \ln \gamma + \beta \ln x + \varepsilon$ é linear, mas
- $y = \gamma x^\beta + \varepsilon$ é não linear.

Os seguintes modelos são lineares:

$$y = \alpha + \beta x + \varepsilon$$

$$y = \alpha + \beta \cos x + \varepsilon$$

$$y = \alpha + \beta \frac{1}{x} + \varepsilon$$

$$y = \alpha + \beta \ln x + \varepsilon$$

$$\ln y = \alpha + \beta \ln x + \varepsilon$$

$$z = \ln \gamma + \beta \ln x + \varepsilon = \alpha + \beta w + \varepsilon.$$

O modelo log-linear

$$\ln y = \beta_1 + \beta_2 \ln x_2 + \cdots + \beta_p \ln x_p + \varepsilon$$

(Elasticidade constante)

A **elasticidade de y com respeito a mudanças em x_k** é dado por

$$\epsilon = \frac{\partial y/y}{\partial x/x} = \frac{\partial y}{\partial x} \times \frac{\partial \ln y / \partial y}{\partial \ln x / \partial x} = \frac{\partial \ln y}{\partial \ln x_k} = \beta_k.$$

Exemplo - mercado de gasolina dos EUA - 1953-2004

Considere o seguinte modelo para o consumo de gasolina per capita:

$$\ln(G/Pop) = \beta_1 + \beta_2 \ln(Renda/Pop) + \beta_3 \ln \text{Preço} + \beta_4 \ln(PCN) + \beta_5 \ln(PCU) + \varepsilon,$$

sendo PCN o preço dos carros novos e PCU o preço dos carros usados.

- Elasticidade renda e preço para a gasolina; e
- Elasticidade para demanda com respeito a preço de carros novos e usados.

Em geral, utiliza-se um modelo **semi-log** para modelar taxas de crescimento:

$$\ln y_t = \mathbf{x}_t^\top \boldsymbol{\beta} + \gamma t + \varepsilon_t \quad \text{com} \quad \partial \ln y_t / \partial t = \gamma,$$

sendo a taxa de crescimento por período.

O modelo translog

- Se $y = g(x_1, x_2, \dots, x_k)$, então podemos transformar para $\ln y = \ln g(x_1, x_2, \dots, x_k) = f(\ln x_1, \ln x_2, \dots, \ln x_k)$, pois $x_k = \exp(\ln x_k)$.
- Daí, aproximar por uma expansão de Taylor de segunda ordem em torno de $\mathbf{x} = \mathbf{1}$. Definimos $\ln \mathbf{x} = (\ln x_1, \ln x_2, \dots, \ln x_k)^\top$. Logo,

$$\begin{aligned} \ln y &= f(\mathbf{0}) + \sum_{k=1}^p \frac{\partial f(\ln \mathbf{x})}{\partial \ln x_k} \Big|_{\ln \mathbf{x}=\mathbf{0}} \ln x_k \\ &+ \frac{1}{2} \sum_{k=1}^p \sum_{\ell=1}^p \frac{\partial^2 f(\ln \mathbf{x})}{\partial \ln x_k \partial \ln x_\ell} \Big|_{\ln \mathbf{x}=\mathbf{0}} \ln x_k \ln x_\ell + \varepsilon. \end{aligned}$$

- Segue-se que

$$\ln y = \beta_0 + \sum_{k=1}^p \beta_k \ln x_k + \frac{1}{2} \sum_{k=1}^p \sum_{\ell=1}^p \gamma_{k\ell} \ln x_k \ln x_\ell + \varepsilon.$$

Distúrbios do modelo de regressão

Temos:

- $E(\varepsilon_i|\mathbf{X}) = 0$ por hipótese, ou seja,

$$E(\varepsilon|\mathbf{X}) = \begin{bmatrix} E(\varepsilon_1|\mathbf{X}) \\ E(\varepsilon_2|\mathbf{X}) \\ \vdots \\ E(\varepsilon_n|\mathbf{X}) \end{bmatrix} = \mathbf{0};$$

- $E(\varepsilon_i|\varepsilon_1, \dots, \varepsilon_{i-1}, \varepsilon_{i+1}, \dots, \varepsilon_n) = 0$;
- Por hipótese, todos os distúrbios são simplesmente uma amostra aleatória de alguma população;
- $E(\varepsilon_i) = E_{\mathbf{X}}(E(\varepsilon_i|\mathbf{X})) = E_{\mathbf{X}}(0) = 0$;
- Para cada ε_i , $\text{Cov}(\varepsilon_i, \mathbf{X}) = \text{Cov}(E(\varepsilon_i|\mathbf{X}), \mathbf{X}) = \text{Cov}(0, \mathbf{X}) = \mathbf{0}$;

Note que $E(\varepsilon_i) = 0$ **não implica** que $E(\varepsilon_i|\mathbf{x}_i) = 0$. Isto é, $E(\varepsilon_i) = 0$ **não implica** que $\text{Cov}(0, \mathbf{X}) = \mathbf{0}$.

- $\mathbf{y} = \mathbf{X}\beta + \varepsilon \Rightarrow E(\mathbf{y}|\mathbf{X}) = E(\mathbf{X}\beta|\mathbf{X}) + E(\varepsilon|\mathbf{X}) = \mathbf{X}\beta + \mathbf{0} = \mathbf{X}\beta$.

Distúrbios esféricos

Temos:

- $\text{Var}(\varepsilon_i|\mathbf{X}) = \sigma^2$, para todo $i = 1, 2, \dots, n$;
- $\text{Cov}(\varepsilon_i, \varepsilon_j|\mathbf{X}) = 0$, para todo $i \neq j$ (correlação nula), ou seja,

$$\begin{aligned} E(\varepsilon\varepsilon^\top|\mathbf{X}) &= \begin{bmatrix} E(\varepsilon_1\varepsilon_1|\mathbf{X}) & E(\varepsilon_1\varepsilon_2|\mathbf{X}) & \cdots & E(\varepsilon_1\varepsilon_n|\mathbf{X}) \\ E(\varepsilon_2\varepsilon_1|\mathbf{X}) & E(\varepsilon_2\varepsilon_2|\mathbf{X}) & \cdots & E(\varepsilon_2\varepsilon_n|\mathbf{X}) \\ \vdots & \vdots & \ddots & \vdots \\ E(\varepsilon_n\varepsilon_1|\mathbf{X}) & E(\varepsilon_n\varepsilon_2|\mathbf{X}) & \cdots & E(\varepsilon_n\varepsilon_n|\mathbf{X}) \end{bmatrix} \\ &= \begin{bmatrix} \sigma^2 & 0 & \cdots & 0 \\ 0 & \sigma^2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \sigma^2 \end{bmatrix} = \sigma^2 \mathbf{I}; \end{aligned}$$

- $\text{Var}(\varepsilon) = E_{\mathbf{X}}(\text{Var}(\varepsilon|\mathbf{X})) + \text{Var}_{\mathbf{X}}(E(\varepsilon|\mathbf{X})) = \sigma^2 \mathbf{I}$;
- Por hipótese, $(\varepsilon|\mathbf{X}) \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$.

Exemplo

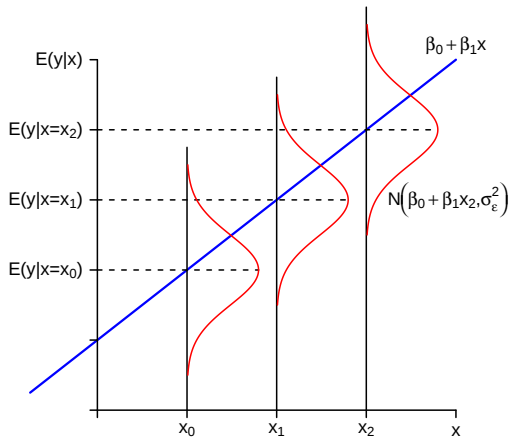


Figura 2: O modelo clássico de regressão.

Ver arquivo exemplo_02.r

Mínimos quadrados

Temos que

$$y_i = \mathbf{x}_i^\top \boldsymbol{\beta} + \varepsilon_i.$$

- $\boldsymbol{\beta}$ e ε_i são quantidades da população;
- A regressão populacional é $E(y_i|\mathbf{x}_i) = \mathbf{x}_i^\top \boldsymbol{\beta}$;
- $\mathbf{b}$ e e_i são estimativas amostrais;
- A estimativa de $E(y_i|\mathbf{x}_i)$ é $\hat{y}_i = \mathbf{x}_i^\top \mathbf{b}$;
- Seja $\varepsilon_i = y_i - \mathbf{x}_i^\top \boldsymbol{\beta}$ o distúrbio da i -ésima observação;
- Para qualquer valor de $\mathbf{b}$, estimamos ε_i com o resíduo

$$e_i = y_i - \mathbf{x}_i^\top \mathbf{b};$$

- Das definições, $y_i = \mathbf{x}_i^\top \boldsymbol{\beta} + \varepsilon_i = \mathbf{x}_i^\top \mathbf{b} + e_i$;
- Podemos estimar $\boldsymbol{\beta}$ com a amostra de $(y_i, \mathbf{x}_i)$, $i = 1, 2, \dots, n$;
- Na estimação de $\boldsymbol{\beta}$ por $\mathbf{b}$, queremos que os pontos observados fiquem perto da reta ajustada;
- Perto neste caso é definido por algum critério de ajuste;
- **Critério de mínimos quadrados.**

Temos

$$S(\mathbf{b}) = \sum_{i=1}^n e_i^2 = \sum_{i=1}^n (y_i - \mathbf{x}_i^\top \mathbf{b})^2,$$

com $\mathbf{b}$ denotando a escolha do vetor de coeficientes.

Queremos minimizar $S(\mathbf{b})$, ou seja,

$$\min_{\mathbf{b}} S(\mathbf{b}) = \min_{\mathbf{b}} \mathbf{e}^\top \mathbf{e} = \min_{\mathbf{b}} (\mathbf{y} - \mathbf{X}\mathbf{b})^\top (\mathbf{y} - \mathbf{X}\mathbf{b}).$$

Mas

$$\begin{aligned} S(\mathbf{b}) &= \mathbf{y}^\top \mathbf{y} - \mathbf{b}^\top \mathbf{X}^\top \mathbf{y} - \mathbf{y}^\top \mathbf{X}\mathbf{b} + \mathbf{b}^\top \mathbf{X}^\top \mathbf{X}\mathbf{b} \\ &= \mathbf{y}^\top \mathbf{y} - 2\mathbf{y}^\top \mathbf{X}\mathbf{b} + \mathbf{b}^\top \mathbf{X}^\top \mathbf{X}\mathbf{b}. \end{aligned}$$

A condição necessária para um mínimo é

$$\frac{\partial S(\mathbf{b})}{\partial \mathbf{b}} = -2\mathbf{X}^\top \mathbf{y} + 2\mathbf{X}^\top \mathbf{X}\mathbf{b} = \mathbf{0}.$$

Seja $\mathbf{b}$ a solução. Então, $\mathbf{b}$ satisfaz as equações normais de mínimos quadrados:

$$\mathbf{X}^\top \mathbf{X} \mathbf{b} = \mathbf{X}^\top \mathbf{y}.$$

Se a inversa de $\mathbf{X}^\top \mathbf{X}$ existir, que segue pela hipótese de posto completo de $\mathbf{X}$, então a solução é

$$\mathbf{b} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}.$$

Para esta solução minimizar a soma de quadrados,

$$\frac{\partial^2 S(\mathbf{b})}{\partial \mathbf{b} \partial \mathbf{b}^\top} = 2\mathbf{X}^\top \mathbf{X}$$

deve ser uma matriz positiva definida.

Seja $q = \mathbf{c}^\top \mathbf{X}^\top \mathbf{X} \mathbf{c}$ para algum vetor arbitrário $\mathbf{c}$ diferente de zero. Então,

$$q = \mathbf{v}^\top \mathbf{v} = \sum_{i=1}^n v_i^2 \quad \text{sendo} \quad \mathbf{v} = \mathbf{X} \mathbf{c}.$$

A menos que cada elemento de $\mathbf{v}$ seja zero, q é positivo.

Exemplo

$$y_i = \beta_1 + \beta_2 T_i + \beta_3 G_i + \beta_4 R_i + \beta_5 P_i + \varepsilon_i$$

Tabela 1: Observações anuais de 1968 a 1982.

Investimento Real (<i>Y</i>)	Constante (1)	Tendência (<i>T</i>)	PIB Real (<i>G</i>)	Taxa de Juros (<i>R</i>)	Taxa de Inflação (<i>P</i>)
0,161	1	1	1,058	5,16	4,40
0,172	1	2	1,088	5,87	5,15
0,158	1	3	1,086	5,95	5,37
0,173	1	4	1,122	4,88	4,99
0,195	1	5	1,186	4,50	4,16
0,217	1	6	1,254	6,44	5,75
0,199	1	7	1,246	7,83	8,82
0,163	1	8	1,232	6,25	9,31
0,195	1	9	1,298	5,50	5,21
0,231	1	10	1,370	5,46	5,83
0,257	1	11	1,439	7,46	7,40
0,259	1	12	1,479	10,28	8,64
0,225	1	13	1,474	11,77	9,31
0,241	1	14	1,503	13,42	9,44
0,204	1	15	1,475	11,02	5,99

Ver arquivo exemplo_03.r

Teorema (regressão ortogonal)

Se as variáveis em um modelo de regressão são não correlacionadas, isto é, são ortogonais, então os coeficientes angulares da regressão são os mesmos que os coeficientes na regressão individual simples.

Aspectos algébricos da solução de mínimos quadrados

- As equações normais são

$$\mathbf{X}^\top \mathbf{X} \mathbf{b} - \mathbf{X}^\top \mathbf{y} = -\mathbf{X}^\top (\mathbf{y} - \mathbf{X} \mathbf{b}) = -\mathbf{X}^\top \mathbf{e} = \mathbf{0}.$$

- Para cada coluna $\mathbf{x}_k$ de $\mathbf{X}$, temos $\mathbf{x}_k^\top \mathbf{e} = 0$.
- Se a primeira coluna de $\mathbf{X}$ for de 1's, então
 1. os resíduos de mínimos quadrados somam zero,

$$\mathbf{x}_1^\top \mathbf{e} = \mathbf{1}^\top \mathbf{e} = \sum_{i=1}^n e_i = 0;$$

2. a regressão no hiperplano passa pelo ponto das médias dos dados; a primeira equação normal implica que $\bar{y} = \bar{\mathbf{x}}^\top \mathbf{b}$;
3. a média dos valores ajustados da regressão é igual a média dos valores atuais.

Isto só vale se a regressão tiver uma constante.

Projeção

O valor dos resíduos de mínimos quadrados é dado por

$$\begin{aligned}\mathbf{e} &= \mathbf{y} - \mathbf{X}\mathbf{b} \\ &= \mathbf{y} - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} \\ &= \mathbf{y} - \mathbf{P}\mathbf{y} \\ &= \mathbf{I}\mathbf{y} - \mathbf{P}\mathbf{y} \\ &= (\mathbf{I} - \mathbf{P})\mathbf{y} \\ &= \mathbf{M}\mathbf{y},\end{aligned}$$

sendo $\mathbf{P} = \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$ e $\mathbf{M} = \mathbf{I} - \mathbf{P}$.

Temos que $\mathbf{P}$ é a matriz de projeção (leva $\mathbf{y}$ em $\hat{\mathbf{y}} = \mathbf{X}\mathbf{b}$):

$$\mathbf{P} = \mathbf{P}^\top, \quad \text{e} \quad \mathbf{P} = \mathbf{P}^2.$$

E $\mathbf{M}$ é uma matriz importante pois leva $\mathbf{y}$ em $\mathbf{e}$:

$$\mathbf{M} = \mathbf{M}^\top, \quad \mathbf{M} = \mathbf{M}^2, \quad \text{e} \quad \mathbf{M}\mathbf{X} = \mathbf{0}.$$

Toda matriz simétrica e idempotente é positiva definida, em particular $\mathbf{P}$ e $\mathbf{M}$.

A matriz de projeção

$$\begin{aligned} \mathbf{y} &= \mathbf{X}\mathbf{b} + \mathbf{e} \quad \Leftrightarrow \quad \mathbf{y} = \hat{\mathbf{y}} + \mathbf{e} \quad \Rightarrow \\ \hat{\mathbf{y}} &= \mathbf{y} - \mathbf{e} = (\mathbf{I} - \mathbf{M})\mathbf{y} = \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = \mathbf{P}\mathbf{y}. \end{aligned}$$

$$\mathbf{P}\mathbf{M} = \mathbf{M}\mathbf{P} = \mathbf{0}, \quad \text{e} \quad \mathbf{P}\mathbf{X} = \mathbf{X}.$$

$$\mathbf{y} = \mathbf{P}\mathbf{y} + \mathbf{M}\mathbf{y} = \text{projeção} + \text{resíduos}.$$

Podemos reescrever a soma de quadrados da seguinte forma:

$$\begin{aligned} \mathbf{y}^\top \mathbf{y} &= \mathbf{y}^\top \mathbf{P}^\top \mathbf{P} \mathbf{y} + \mathbf{y}^\top \mathbf{M}^\top \mathbf{M} \mathbf{y} \\ &= \hat{\mathbf{y}}^\top \hat{\mathbf{y}} + \mathbf{e}^\top \mathbf{e}. \end{aligned}$$

Outras relações úteis:

$$\begin{aligned} \mathbf{e}^\top \mathbf{e} &= \mathbf{y}^\top \mathbf{M}^\top \mathbf{M} \mathbf{y} = \mathbf{y}^\top \mathbf{M} \mathbf{y} = \mathbf{y}^\top \mathbf{e} = \mathbf{e}^\top \mathbf{y} \\ \mathbf{e}^\top \mathbf{e} &= \mathbf{y}^\top \mathbf{y} - \mathbf{b}^\top \mathbf{X}^\top \mathbf{X} \mathbf{b} = \mathbf{y}^\top \mathbf{y} - \mathbf{b}^\top \mathbf{X}^\top \mathbf{y} = \mathbf{y}^\top \mathbf{y} - \mathbf{y}^\top \mathbf{X} \mathbf{b}. \end{aligned}$$

Regressão particionada e regressão parcial

Suponha que a regressão envolva dois conjuntos de variáveis $\mathbf{X}_1$ e $\mathbf{X}_2$. Logo,

$$\mathbf{y} = \mathbf{X}\beta + \varepsilon = \mathbf{X}_1\beta_1 + \mathbf{X}_2\beta_2 + \varepsilon.$$

Qual é a solução algébrica para $\mathbf{b}_2$?

$$\begin{bmatrix} \mathbf{X}_1^\top \mathbf{X}_1 & \mathbf{X}_1^\top \mathbf{X}_2 \\ \mathbf{X}_2^\top \mathbf{X}_1 & \mathbf{X}_2^\top \mathbf{X}_2 \end{bmatrix} \begin{bmatrix} \mathbf{b}_1 \\ \mathbf{b}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1^\top \mathbf{y} \\ \mathbf{X}_2^\top \mathbf{y} \end{bmatrix}$$

Resolvendo para $\mathbf{b}_1$ (primeira equação):

$$\mathbf{X}_1^\top \mathbf{X}_1 \mathbf{b}_1 + \mathbf{X}_1^\top \mathbf{X}_2 \mathbf{b}_2 = \mathbf{X}_1^\top \mathbf{y}$$

$$(\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{X}_1 \mathbf{b}_1 + (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{X}_2 \mathbf{b}_2 = (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{y}$$

$$\begin{aligned} \mathbf{b}_1 &= (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{y} - (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{X}_2 \mathbf{b}_2 \\ &= (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top (\mathbf{y} - \mathbf{X}_2 \mathbf{b}_2). \end{aligned}$$

$\mathbf{b}_1$ é o conjunto de coeficientes na regressão de $\mathbf{y}$ sobre $\mathbf{X}_1$, menos um vetor de correção. Suponha que $\mathbf{X}_1^\top \mathbf{X}_2 = \mathbf{0}$, então $\mathbf{b}_1 = (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{y}$.

Na solução de $\mathbf{b}_2$ temos

$$(\mathbf{X}_2^\top \mathbf{X}_1) \mathbf{b}_1 + (\mathbf{X}_2^\top \mathbf{X}_2) \mathbf{b}_2 = \mathbf{X}_2^\top \mathbf{y}$$

$$(\mathbf{X}_2^\top \mathbf{X}_1)(\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top (\mathbf{y} - \mathbf{X}_2 \mathbf{b}_2) + (\mathbf{X}_2^\top \mathbf{X}_2) \mathbf{b}_2 = \mathbf{X}_2^\top \mathbf{y}$$

$$\mathbf{X}_2^\top \mathbf{P}_1 \mathbf{y} - \mathbf{X}_2^\top \mathbf{P}_1 \mathbf{X}_2 \mathbf{b}_2 + (\mathbf{X}_2^\top \mathbf{X}_2) \mathbf{b}_2 = \mathbf{X}_2^\top \mathbf{y}$$

$$\mathbf{X}_2^\top (\mathbf{I} - \mathbf{P}_1) \mathbf{X}_2 \mathbf{b}_2 = \mathbf{X}_2^\top (\mathbf{I} - \mathbf{P}_1) \mathbf{y}$$

$$\mathbf{X}_2^\top \mathbf{M}_1 \mathbf{X}_2 \mathbf{b}_2 = \mathbf{X}_2^\top \mathbf{M}_1 \mathbf{y}$$

$$\mathbf{b}_2 = (\mathbf{X}_2^\top \mathbf{M}_1 \mathbf{X}_2)^{-1} \mathbf{X}_2^\top \mathbf{M}_1 \mathbf{y} = (\mathbf{X}_2^{*\top} \mathbf{X}_2^*)^{-1} \mathbf{X}_2^{*\top} \mathbf{y}^*,$$

sendo $\mathbf{X}_2^* = \mathbf{M}_1 \mathbf{X}_2$ e $\mathbf{y}^* = \mathbf{M}_1 \mathbf{y}$.

$\mathbf{M}_1$ é a matriz construtora de resíduos definida para a regressão sobre as colunas de $\mathbf{X}_1$.

Logo, $\mathbf{M}_1 \mathbf{X}_2$ é uma matriz de resíduos.

Cada coluna de $\mathbf{M}_1 \mathbf{X}_2$ é um vetor de resíduos da regressão correspondente da coluna $\mathbf{X}_2$ sobre as variáveis em $\mathbf{X}_1$.

Teorema (Frisch-Waugh, 1933; e Lovell, 1963)

Na regressão linear de mínimos quadrados do vetor $\mathbf{y}$ sobre dois conjuntos de variáveis $\mathbf{X}_1$ e $\mathbf{X}_2$, o subvetor $\mathbf{b}_2$ é o conjunto de coeficientes obtidos quando os resíduos da regressão de $\mathbf{y}$ sobre $\mathbf{X}_1$ sozinho são regredidos sobre o conjunto de resíduos obtidos quando cada coluna de $\mathbf{X}_2$ é regredida sobre $\mathbf{X}_1$.

Considere uma regressão de $\mathbf{y}$ sobre um conjunto de variáveis $\mathbf{X}$ e uma variável adicional $\mathbf{z}$, e seus respectivos coeficientes por $\mathbf{b}$ e c .

Corolário

O coeficiente sobre $\mathbf{z}$ em uma regressão múltipla de $\mathbf{y}$ sobre $\mathbf{W} = [\mathbf{X}, \mathbf{z}]$ é calculada como

$$c = (\mathbf{z}^\top \mathbf{M} \mathbf{z})^{-1} \mathbf{z}^\top \mathbf{M} \mathbf{y} = (\mathbf{z}^{\star\top} \mathbf{z}^{\star})^{-1} \mathbf{z}^{\star\top} \mathbf{y}^{\star},$$

sendo $\mathbf{z}^{\star} = \mathbf{M} \mathbf{z}$ e $\mathbf{y}^{\star} = \mathbf{M} \mathbf{y}$ os vetores de resíduos da regressão de mínimos quadrados de $\mathbf{z}$ e $\mathbf{y}$ sobre $\mathbf{X}$.

Lembrando que $\mathbf{M} = \mathbf{I} - \mathbf{P} = \mathbf{I} - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$.

Coeficiente de correlação parcial

Seja a regressão $y = \beta_1 + \beta_2 z + \beta_3 x + \varepsilon$. Além disso,

(a) y_i^* os resíduos da regressão de y sobre x , isto é, $y_i^* = y_i - a - bx_i$; e

(b) z_i^* os resíduos da regressão de z sobre x , isto é, $z_i^* = z_i - c - dx_i$,

para $i = 1, 2, \dots, n$.

Então, a correlação parcial r_{yz}^* é a correlação simples entre y^* e z^* , isto é,

$$r_{yz}^* = \frac{\mathbf{z}^{*\top} \mathbf{y}^*}{\sqrt{(\mathbf{z}^{*\top} \mathbf{z}^*)(\mathbf{y}^{*\top} \mathbf{y}^*)}}.$$

pois o modelo de regressão tem a constante β_1 . Daí, $\mathbf{1}^\top \mathbf{z}^* = \mathbf{1}^\top \mathbf{y}^* = 0$.

Notemos que $\mathbf{y}^* = \mathbf{M}\mathbf{y}$ e $\mathbf{z}^* = \mathbf{M}\mathbf{z}$ com $\mathbf{M} = \mathbf{I} - \mathbf{P} = \mathbf{I} - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$.

Soma de quadrados

Considere agora os resíduos $\mathbf{u}$ da regressão: $\mathbf{y} = \mathbf{X}\mathbf{d} + \mathbf{z}\mathbf{c} + \mathbf{u}$.

A menos que $\mathbf{X}^\top \mathbf{z} = \mathbf{0}$, $\mathbf{d}$ não será igual a $\mathbf{b} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$.

A menos que $\mathbf{c} = 0$, $\mathbf{u}$ não será igual a $\mathbf{e} = \mathbf{y} - \mathbf{X}\mathbf{b}$.

Notemos que

$$\begin{aligned}c &= (\mathbf{z}^{\star\top} \mathbf{z}^{\star})^{-1} \mathbf{z}^{\star\top} \mathbf{y}^{\star}, \quad \text{e} \\ \mathbf{d} &= (\mathbf{X}^{\top} \mathbf{X})^{-1} \mathbf{X}^{\top} (\mathbf{y} - \mathbf{z}c) = \mathbf{b} - (\mathbf{X}^{\top} \mathbf{X})^{-1} \mathbf{X}^{\top} \mathbf{z}c.\end{aligned}$$

Segue-se que

$$\mathbf{u} = \mathbf{y} - \mathbf{X}\mathbf{b} + \mathbf{X}(\mathbf{X}^{\top} \mathbf{X})^{-1} \mathbf{X}^{\top} \mathbf{z}c - \mathbf{z}c = \mathbf{e} - \mathbf{M}\mathbf{z}c = \mathbf{e} - \mathbf{z}^{\star}c.$$

Portanto,

$$\mathbf{u}^{\top} \mathbf{u} = \mathbf{e}^{\top} \mathbf{e} + c^2 (\mathbf{z}^{\star\top} \mathbf{z}^{\star}) - 2c \mathbf{z}^{\star\top} \mathbf{e}.$$

Mas

$$\begin{aligned}\mathbf{e} &= \mathbf{M}\mathbf{y} = \mathbf{y}^{\star} \\ \mathbf{z}^{\star\top} \mathbf{e} &= \mathbf{z}^{\star\top} \mathbf{M}\mathbf{y} = \mathbf{z}^{\top} \mathbf{M}^{\top} \mathbf{M}\mathbf{y} = \mathbf{z}^{\top} \mathbf{M}^{\top} \mathbf{y} = \mathbf{z}^{\star\top} \mathbf{y}^{\star} \\ \mathbf{z}^{\star\top} \mathbf{e} &= \mathbf{z}^{\star\top} \mathbf{y}^{\star} = c(\mathbf{z}^{\star\top} \mathbf{z}^{\star}),\end{aligned}$$

pois pela equação normal $(\mathbf{z}^{\star\top} \mathbf{z}^{\star})c = \mathbf{z}^{\star\top} \mathbf{y}^{\star}$.

Assim, $\mathbf{u}^{\top} \mathbf{u} = \mathbf{e}^{\top} \mathbf{e} - c^2 (\mathbf{z}^{\star\top} \mathbf{z}^{\star})$.

Teorema (mudança na soma de quadrados quando uma variável é adicionada ao modelo)

Se $\mathbf{e}^\top \mathbf{e}$ é a soma de quadrados dos resíduos quando $\mathbf{y}$ é regredido sobre $\mathbf{X}$ e $\mathbf{u}^\top \mathbf{u}$ é a soma de quadrados dos resíduos quando $\mathbf{y}$ é regredido sobre $\mathbf{X}$ e $\mathbf{z}$, então

$$\mathbf{u}^\top \mathbf{u} = \mathbf{e}^\top \mathbf{e} - c^2(\mathbf{z}^{*\top} \mathbf{z}^*) \leq \mathbf{e}^\top \mathbf{e},$$

sendo c o coeficiente de $\mathbf{z}$ na regressão completa e $\mathbf{z}^* = [\mathbf{I} - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top] \mathbf{z}$ o vetor de resíduos quando $\mathbf{z}$ é regredido sobre $\mathbf{X}$.

Temos que $\mathbf{e}^\top \mathbf{e} = \mathbf{y}^{*\top} \mathbf{y}^*$ e pela equação normal

$$c(\mathbf{z}^{*\top} \mathbf{z}^*) = \mathbf{z}^{*\top} \mathbf{y}^* \Rightarrow c^2(\mathbf{z}^{*\top} \mathbf{z}^*) = (\mathbf{z}^{*\top} \mathbf{y}^*)^2 / (\mathbf{z}^{*\top} \mathbf{z}^*).$$

Portanto,

$$\frac{\mathbf{u}^\top \mathbf{u}}{\mathbf{y}^{*\top} \mathbf{y}^*} = \frac{\mathbf{y}^{*\top} \mathbf{y}^* - (\mathbf{z}^{*\top} \mathbf{y}^*)^2 / \mathbf{z}^{*\top} \mathbf{z}^*}{\mathbf{y}^{*\top} \mathbf{y}^*} = 1 - r_{yz}^2.$$

Exemplo (continuação)

$$y_i = \beta_1 + \beta_2 T_i + \beta_3 G_i + \beta_3 R_i + \beta_4 P_i + \varepsilon_i$$

Seguimos a análise dos dados da equação de investimento.

Calculamos a correlação (simples) entre a variável investimento e as variáveis de tendência temporal (T), produto interno bruto (G), taxa de juros (R) e taxa de inflação (P).

Calculamos também a correlação parcial entre investimento e as quatro regressoras dado as demais regressoras.

Tabela 2: Correlação entre investimento e as demais variáveis.

	Correlação Simples	Correlação Parcial
Tendência	0,7496	-0,9360
PIB	0,8632	0,9680
Taxa de juros	0,5871	-0,5167
Taxa de inflação	0,4777	-0,0221

Ver arquivo exemplo_04.r

Qualidade de ajuste e análise de variância

Podemos perguntar se a variação em $\mathbf{x}$ é um bom preditor da variação em y .

A variação total é dada por:

$$SQTot = \sum_{i=1}^n (y_i - \bar{y})^2.$$

Já sabemos que $\mathbf{y} = \mathbf{X}\mathbf{b} + \mathbf{e} = \hat{\mathbf{y}} + \mathbf{e}$. Logo, para uma única observação, temos que

$$\begin{aligned} y_i &= \hat{y}_i + e_i = \mathbf{x}_i^\top \mathbf{b} + e_i \\ y_i - \bar{y} &= \hat{y}_i - \bar{y} + e_i = (\mathbf{x}_i - \bar{\mathbf{x}})^\top \mathbf{b} + e_i. \end{aligned}$$

Exemplo

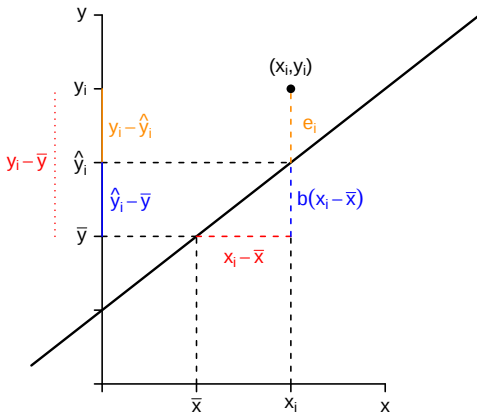


Figura 3: Decomposição de y_i .

Ver arquivo exemplo_05.r

Algumas propriedades, notações e somas de quadrados

$$\mathbf{1}\bar{x} = \mathbf{1}\frac{1}{n}\mathbf{1}^\top \mathbf{x} = \begin{bmatrix} \bar{x} \\ \bar{x} \\ \vdots \\ \bar{x} \end{bmatrix} = \frac{1}{n}\mathbf{1}\mathbf{1}^\top \mathbf{x}$$

$$\begin{bmatrix} x_1 - \bar{x} \\ x_2 - \bar{x} \\ \vdots \\ x_n - \bar{x} \end{bmatrix} = [\mathbf{x} - \mathbf{1}\bar{x}] = [\mathbf{x} - \frac{1}{n}\mathbf{1}\mathbf{1}^\top \mathbf{x}] = [I\mathbf{x} - \frac{1}{n}\mathbf{1}\mathbf{1}^\top \mathbf{x}] = [I - \frac{1}{n}\mathbf{1}\mathbf{1}^\top]\mathbf{x} = \tilde{\mathbf{M}}\mathbf{x}.$$

Definimos $\tilde{\mathbf{M}} = I - \frac{1}{n}\mathbf{1}\mathbf{1}^\top$. Esta também é uma matriz idempotente.

Temos que $\tilde{\mathbf{M}}\mathbf{1} = \mathbf{0} \Leftrightarrow \mathbf{1}^\top \tilde{\mathbf{M}} = \mathbf{0}^\top$. Logo,

$$\sum_{i=1}^n (x_i - \bar{x}) = \mathbf{1}^\top [\tilde{\mathbf{M}}\mathbf{x}] = \mathbf{0}^\top \mathbf{x} = 0.$$

Além disso,

$$\sum_{i=1}^n (x_i - \bar{x})^2 = (\mathbf{x} - \mathbf{1}\bar{x})^\top (\mathbf{x} - \mathbf{1}\bar{x}) = (\tilde{\mathbf{M}}\mathbf{x})^\top (\tilde{\mathbf{M}}\mathbf{x}) = \mathbf{x}^\top \tilde{\mathbf{M}}^\top \tilde{\mathbf{M}}\mathbf{x} = \mathbf{x}^\top \tilde{\mathbf{M}}\mathbf{x}$$

$$\begin{aligned}\tilde{\mathbf{M}}\mathbf{y} &= \tilde{\mathbf{M}}\mathbf{X}\mathbf{b} + \tilde{\mathbf{M}}\mathbf{e} \\ \tilde{\mathbf{M}}\mathbf{e} &= \mathbf{e} \quad \mathbf{e}^\top \tilde{\mathbf{M}}\mathbf{X} = \mathbf{e}^\top \mathbf{X} = \mathbf{0}^\top.\end{aligned}$$

Finalmente,

$$\underbrace{\mathbf{y}^\top \tilde{\mathbf{M}}\mathbf{y}}_{SQTot} = \underbrace{\mathbf{b}^\top \mathbf{X}^\top \tilde{\mathbf{M}}\mathbf{X}\mathbf{b}}_{SQReg} + \underbrace{\mathbf{e}^\top \mathbf{e}}_{SQRes}.$$

- SQTot é a soma de quadrados total;
- SQReg é a soma de quadrados da regressão; e
- SQRes é a soma de quadrados dos resíduos.

Coeficiente de determinação

$$R^2 = \frac{SQReg}{SQTot} = \frac{\mathbf{b}^\top \mathbf{X}^\top \tilde{\mathbf{M}}\mathbf{X}\mathbf{b}}{\mathbf{y}^\top \tilde{\mathbf{M}}\mathbf{y}} = 1 - \frac{\mathbf{e}^\top \mathbf{e}}{\mathbf{y}^\top \tilde{\mathbf{M}}\mathbf{y}}.$$

Tabela 3: Análise de variância.

Soma de Quadrados	Fonte	Graus de Liberdade	Média Quadrática
Regressão	$\mathbf{b}^\top \mathbf{X}^\top \mathbf{y} - n\bar{y}^2$	$p - 1$	
Resíduo	$\mathbf{e}^\top \mathbf{e}$	$n - p$	s^2
Total	$\mathbf{y}^\top \mathbf{y} - n\bar{y}^2$	$n - 1$	$S_{yy}/(n - 1) = s_y^2$

$$\begin{aligned}
 \mathbf{b}^\top \mathbf{X}^\top \tilde{\mathbf{M}} \mathbf{X} \mathbf{b} &= \mathbf{b}^\top \mathbf{X}^\top \tilde{\mathbf{M}} [\mathbf{y} - \mathbf{e}] = \mathbf{b}^\top \mathbf{X}^\top \tilde{\mathbf{M}} \mathbf{y} - \mathbf{b}^\top \mathbf{X}^\top \tilde{\mathbf{M}} \mathbf{e} \\
 &= \mathbf{b}^\top \mathbf{X}^\top \left[\mathbf{I} - \frac{1}{n} \mathbf{1} \mathbf{1}^\top \right] \mathbf{y} \quad (\text{porque } \mathbf{X}^\top \tilde{\mathbf{M}} \mathbf{e} = \mathbf{0}) \\
 &= \mathbf{b}^\top \mathbf{X}^\top \mathbf{y} - \frac{1}{n} \mathbf{b}^\top \mathbf{X}^\top \mathbf{1} \mathbf{1}^\top \mathbf{y} = \mathbf{b}^\top \mathbf{X}^\top \mathbf{y} - \frac{1}{n} \hat{\mathbf{y}}^\top \mathbf{1} \mathbf{1}^\top \mathbf{y} \\
 &= \mathbf{b}^\top \mathbf{X}^\top \mathbf{y} - \frac{1}{n} (\mathbf{y} - \mathbf{e}) \mathbf{1} \mathbf{1}^\top \mathbf{y} \\
 &= \mathbf{b}^\top \mathbf{X}^\top \mathbf{y} - \frac{1}{n} \left[\sum_{i=1}^n y_i - \sum_{i=1}^n e_i \right] \sum_{i=1}^n y_i \\
 &= \mathbf{b}^\top \mathbf{X}^\top \mathbf{y} - n\bar{y}^2.
 \end{aligned}$$

Exemplo

$$y_i = \beta_1 + \beta_2 T_i + \beta_3 G_i + \beta_4 R_i + \beta_5 P_i + \varepsilon_i$$

Tabela 4: Análise de variância - equação de investimento.

Soma de Quadrados	Fonte	Graus de Liberdade	Média Quadrática	R^2
Regressão	0,0159025	4	0,00397563	0,9724
Resíduo	0,0004508	10	0,00004508	
Total	0,0163533	14	0,00116810	

Ver arquivo exemplo_06.r

Teorema (Mudança no R^2 quando uma variável é adicionada ao modelo)

Seja R_{XZ}^2 o coeficiente de determinação na regressão de $\mathbf{y}$ sobre $\mathbf{X}$ e uma variável adicional $\mathbf{z}$ (**u os resíduos**), seja R_X^2 o mesmo para a regressão de $\mathbf{y}$ sobre $\mathbf{X}$ sozinha (**e os resíduos**), e seja r_{yz}^{*2} o coeficiente de correlação parcial entre y e z . Então,

$$R_{XZ}^2 = R_X^2 + (1 - R_X^2)r_{yz}^{*2} \geq R_X^2.$$

Prova: Vimos que $\mathbf{z}^* = \mathbf{Mz}$, $\mathbf{y}^* = \mathbf{My} = \mathbf{e} \Rightarrow \mathbf{e}^\top \mathbf{e} = \mathbf{y}^{*\top} \mathbf{y}^*$, e

$$\mathbf{u}^\top \mathbf{u} = \mathbf{e}^\top \mathbf{e} - c^2(\mathbf{z}^{*\top} \mathbf{z}^*) = \mathbf{y}^{*\top} \mathbf{y}^* - \frac{(\mathbf{z}^{*\top} \mathbf{y}^*)^2}{\mathbf{z}^{*\top} \mathbf{z}^*} = \mathbf{e}^\top \mathbf{e}(1 - r_{yz}^{*2}).$$

Logo,

$$R_{XZ}^2 = 1 - \frac{\mathbf{u}^\top \mathbf{u}}{\mathbf{y}^\top \tilde{\mathbf{M}}\mathbf{y}} = 1 - \frac{\mathbf{e}^\top \mathbf{e}}{\mathbf{y}^\top \tilde{\mathbf{M}}\mathbf{y}} + \frac{\mathbf{e}^\top \mathbf{e}}{\mathbf{y}^\top \tilde{\mathbf{M}}\mathbf{y}} r_{yz}^{*2} = R_X^2 + (1 - R_X^2)r_{yz}^{*2} \geq R_X^2.$$

Exemplo - ajuste de uma função consumo

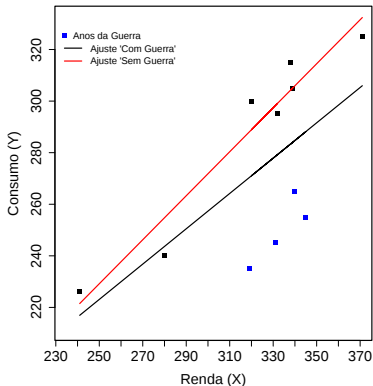


Figura 4: Dados de consumo e renda de 1940-1950 dos EUA. Fonte: *Economic Report of the President*, U.S. Government Printing Office, Washington, D.C., 1983.

Ver arquivo exemplo_07.r

Exemplo (continuação) - ajuste de uma função consumo

Tabela 5: Análise de variância: consumo e renda.

Soma de Quadrados	Fonte	Graus de Liberdade	Média Quadrática	R^2
Regressão	5768,2068	1	5768,2068	0,4571
Resíduo	6849,9751	9	761,1083	
Total	12618,1818	10	1261,8182	

Tabela 6: Análise de variância: consumo, renda e dicotômica para anos de guerra.

Soma de Quadrados	Fonte	Graus de Liberdade	Média Quadrática	R^2
Regressão x	5768,2068	1	5768,2068	0,9464
Regressão w	6173,5189	1	6173,5189	
Regressão	11941,7257	2	5970,8629	
Resíduo	676,4562	8	84,5570	
Total	12618,1818	10	1261,8182	

O R^2 ajustado para os graus de liberdade

R^2 nunca decresce quando uma nova variável regressora é adicionado a equação de regressão. Portanto, é tentador adicionar diversas variáveis regressoras no modelo para que R^2 se aproxime de 1.

O R^2 **ajustado** para os graus de liberdade é dado por

$$\bar{R}^2 = 1 - \frac{\mathbf{e}^\top \mathbf{e} / (n - p)}{\mathbf{y}^\top \tilde{\mathbf{M}} \mathbf{y} / (n - 1)}.$$

Para efeitos computacionais, temos

$$\bar{R}^2 = 1 - \frac{n - 1}{n - p} \times (1 - R^2).$$

- O $\bar{R}^2$ **pode decrescer** quando uma nova variável é adicionada ao conjunto de variáveis independentes.
- Na verdade, $\bar{R}^2$ **pode até ser negativo**.

- Uma alternativa para comparar modelos com uma maior penalização pela perda dos graus de liberdade é dada por

$$\bar{R}^2 = 1 - \frac{n+p}{n-p} \times (1 - R^2).$$

O modelo que tem $\bar{R}^2$ máximo é o mesmo que tem o valor mínimo do **critério de predição (CP)**:

$$\text{CP} = \frac{\mathbf{e}^\top \mathbf{e}}{n-p} \left(1 + \frac{p}{n}\right) = s^2 \left(1 + \frac{p}{n}\right).$$

- Veremos outros critérios de comparação de modelos baseados no máximo da função de verossimilhança com penalização pelo número de parâmetros (AIC, BIC, HQIC, etc.).

R^2 e o termo constante do modelo

A prova que $0 \leq R^2 \leq 1$ requer que $\mathbf{X}$ contenha uma coluna de 1's.

Se isto não ocorrer, então $\tilde{\mathbf{M}}\mathbf{e} \neq \mathbf{e}$, $\mathbf{e}^\top \tilde{\mathbf{M}}\mathbf{X} \neq \mathbf{0}$ e o termo $2\mathbf{e}^\top \tilde{\mathbf{M}}\mathbf{X}\mathbf{b}$ em

$$\mathbf{y}^\top \tilde{\mathbf{M}}\mathbf{y} = (\tilde{\mathbf{M}}\mathbf{X}\mathbf{b} + \tilde{\mathbf{M}}\mathbf{e})^\top (\tilde{\mathbf{M}}\mathbf{X}\mathbf{b} + \tilde{\mathbf{M}}\mathbf{e})$$

não se cancela.

Daí,

$$R^2 = 1 - \frac{\mathbf{e}^\top \mathbf{e}}{\mathbf{y}^\top \tilde{\mathbf{M}}\mathbf{y}}$$

pode resultar em qualquer valor.

Portanto, devemos ter cuidado ao analisar o R^2 de um modelo de regressão que não contenha o termo constante.

Este é um dos motivos que, em geral, deixa-se o termo constante no modelo de regressão linear mesmo que testes de hipóteses indiquem o contrário.

Regressão transformada linearmente

Na regressão de $\mathbf{y}$ sobre $\mathbf{X}$, suponha que as colunas de $\mathbf{X}$ são transformadas linearmente.

Exemplo - precificação de arte

A teoria I da determinação de preços de leilão em pinturas de Monet diz que o preço é determinado pelas dimensões (largura W e altura H) da pintura,

$$\begin{aligned}\ln \text{preço} &= \beta_1 + \beta_2 \ln W + \beta_3 \ln H + \varepsilon \\ &= \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + \varepsilon.\end{aligned}$$

A teoria II diz que os compradores estão interessados na área da pintura e a razão do aspecto,

$$\begin{aligned}\ln \text{preço} &= \gamma_1 + \gamma_2 \ln(WH) + \gamma_3 \ln(W/H) + \xi \\ &= \gamma_1 z_1 + \gamma_2 z_2 + \gamma_3 z_3 + \xi.\end{aligned}$$

É evidente que $z_1 = x_1$, $z_2 = x_2 + x_3$ e $z_3 = x_2 - x_3$.

Em termos matriciais, temos $\mathbf{Z} = \mathbf{XP}$ sendo

$$\mathbf{P} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 1 \\ 0 & 1 & -1 \end{bmatrix}.$$

Teorema

Na regressão linear de $\mathbf{y}$ sobre $\mathbf{Z} = \mathbf{XP}$ sendo $\mathbf{P}$ uma matriz não singular que transforma as colunas de $\mathbf{X}$, os coeficientes são iguais a $\mathbf{P}^{-1}\mathbf{b}$ sendo $\mathbf{b}$ o vetor de coeficientes da regressão linear de $\mathbf{y}$ sobre $\mathbf{X}$. Além disto, os R^2 das duas regressões são idênticos.

Prova: Os coeficientes são

$$\begin{aligned} \mathbf{d} &= (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{y} = [(\mathbf{XP})^\top (\mathbf{XP})]^{-1} (\mathbf{XP})^\top \mathbf{y} = (\mathbf{P}^\top \mathbf{X}^\top \mathbf{XP})^{-1} \mathbf{P}^\top \mathbf{X}^\top \mathbf{y} \\ &= \mathbf{P}^{-1} (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{P}^{\top -1} \mathbf{P}^\top \mathbf{X}^\top \mathbf{y} = \mathbf{P}^{-1} (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = \mathbf{P}^{-1} \mathbf{b}. \end{aligned}$$

Temos também que

$$\mathbf{u} = \mathbf{y} - \mathbf{Zd} = \mathbf{y} - \mathbf{ZP}^{-1}\mathbf{b} = \mathbf{y} - \mathbf{XPP}^{-1}\mathbf{b} = \mathbf{y} - \mathbf{Xb} = \mathbf{e}.$$

Logo, $\mathbf{u}^\top \mathbf{u} = \mathbf{e}^\top \mathbf{e} \Rightarrow R_{y|z}^2 = R_{y|x}^2$.

Hipóteses do modelo de regressão linear

- HP.1:** Linearidade: $y_i = x_{i1}\beta_1 + x_{i2}\beta_2 + \cdots + x_{ip}\beta_p + \varepsilon_i$;
- HP.2:** Posto completo: a matriz $\mathbf{X}$ ($n \times p$) é de posto completo;
- HP.3:** Exogeneidade das variáveis independentes: $E(\varepsilon_i | \mathbf{x}_i) = 0$;
- HP.4:** Homocedasticidade e autocorrelação nula: cada distúrbio, ε_i , tem a mesma variância, σ^2 , e é não correlacionado com qualquer outro distúrbio ε_j ;
- HP.5:** Geração dos dados: os dados em $(x_{i1}, x_{i2}, \dots, x_{ip})$ podem ser qualquer mistura de constantes e variáveis aleatórias; e
- HP.6:** Distribuição normal: os distúrbios são normalmente distribuídos.

O estimador de mínimos quadrados

Por hipótese,

$$E(\varepsilon|\mathbf{x}) = 0 \Rightarrow \text{Cov}(\mathbf{x}, \varepsilon) = \text{Cov}(E(\varepsilon|\mathbf{x}), \mathbf{x}) = \mathbf{0} \quad \text{e} \quad E(\varepsilon) = E_{\mathbf{x}}(E(\varepsilon|\mathbf{x})) = 0.$$

Mas,

$$\begin{aligned}\text{Cov}(\mathbf{x}, \varepsilon) &= E_{\mathbf{x}, \varepsilon}(\mathbf{x}\varepsilon) = E_{\mathbf{x}, y}(\mathbf{x}(y - \mathbf{x}^T\beta)) = \mathbf{0} \Rightarrow \\ E_{\mathbf{x}, y}(\mathbf{x}y) &= E_{\mathbf{x}}(\mathbf{x}\mathbf{x}^T)\beta \quad \text{(população)}.\end{aligned}$$

Tomemos a equação normal de mínimos quadrados $\mathbf{X}^T\mathbf{y} = \mathbf{X}^T\mathbf{X}\mathbf{b}$ que dividida por n , podemos reescrever como

$$\left(\frac{1}{n}\sum_{i=1}^n \mathbf{x}_i y_i\right) = \left(\frac{1}{n}\sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T\right) \mathbf{b} \quad \text{(amostra)}.$$

Assim, o estimador de mínimos está parecido com sua respectiva quantidade populacional. Podemos aplicar a lei dos grandes números sob certas condições.

Preditor linear ótimo

Considere o problema de encontrar um preditor linear ótimo para y .

Critério: o erro quadrático médio (EQM) do preditor linear y .

Ignoraremos a **HP.6** e desconsideraremos **HP.1** que $E(y|\mathbf{x})$ é linear.

Queremos minimizar

$$\text{EQM} = E_{\mathbf{x},y}([y - \mathbf{x}^\top \boldsymbol{\gamma}]^2).$$

Isto pode ser reescrito como

$$\text{EQM} = E_{\mathbf{x},y}([y - E(y|\mathbf{x})]^2) + E_{\mathbf{x},y}([E(y|\mathbf{x}) - \mathbf{x}^\top \boldsymbol{\gamma}]^2).$$

A condição necessária é

$$\begin{aligned} \frac{\partial E_{\mathbf{x},y}([E(y|\mathbf{x}) - \mathbf{x}^\top \boldsymbol{\gamma}]^2)}{\partial \boldsymbol{\gamma}} &= E_{\mathbf{x},y} \left(\frac{\partial [E(y|\mathbf{x}) - \mathbf{x}^\top \boldsymbol{\gamma}]^2}{\partial \boldsymbol{\gamma}} \right) \\ &= -2E_{\mathbf{x},y}(\mathbf{x}[E(y|\mathbf{x}) - \mathbf{x}^\top \boldsymbol{\gamma}]) = \mathbf{0} \Rightarrow \\ E_{\mathbf{x},y}(\mathbf{x}E(y|\mathbf{x})) &= E_{\mathbf{x},y}(\mathbf{x}\mathbf{x}^\top)\boldsymbol{\gamma}. \end{aligned}$$

Mas

$$\begin{aligned} E_{\mathbf{x},y}(\mathbf{x}E(y|\mathbf{x})) &= \text{Cov}(\mathbf{x}, E(y|\mathbf{x})) + E_{\mathbf{x}}(\mathbf{x})E_{\mathbf{x}}(E(y|\mathbf{x})) \\ &= \text{Cov}(\mathbf{x}, y) + E_{\mathbf{x}}(\mathbf{x})E_y(y) = E_{\mathbf{x},y}(\mathbf{x}y). \end{aligned}$$

Portanto, a condição necessária para minimizar o preditor linear, EQM, é

$$E_{\mathbf{x},y}(\mathbf{x}y) = E_{\mathbf{x},y}(\mathbf{x}\mathbf{x}^\top)\gamma.$$

Teorema

Se o processo de geração dos dados $(y_i, \mathbf{x}_i)$, $i = 1, 2, \dots, n$, é tal que a lei dos grandes números se aplica ao estimador

$$\left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i y_i \right) = \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^\top \right) \mathbf{b}$$

das matrizes

$$E_{\mathbf{x},y}(\mathbf{x}y) = E_{\mathbf{x}}(\mathbf{x}\mathbf{x}^\top)\beta,$$

então o mínimo do preditor linear do erro quadrático esperado de y_i é estimado pela regressão de mínimos quadrados.

Estimação não tendenciosa

O estimador de mínimos quadrados é não tendencioso para cada amostra.

Temos

$$\mathbf{b} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top (\mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}) = \boldsymbol{\beta} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon} \quad \text{e}$$

$$E(\mathbf{b}|\mathbf{X}) = \boldsymbol{\beta} + E((\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon}|\mathbf{X}) = \boldsymbol{\beta} \quad (\text{pela HP.3}).$$

Portanto,

$$E(\mathbf{b}) = E_X(E(\mathbf{b}|\mathbf{X})) = E_X(\boldsymbol{\beta}) = \boldsymbol{\beta}.$$

Viés: omissão de variáveis relevantes

Suponha que o modelo correto seja

$$\mathbf{y} = \mathbf{X}_1\beta_1 + \mathbf{X}_2\beta_2 + \varepsilon$$

sendo p_1 e p_2 os números de colunas em $\mathbf{X}_1$ e $\mathbf{X}_2$, respectivamente.

A regressão de $\mathbf{y}$ sobre $\mathbf{X}_1$ (sem incluir $\mathbf{X}_2$) resulta no estimador

$$\mathbf{b}_1 = (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{y} = \beta_1 + (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{X}_2 \beta_2 + (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \varepsilon.$$

Logo,

$$E(\mathbf{b}_1 | \mathbf{X}) = \beta_1 + (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{X}_2 \beta_2 = \beta_1 + \mathbf{P}_{1.2} \beta_2,$$

sendo $\mathbf{P}_{1.2} = (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{X}_2$.

Portanto, $\mathbf{b}_1$ é um estimador tendencioso de β_1 , exceto para os casos que $\mathbf{X}_1^\top \mathbf{X}_2 = \mathbf{0}$ ou $\beta_2 = \mathbf{0}$.

A estimação do modelo pode ser vista ao impor **incorretamente** a restrição **$\beta_2 = \mathbf{0}$** . Isto introduz um viés na estimação de $\mathbf{b}_1$.

Inclusão de variáveis irrelevantes

Suponha que o modelo de regressão correto seja

$$\mathbf{y} = \mathbf{X}_1\boldsymbol{\beta}_1 + \varepsilon,$$

mas que estimamos como a modelo de regressão (variáveis extras)

$$\mathbf{y} = \mathbf{X}_1\boldsymbol{\beta}_1 + \mathbf{X}_2\boldsymbol{\beta}_2 + \varepsilon.$$

Aqui, a inclusão de $\mathbf{X}_2$ no modelo de regressão pode ser vista como uma falha em impor que $\boldsymbol{\beta}_2 = \mathbf{0}$. Portanto, não a nada a se provar aqui.

Temos que

$$E(\mathbf{b}|\mathbf{X}) = \begin{bmatrix} \boldsymbol{\beta}_1 \\ \boldsymbol{\beta}_2 \end{bmatrix} = \begin{bmatrix} \boldsymbol{\beta}_1 \\ \mathbf{0} \end{bmatrix}.$$

É possível provar que ao aumentar o número de variáveis desnecessárias no modelo faz com que a precisão das estimativas diminua (a menos que $\mathbf{X}_1^\top \mathbf{X}_2 = \mathbf{0}$).

Exemplo

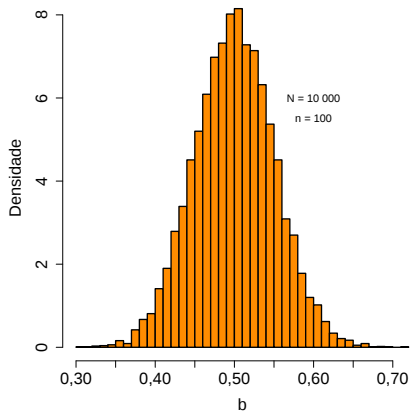


Figura 5: Distribuição amostral do estimador b de β na regressão $y = \alpha + \beta x + \varepsilon$.

Ver arquivo exemplo_08.r

Variância amostral do estimador de mínimos quadrados

$$\mathbf{b} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top (\mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}) = \boldsymbol{\beta} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon}.$$

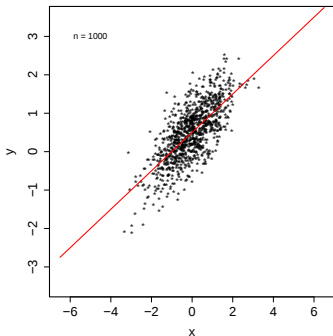
Logo,

$$\begin{aligned}\text{Var}(\mathbf{b}|\mathbf{X}) &= \text{E}([\mathbf{b} - \boldsymbol{\beta}][\mathbf{b} - \boldsymbol{\beta}]^\top | \mathbf{X}) \\&= \text{E}((\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^\top \mathbf{X} (\mathbf{X}^\top \mathbf{X})^{-1} | \mathbf{X}) \\&= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \text{E}(\boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^\top | \mathbf{X}) \mathbf{X} (\mathbf{X}^\top \mathbf{X})^{-1} \\&= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \sigma^2 \mathbf{I} \mathbf{X} (\mathbf{X}^\top \mathbf{X})^{-1} \quad (\text{pela HP.4}) \\&= \sigma^2 (\mathbf{X}^\top \mathbf{X})^{-1}.\end{aligned}$$

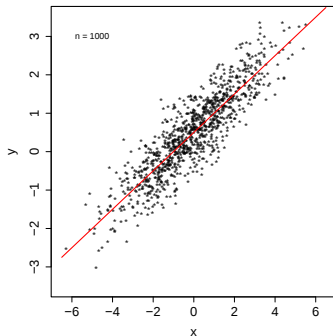
Suponha uma regressão linear simples da forma $y = \beta_1 + \beta_2 x + \varepsilon$. Assim, temos

$$\text{Var}(b_2|\mathbf{X}) = \frac{\sigma^2}{\sum_{i=1}^n (x_i - \bar{x})^2}.$$

Exemplo



(a) Menor variância em x .



(b) Maior variância em x .

Figura 6: Variância do estimador b de β na regressão $y = \alpha + \beta x + \varepsilon$.

Ver arquivo exemplo_09.r

Tabela 7: Estimativas para os dados com **menor** dispersão em x .

Parâmetro	Estimativa	Erro padrão	Estat. t	Valor p
β_0	0,52368	0,01579	33,16	$< 10^{-15}$
β_1	0,51920	0,01566	33,15	$< 10^{-15}$

$$\hat{\sigma}_1 = 0,4992 \text{ e } R_1^2 = 0,5240.$$

Tabela 8: Estimativas para os dados com **maior** dispersão em x .

Parâmetro	Estimativa	Erro padrão	Estat. t	Valor p
β_0	0,52294	0,01580	33,09	$< 10^{-15}$
β_1	0,49683	0,00826	60,15	$< 10^{-15}$

$$\hat{\sigma}_2 = 0,4996 \text{ e } R_2^2 = 0,7838.$$

Agora, seja $\mathbf{b}_\star = \mathbf{C}\mathbf{y}$ um outro estimador não tendencioso e linear de β , em que $\mathbf{C}$ é uma matriz $p \times n$. Então,

$$\beta = E(\mathbf{b}_\star|\mathbf{X}) = E(\mathbf{C}\mathbf{y}|\mathbf{X}) = E(\mathbf{C}\mathbf{X}\beta + \mathbf{C}\varepsilon|\mathbf{X}) \Rightarrow \mathbf{C}\mathbf{X} = \mathbf{I}.$$

(existem vários candidatos para $\mathbf{C}$ que satisfazem a igualdade) e

$$\text{Var}(\mathbf{b}_\star|\mathbf{X}) = \sigma^2 \mathbf{C}\mathbf{C}^\top.$$

Agora, seja $\mathbf{D} = \mathbf{C} - (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$ tal que $\mathbf{D}\mathbf{y} = \mathbf{b}_\star - \mathbf{b}$. Então, $\mathbf{b}_\star = \mathbf{D}\mathbf{y} + \mathbf{b} = (\mathbf{D} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top)\mathbf{y}$ e

$$\text{Var}(\mathbf{b}_\star|\mathbf{X}) = \sigma^2 \left[(\mathbf{D} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top)(\mathbf{D} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top)^\top \right].$$

Sabemos que $\mathbf{C}\mathbf{X} = \mathbf{I} = \mathbf{D}\mathbf{X} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{X} \Rightarrow \mathbf{D}\mathbf{X} = \mathbf{0}$. Portanto,

$$\text{Var}(\mathbf{b}_\star|\mathbf{X}) = \sigma^2 (\mathbf{X}^\top \mathbf{X})^{-1} + \sigma^2 \mathbf{D}\mathbf{D}^\top = \text{Var}(\mathbf{b}|\mathbf{X}) + \sigma^2 \mathbf{D}\mathbf{D}^\top.$$

e $\mathbf{D}\mathbf{D}^\top$ é uma forma quadrática ($\mathbf{a}^\top \mathbf{D}\mathbf{D}^\top \mathbf{a} \geq 0$).

Logo, $\text{Var}(\mathbf{b}|\mathbf{X}) \leq \text{Var}(\mathbf{b}_\star|\mathbf{X})$ para qualquer $\mathbf{b}_\star \neq \mathbf{b}$ e $\mathbf{X}$ fixo.

Teorema de Gauss-Markov

No modelo de regressão linear clássico com matriz de variáveis regressoras $\mathbf{X}$, o estimador de mínimos quadrados $\mathbf{b}$ é o estimador não tendencioso e linear de variância mínima de β . Para qualquer vetor de constantes $\mathbf{v}$, o estimador não tendencioso e linear de variância mínima de $\mathbf{v}^\top \beta$ no modelo de regressão clássico é $\mathbf{v}^\top \mathbf{b}$, sendo $\mathbf{b}$ o estimador de mínimos quadrados.

O teorema também implica que o estimador de mínimos quadrados para cada coeficiente também tem variância mínima. Basta tomar um vetor $\mathbf{v} = (0, \dots, 0, 1, 0, \dots, 0)$ com o valor 1 no coeficiente a ser estimador.

A variância incondicional é dada por

$$\begin{aligned}\text{Var}(\mathbf{b}) &= \text{E}_X(\text{Var}(\mathbf{b}|\mathbf{X})) + \text{Var}_X(\text{E}(\mathbf{b}|\mathbf{X})) \\ &= \sigma^2 \text{E}_X((\mathbf{X}^\top \mathbf{X})^{-1}),\end{aligned}$$

pois $\text{E}(\mathbf{b}|\mathbf{X}) = \beta \Rightarrow \text{Var}_X(\text{E}(\mathbf{b}|\mathbf{X})) = \mathbf{0}$.

Estimação da variância

Sabemos que $\text{Var}(\mathbf{b}|\mathbf{X}) = \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}$. Precisamos estimar σ^2 .

Temos que $\text{Var}(\varepsilon_i) = \text{E}(\varepsilon_i^2) = \sigma^2$ e e_i é uma estimativa de ε_i . Então,

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n e_i^2$$

parece um estimador natural de σ^2 .

Os resíduos de mínimos quadrados são

$$\mathbf{e} = \mathbf{M}\mathbf{y} = \mathbf{M}[\mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}] = \mathbf{M}\boldsymbol{\varepsilon}, \quad \text{pois } \mathbf{M}\mathbf{X} = \mathbf{0}.$$

Este estimador de σ^2 é baseado em $\mathbf{e}^\top \mathbf{e} = \boldsymbol{\varepsilon}^\top \mathbf{M}\boldsymbol{\varepsilon}$.

Note que os resíduos são correlacionados:

$$\text{Var}(\mathbf{e}|\mathbf{X}) = \text{E}(\mathbf{e}\mathbf{e}^\top|\mathbf{X}) = \text{E}(\mathbf{M}\boldsymbol{\varepsilon}\boldsymbol{\varepsilon}^\top \mathbf{M}|\mathbf{X}) = \mathbf{M}\text{E}(\boldsymbol{\varepsilon}\boldsymbol{\varepsilon}^\top|\mathbf{X})\mathbf{M} = \mathbf{M}\sigma^2 \mathbf{I}\mathbf{M} = \sigma^2 \mathbf{M}.$$

Note que $\mathbf{M}$ é uma matriz positiva definida.

Agora, estamos interessados em

$$E(\mathbf{e}^\top \mathbf{e} | \mathbf{X}) = E(\boldsymbol{\varepsilon}^\top \mathbf{M} \boldsymbol{\varepsilon} | \mathbf{X}).$$

A quantidade escalar $\boldsymbol{\varepsilon}^\top \mathbf{M} \boldsymbol{\varepsilon}$ é uma matriz 1×1 , logo é igual ao seu traço.

Segue-se que

$$\begin{aligned} E(\text{tr}(\boldsymbol{\varepsilon}^\top \mathbf{M} \boldsymbol{\varepsilon}) | \mathbf{X}) &= E(\text{tr}(\mathbf{M} \boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^\top) | \mathbf{X}) = \text{tr}(E(\mathbf{M} \boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^\top | \mathbf{X})) = \text{tr}(\mathbf{M} E(\boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^\top | \mathbf{X})) \\ &= \text{tr}(\mathbf{M} \sigma^2 \mathbf{I}) = \sigma^2 \text{tr}(\mathbf{M}). \end{aligned}$$

O traço de $\mathbf{M}$ é

$$\text{tr}(\mathbf{I}_n - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top) = \text{tr}(\mathbf{I}_n) - \text{tr}((\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{X}) = \text{tr}(\mathbf{I}_n) - \text{tr}(\mathbf{I}_p) = n - p.$$

Portanto, $E(\mathbf{e}^\top \mathbf{e} | \mathbf{X}) = (n - p)\sigma^2$.

Logo, o estimador natural de σ^2 é tendencioso. Um estimador não tendencioso de σ^2 é

$$s^2 = \frac{\mathbf{e}^\top \mathbf{e}}{n - p}.$$

Este estimador é não tendencioso incondicionalmente também,

$$E(s^2) = E_X(E(s^2|\mathbf{X})) = E_X(\sigma^2) = \sigma^2.$$

Podemos estimar $\text{Var}(\mathbf{b}|\mathbf{X})$ por $\widehat{\text{Var}}(\mathbf{b}|\mathbf{X}) = s^2(\mathbf{X}^\top \mathbf{X})^{-1}$.

A raiz quadrada do k -ésimo elemento da diagonal de $s^2(\mathbf{X}^\top \mathbf{X})^{-1}$ é o erro padrão do estimador b_k .

Normalidade (HP.6)

Se tivermos a hipótese $\varepsilon \sim \mathcal{N}_n(\mathbf{0}, \sigma^2 \mathbf{I})$, então

$$\begin{aligned}(\mathbf{b}|\mathbf{X}) &\sim \mathcal{N}_p(\boldsymbol{\beta}, \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}), \quad \text{e} \\(b_k|\mathbf{X}) &\sim \mathcal{N}(\beta_k, \sigma^2(\mathbf{X}^\top \mathbf{X})_{kk}^{-1}).\end{aligned}$$

Lembre-se que $\mathbf{b} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top (\mathbf{X}\boldsymbol{\beta} + \varepsilon) = \boldsymbol{\beta} + \mathbf{A}\varepsilon$ com $\mathbf{A} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$, ou seja, sob **HP.6** $\mathbf{b}$ é uma combinação linear de normais. Portanto, $\mathbf{b}$ (condicional a $\mathbf{X}$) tem distribuição normal.

Propriedades dos estimadores de mínimos quadrados para amostras grandes

Utilizaremos **HP.1** até **HP.5** (sem normalidade).

Consistência de $\mathbf{b}$

Suponha que **HP.5a** $(\mathbf{x}_i, \varepsilon_i)$, $i = 1, \dots, n$, seja uma sequência de observações independentes e que

$$\text{plim} \frac{\mathbf{X}^\top \mathbf{X}}{n} = \mathbf{Q}, \quad \text{uma matriz positiva definida.}$$

Temos que

$$\mathbf{b} = \beta + \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{X}^\top \varepsilon}{n} \right).$$

Se $\mathbf{Q}^{-1}$ existir, então

$$\text{plim} \mathbf{b} = \beta + \mathbf{Q}^{-1} \text{plim} \left(\frac{\mathbf{X}^\top \varepsilon}{n} \right).$$

pois a inversa é uma função contínua da matriz original.

Segue-se

$$\frac{1}{n} \mathbf{X}^\top \boldsymbol{\varepsilon} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \varepsilon_i = \frac{1}{n} \sum_{i=1}^n \mathbf{w}_i = \bar{\mathbf{w}}.$$

Portanto,

$$\text{plim} \mathbf{b} = \boldsymbol{\beta} + \mathbf{Q}^{-1} \text{plim} \bar{\mathbf{w}}.$$

Temos que

$$E(\mathbf{w}_i) = E_{\mathbf{x}}(E(\mathbf{w}_i | \mathbf{X})) = E_{\mathbf{x}}(\mathbf{x}_i E(\varepsilon_i | \mathbf{X})) = \mathbf{0} \Rightarrow E(\bar{\mathbf{w}}) = \mathbf{0}.$$

Além disso,

$$\text{Var}(\bar{\mathbf{w}} | \mathbf{X}) = E(\bar{\mathbf{w}} \bar{\mathbf{w}}^\top | \mathbf{X}) = \frac{1}{n} \mathbf{X}^\top E(\boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^\top | \mathbf{X}) \mathbf{X} \frac{1}{n} = \left(\frac{\sigma^2}{n} \right) \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right) \Rightarrow$$

$$\text{Var}(\bar{\mathbf{w}}) = E_{\mathbf{x}}(\text{Var}(\bar{\mathbf{w}} | \mathbf{X})) + \text{Var}_{\mathbf{x}}(E(\bar{\mathbf{w}} | \mathbf{X})) = \left(\frac{\sigma^2}{n} \right) E_{\mathbf{x}} \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right).$$

Se $\lim_{n \rightarrow \infty} E_{\mathbf{x}} \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right) = \mathbf{Q}$, então $\lim_{n \rightarrow \infty} \text{Var}(\bar{\mathbf{w}}) = \mathbf{0}$ e

$$\text{plim} \frac{\mathbf{X}^\top \boldsymbol{\varepsilon}}{n} = \mathbf{0} \Rightarrow$$

$$\text{plim} \mathbf{b} = \boldsymbol{\beta} + \mathbf{Q}^{-1} \mathbf{0} = \boldsymbol{\beta}.$$

Assim, $\mathbf{b}$ é um estimador consistente de $\boldsymbol{\beta}$ no modelo de regressão linear.

Condições de Grenander

- HG.1:** Para cada coluna de $\mathbf{X}$, $\mathbf{x}_k$, se $d_{nk}^2 = \mathbf{x}_k^\top \mathbf{x}_k$, então $\lim_{n \rightarrow \infty} d_{nk}^2 = +\infty$. Assim, $\mathbf{x}_k$ não degenera para uma sequência de zeros. A soma de quadrados continuará a crescer com o tamanho da amostra.
- HG.2:** $\lim_{n \rightarrow \infty} x_{ik}^2 / d_{nk}^2 = 0$ para todo $i = 1, \dots, n$. Esta condição implica que nenhuma observação dominará $\mathbf{x}_k^\top \mathbf{x}_k$ e quando $n \rightarrow \infty$, observações individuais serão menos importantes.
- HG.3:** Seja $\mathbf{R}_n$ a correlação amostral das colunas de $\mathbf{X}$, excluindo a variável constante. Então, $\lim_{n \rightarrow \infty} \mathbf{R}_n = \mathbf{C}$, uma matriz positiva definida. Isto implica que a condição de posto completo sempre será satisfeita. Nós já assumimos que $\mathbf{X}$ tem posto completo para uma amostra finita, então esta hipótese assegura que a condição nunca será violada.

Estas condições garantem que a matriz de dados são “bem comportadas” em amostras grandes mesmo nos casos envolvendo séries temporais com tendência temporal, séries polinômiais e variáveis de tendência (as hipóteses anteriores eram muito restritivas).

Normalidade assintótica de $\mathbf{b}$

Utilizaremos a **HP.3** (exogeneidade) e adicionaremos a hipótese que as observações são independentes.

Sabemos que

$$\begin{aligned}\sqrt{n}(\mathbf{b} - \boldsymbol{\beta}) &= \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)^{-1} \left(\frac{1}{\sqrt{n}} \right) \mathbf{X}^\top \boldsymbol{\varepsilon} \Rightarrow \\ \left[\text{plim} \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)^{-1} \right] \left(\frac{1}{\sqrt{n}} \right) \mathbf{X}^\top \boldsymbol{\varepsilon} &= \mathbf{Q}^{-1} \left(\frac{1}{\sqrt{n}} \right) \mathbf{X}^\top \boldsymbol{\varepsilon}.\end{aligned}$$

Mas

$$\left(\frac{1}{\sqrt{n}} \right) \mathbf{X}^\top \boldsymbol{\varepsilon} = \sqrt{n}(\bar{\mathbf{w}} - \mathbf{E}(\bar{\mathbf{w}})) \quad \text{sendo} \quad \mathbf{E}(\bar{\mathbf{w}}) = \mathbf{0}.$$

Podemos utilizar a versão multivariada do teorema central do limite de Lindeberg-Feller para obter a distribuição limite de $\sqrt{n}\bar{\mathbf{w}}$.

Mas $\bar{\mathbf{w}}$ é a média de n vetores aleatórios independentes $\mathbf{w}_i = \mathbf{x}_i \varepsilon_i$ com média $\mathbf{0}$ e variância $\text{Var}(\mathbf{w}_i) = \text{Var}(\mathbf{x}_i \varepsilon_i) = \sigma^2 \text{E}(\mathbf{x}_i \mathbf{x}_i^\top) = \sigma^2 \mathbf{Q}_i$. Segue-se que a variância de $\sqrt{n}\bar{\mathbf{w}}$ é

$$\sigma^2 \overline{\mathbf{Q}}_n = \sigma^2 \left(\frac{1}{n} \right) [\mathbf{Q}_1 + \mathbf{Q}_2 + \cdots + \mathbf{Q}_n].$$

Se a soma não for dominada por nenhum termo e se os regressores forem bem comportados,

$$\lim_{n \rightarrow \infty} \sigma^2 \overline{\mathbf{Q}}_n = \sigma^2 \mathbf{Q}.$$

Portanto, se $[\mathbf{x}_i \varepsilon_i]$, $i = 1, \dots, n$, forem vetores independentes distribuídos com média $\mathbf{0}$ e variância $\sigma^2 \mathbf{Q}_i$ e se $\text{plim} \frac{\mathbf{X}^\top \mathbf{X}}{n} = \mathbf{Q}$, então

$$\frac{1}{\sqrt{n}} \mathbf{X}^\top \boldsymbol{\varepsilon} \xrightarrow{d} \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{Q}).$$

Segue-se que

$$\mathbf{Q}^{-1} \left(\frac{1}{\sqrt{n}} \right) \mathbf{X}^\top \boldsymbol{\varepsilon} \xrightarrow{d} \mathcal{N}(\mathbf{Q}^{-1} \mathbf{0}, \mathbf{Q}^{-1} \sigma^2 \mathbf{Q} \mathbf{Q}^{-1}) \Rightarrow \sqrt{n}(\mathbf{b} - \boldsymbol{\beta}) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{Q}^{-1}).$$

Teorema

Se $\{\varepsilon_i\}$ são independentemente distribuídos com média 0 e variância σ^2 e x_{ik} é tal que as condições de Grenander são satisfeitas, então

$$\mathbf{b} \stackrel{a}{\sim} \mathcal{N} \left(\boldsymbol{\beta}, \frac{\sigma^2}{n} \mathbf{Q}^{-1} \right).$$

Na prática, é necessário estimar $(1/n)\mathbf{Q}^{-1}$ com $(\mathbf{X}^\top \mathbf{X})^{-1}$ e σ^2 com $\mathbf{e}^\top \mathbf{e} / (n - p)$.

Consistência de s^2

Temos a variância assintótica de $\mathbf{b}$: $(\sigma^2/n)\mathbf{Q}^{-1}$.

Podemos restringir a atenção a s^2 como estimador consistente de σ^2 .

$$\begin{aligned}s^2 &= \frac{1}{n-p} \boldsymbol{\varepsilon}^\top \mathbf{M} \boldsymbol{\varepsilon} = \frac{1}{n-p} \left[\boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon} - \boldsymbol{\varepsilon}^\top \mathbf{X} (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon} \right] \\ &= \frac{n}{n-p} \left[\frac{\boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon}}{n} - \left(\frac{\boldsymbol{\varepsilon}^\top \mathbf{X}}{n} \right) \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{X}^\top \boldsymbol{\varepsilon}}{n} \right) \right].\end{aligned}$$

Mas

$$\begin{aligned}\text{plim} \left(\frac{\boldsymbol{\varepsilon}^\top \mathbf{X}}{n} \right) \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{X}^\top \boldsymbol{\varepsilon}}{n} \right) &= 0, \quad \text{pois} \\ \text{plim} \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)^{-1} &= \mathbf{Q}^{-1} \quad \text{e} \quad \text{plim} \left(\frac{\mathbf{X}^\top \boldsymbol{\varepsilon}}{n} \right) = 0.\end{aligned}$$

Agora, vamos analisar

$$\overline{\varepsilon^2} = \frac{\boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon}}{n} = \frac{1}{n} \sum_{i=1}^n \varepsilon_i^2.$$

Neste caso, $E(\varepsilon_i^2) = \sigma^2$. Pelo teorema de Markov, precisamos de $E|\varepsilon_i|^2|^{1+\delta} < \infty$, $\delta > 0$, tal que a hipótese mínima até agora é que ε_i tenha momentos finitos um pouco maior que 2.

Além disso, se assumirmos que cada ε_i tem a mesma distribuição, então pelo teorema de Khinchine, momentos finitos até 2 é suficiente.

Convergência em média quadrática exigiria que $E(\varepsilon_i^4) = \phi_\varepsilon < \infty$.

Então, os termos da soma são independentes, com média σ^2 e variância $\phi_\varepsilon - \sigma^4$. Logo, sob condições fracas,

$$\text{plim} \frac{\mathbf{e}^\top \mathbf{e}}{n-p} = \text{plim} \frac{\varepsilon^\top \mathbf{M} \varepsilon}{n-p} = \text{plim} s_n^2 = \sigma^2 \Rightarrow \text{plim} s_n^2 \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)^{-1} = \sigma^2 \mathbf{Q}^{-1}.$$

O estimador apropriado da matriz de covariância assintótica de $\mathbf{b}$ é

$$\mathbf{S}_b = s^2 (\mathbf{X}^\top \mathbf{X})^{-1}.$$

Distribuição assintótica de uma função de $\mathbf{b}$

Seja $\mathbf{f}(\mathbf{b})$ um conjunto de J funções contínuas, e continuamente diferenciáveis do estimador de mínimos quadrados, e seja

$$\mathbf{C}(\mathbf{b}) = \frac{\partial \mathbf{f}(\mathbf{b})}{\partial \mathbf{b}^\top},$$

sendo $\mathbf{C}$ uma matriz $J \times p$ cuja j -ésima linha é o vetor de derivadas da j -ésima função com respeito a $\mathbf{b}^\top$.

Pelo teorema de Slutsky,

$$\text{plim} \mathbf{f}(\mathbf{b}) = \mathbf{f}(\boldsymbol{\beta}) \quad \text{e} \quad \text{plim} \mathbf{C}(\mathbf{b}) = \frac{\partial \mathbf{f}(\boldsymbol{\beta})}{\partial \boldsymbol{\beta}^\top} = \boldsymbol{\Gamma}.$$

Teorema

Se $\mathbf{f}(\mathbf{b})$ é um conjunto de funções contínuas e continuamente diferenciável de $\mathbf{b}$ tal que $\boldsymbol{\Gamma} = \frac{\partial \mathbf{f}(\boldsymbol{\beta})}{\partial \boldsymbol{\beta}^\top}$, e se $\mathbf{b} \overset{a}{\sim} \mathcal{N}\left(\boldsymbol{\beta}, \frac{\sigma^2}{n} \mathbf{Q}^{-1}\right)$, então

$$\mathbf{f}(\mathbf{b}) \overset{a}{\sim} \mathcal{N}\left(\mathbf{f}(\boldsymbol{\beta}), \boldsymbol{\Gamma} \left(\frac{\sigma^2}{n} \mathbf{Q}^{-1}\right) \boldsymbol{\Gamma}^\top\right).$$

A prova do teorema utiliza a aproximação de Taylor:

$$\mathbf{f}(\mathbf{b}) = \mathbf{f}(\boldsymbol{\beta}) + \boldsymbol{\Gamma}(\mathbf{b} - \boldsymbol{\beta}) + \text{outros termos.}$$

Na prática, estimaremos a variância de $\mathbf{f}(\mathbf{b})$ por $\mathbf{S}_{f(\mathbf{b})} = s^2 \mathbf{C}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{C}^\top$.

Exemplo - demanda por gasolina

$$\begin{aligned} \ln(\text{G/Pop})_t &= \beta_1 + \gamma \ln(\text{G/Pop})_{t-1} + \beta_2 \ln \text{Preço}_t + \beta_3 \ln(\text{Renda/Pop})_t \\ &+ \beta_4 \ln(\text{PCN})_t + \beta_5 \ln(\text{PCU})_t + \varepsilon_t, \end{aligned}$$

sendo PCN o preço dos carros novos, PCU o preço dos carros usados e G o consumo de gasolina (gasto total (GasExp) dividido pelo preço (GasP)).

- A elasticidade a longo prazo do preço é $\epsilon_2 = \frac{\beta_2}{1 - \gamma}$.
- A elasticidade a longo prazo da renda é $\epsilon_3 = \frac{\beta_3}{1 - \gamma}$.

Utilizaremos mínimos quadrados para estimar os parâmetros do modelo.

Ver arquivo exemplo_10.r

Sejam $\mathbf{x}_t = (1, \ln \text{Preço}_t, \ln(\text{Renda/Pop})_t, \ln(\text{PCN})_t, \ln(\text{PCU})_t)^\top$, $y_t = \ln(\text{G/Pop})_t$ e $\beta = (\beta_1, \dots, \beta_5)^\top$.

Então, podemos reescrever o modelo como:

$$\begin{aligned} y_t &= \mathbf{x}_t^\top \beta + \gamma y_{t-1} + \varepsilon_t \\ y_t - \frac{1}{1-\gamma} \mathbf{x}_t^\top \beta &= \gamma \left(y_{t-1} - \frac{1}{1-\gamma} \mathbf{x}_{t-1}^\top \beta \right) + \varepsilon_t \\ y_t - \mu_t &= \gamma (y_{t-1} - \mu_{t-1}) + \varepsilon_t \\ z_t &= \gamma z_{t-1} + \varepsilon_t, \end{aligned}$$

sendo $z_t = y_t - \mu_t$, $z_{t-1} = y_{t-1} - \mu_{t-1}$ e $\mu_t = \frac{1}{1-\gamma} \mathbf{x}_t^\top \beta$.

Temos, portanto, um modelo autoregressivo de ordem 1 (AR(1)) estacionário (por hipótese pois estaremos removendo todas as tendências com os regressores em $\mathbf{x}_t$).

Da estacionariedade do modelo AR(1), temos que

$E(z_t) = \gamma E(z_{t-1}) + E(\varepsilon) \Rightarrow E(z_t) = 0 \Rightarrow E(y_t) = \mu_t$ (média de longo prazo).

Agora, podemos calcular as elasticidade de longo prazo do consumo em relação ao preço ou a renda.

A elasticidade a longo prazo do consumo em relação ao preço:

$$\epsilon_2 = \frac{\partial \ln(G/Pop)_t}{\partial \ln \text{Preço}_t} = \frac{\partial \mu_t}{\partial \ln \text{Preço}_t} = \frac{\beta_2}{1 - \gamma}.$$

A elasticidade a longo prazo do consumo em relação a renda:

$$\epsilon_3 = \frac{\partial \ln(G/Pop)_t}{\partial \ln(Renda/Pop)_t} = \frac{\partial \mu_t}{\partial \ln(Renda/Pop)_t} = \frac{\beta_3}{1 - \gamma}.$$

Seja $\theta = (\beta_1, \gamma, \beta_2, \dots, \beta_5)^\top$. Então,

$$\begin{aligned}\frac{\partial \epsilon_2(\theta)}{\partial \theta} &= \left(0, \frac{\beta_2}{(1 - \gamma)^2}, \frac{1}{1 - \gamma}, 0, 0, 0\right)^\top \\ \frac{\partial \epsilon_3(\theta)}{\partial \theta} &= \left(0, \frac{\beta_3}{(1 - \gamma)^2}, 0, \frac{1}{1 - \gamma}, 0, 0\right)^\top.\end{aligned}$$

Estimação por máxima verossimilhança

Suponha **HP.6** seja verdade. Então, o estimador dos coeficientes da regressão por mínimos quadrados coincide com o estimador dos coeficientes por máxima verossimilhança.

O logaritmo da função de verossimilhança baseada em uma amostra aleatória de tamanho n com erros normais é dada por

$$\ell(\beta, \sigma^2) = \ln L(\beta, \sigma^2) = -\frac{n}{2} \ln 2\pi - \frac{n}{2} \ln \sigma^2 - \frac{1}{2\sigma^2} (\mathbf{y} - \mathbf{X}\beta)^\top (\mathbf{y} - \mathbf{X}\beta).$$

O estimador de máxima verossimilhança de (β, σ^2) é dado por

$$\begin{aligned} \mathbf{b}_{EMV} &= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}, \quad \text{e} \\ \hat{\sigma}_{EMV}^2 &= \frac{1}{n} (\mathbf{y} - \mathbf{X}\mathbf{b})^\top (\mathbf{y} - \mathbf{X}\mathbf{b}) = \frac{1}{n} \mathbf{e}^\top \mathbf{e}. \end{aligned}$$

Por ser um estimador de máxima verossimilhança, $\mathbf{b}_{EMV}$ é assintoticamente eficiente entre os estimadores consistentes, e assintoticamente normal.

Veremos estimadores de máxima verossimilhança mais adiante.

Estimação por intervalo

Começaremos utilizando todas as hipóteses do modelo de regressão linear, inclusive **HP.6**.

Dado $\varepsilon \sim \mathcal{N}_n(\mathbf{0}, \sigma^2 \mathbf{I})$, temos que

$$\begin{aligned}(\mathbf{b}|\mathbf{X}) &\sim \mathcal{N}_p(\boldsymbol{\beta}, \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}), \quad \text{e} \\(b_k|\mathbf{X}) &\sim \mathcal{N}(\beta_k, \sigma^2(\mathbf{X}^\top \mathbf{X})_{kk}^{-1}).\end{aligned}$$

Definamos s_{kk} o k -ésimo elemento da diagonal de $(\mathbf{X}^\top \mathbf{X})^{-1}$.

Daí, temos

$$Z_k = \frac{b_k - \beta_k}{\sqrt{\sigma^2 s_{kk}}} \sim \mathcal{N}(0, 1).$$

Mas σ^2 é desconhecido.

Resultado 1

Os autovalores de uma matriz idempotente são compostos de 0 (zeros) ou 1 (uns).

Resultado 2

O posto de uma matriz idempotente é igual ao seu traço (um número inteiro).

Resultado 3

A única matriz idempotente não singular é a matriz identidade.

Resultado 4

Se $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ e $\mathbf{C}$ é uma matriz tal que $\mathbf{C}^\top \mathbf{C} = \mathbf{I}$, então $\mathbf{C}^\top \mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$.

Resultado 5

Se $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ e $\mathbf{A}$ é uma matriz idempotente, então $\mathbf{x}^\top \mathbf{A} \mathbf{x} \sim \chi_m^2$ sendo m o número de raízes unitárias de $\mathbf{A}$ (que é igual ao posto de $\mathbf{A}$).

Resultado 6

Se $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, $\mathbf{L}\mathbf{x}$ uma função linear e $\mathbf{x}^\top \mathbf{A} \mathbf{x}$ uma forma quadrática idempotente, então $\mathbf{L}\mathbf{x}$ e $\mathbf{x}^\top \mathbf{A} \mathbf{x}$ são independentes se $\mathbf{L}\mathbf{A} = \mathbf{0}$.

A quantidade

$$\frac{(n-p)S^2}{\sigma^2} = \frac{\mathbf{e}^\top \mathbf{e}}{\sigma^2} = \left(\frac{\varepsilon}{\sigma}\right)^\top \mathbf{M} \left(\frac{\varepsilon}{\sigma}\right)$$

é uma forma quadrática idempotente do vetor normal (ε/σ) com $\text{tr}(\mathbf{M}) = n - p$ e a forma linear dada por

$$\left(\frac{\mathbf{b} - \boldsymbol{\beta}}{\sigma}\right) = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \left(\frac{\varepsilon}{\sigma}\right).$$

Temos que $(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{M} = \mathbf{0}$ pois $\mathbf{M}\mathbf{X} = \mathbf{0}$. Daí, temos que a independência.

Teorema (independência de $\mathbf{b}$ e s^2)

Se ε é um vetor com distribuição normal, então o estimador $\mathbf{b}$ de mínimos quadrados dos coeficientes é estatisticamente independente do vetor de resíduos $\mathbf{e}$ e, portanto, todas as funções de $\mathbf{e}$, incluindo s^2 .

Destes resultados, temos que

$$\frac{b_k - \beta_k}{\sqrt{s^2 s_{kk}}} \sim t_{n-p}.$$

Consequentemente,

$$IC_{100(1-\alpha)\%}(\beta_k) = (b_k - t_{n-p, 1-\alpha/2} \sqrt{s^2 s_{kk}}, b_k + t_{n-p, 1-\alpha/2} \sqrt{s^2 s_{kk}}).$$

Exemplo (continuação) - demanda por Gasolina

$$\ln(G/Pop)_t = \beta_1 + \beta_2 \ln \text{Preço}_t + \beta_3 \ln(\text{Renda/Pop})_t + \beta_4 \ln(PCN)_t + \beta_5 \ln(PCU)_t + \varepsilon_t.$$

Ver arquivo exemplo_10.r

Tabela 9: Intervalos de confiança de 95%.

Variável	2,5%	97,5%
Constante	-22,7264	-19,6958
ln Preço	-0,1093	0,0668
ln Renda/Pop	0,9395	1,2522
ln PCN	-0,6896	-0,0576
ln PCU	-0,1878	0,2279

Os intervalos de confiança para ϵ_2 e ϵ_3 baseados em teoria assintótica:

$$IC_{95\%}(\epsilon_2) = (-0,7099; -0,1129) \quad \text{e} \quad IC_{95\%}(\epsilon_3) = (0,6523; 1,2888).$$

Combinação linear dos coeficientes

- Seja $\mathbf{w}$ um vetor $p \times 1$ de constantes.
- Estamos interessados na combinação linear $\alpha = \mathbf{w}^\top \boldsymbol{\beta}$.
- Estimamos α com $a = \mathbf{w}^\top \mathbf{b}$ sendo $\mathbf{b}$ o estimador de mínimos quadrados.
- Como $\mathbf{b}$ tem distribuição normal, temos

$$a \sim \mathcal{N}(\mathbf{w}^\top \boldsymbol{\beta}, \sigma^2 \mathbf{w}^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{w}).$$

- Ainda estimamos σ^2 com s^2 .
- Estimamos a variância de a com $s_a^2 = s^2 \mathbf{w}^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{w}$.
- Segue-se que $(a - \alpha)/s_a \sim t_{n-p}$ (pois a e s_a^2 são independentes).
- Daí, podemos construir intervalos de confiança para α .
- Podemos testar somas ou diferenças entre os coeficientes.

Predição

- Suponha que desejamos prever o valor de y_* associado ao vetor de regressores $\mathbf{x}_*$.
- Este valor seria $y_* = \mathbf{x}_*^\top \boldsymbol{\beta} + \varepsilon_*$.
- Segue do teorema de Gauss-Markov que $\hat{y}_* = \mathbf{x}_*^\top \mathbf{b}$ é o estimador não tendencioso linear de variância mínima de $E(y_*|\mathbf{x}_*)$.
- O erro de previsão é $e_* = y_* - \hat{y}_* = (\boldsymbol{\beta} - \mathbf{b})^\top \mathbf{x}_* + \varepsilon_*$.
- A variância de predição para este estimador é

$$\text{Var}(e_*|\mathbf{X}, \mathbf{x}_*) = \sigma^2 + \text{Var}((\boldsymbol{\beta} - \mathbf{b})^\top \mathbf{x}_*|\mathbf{X}, \mathbf{x}_*) = \sigma^2(1 + \mathbf{x}_*^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{x}_*).$$

- Podemos estimar esta variância ao substituir σ^2 pelo estimador s^2 .
- Assim, um intervalo de predição de y_* é dada por

$$\text{IC}_{100(1-\alpha)\%}(y_*) = (\hat{y}_* - t_{\alpha/2;n-p} \times \text{ep}(e_*), \hat{y}_* + t_{\alpha/2;n-p} \times \text{ep}(e_*)).$$

Exemplo

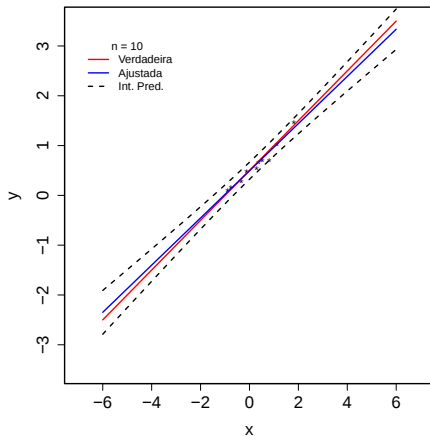


Figura 7: Intervalo de predição.

Ver arquivo exemplo_11.r

Multicolinearidade

Quando os regressores são altamente correlacionados, temos que

- pequenas mudanças nos dados podem produzir grandes variações na estimativas dos parâmetros;
- coeficientes podem ter erros padrões grandes e baixos níveis de significância mesmo que eles sejam conjuntamente significativos e o R^2 para a regressão seja elevado; e
- coeficientes podem ter um sinal equivocado ou magnitudes implausíveis.

Seja $\mathbf{X}$ a matriz de variáveis regressoras com uma constante e $p - 1$ variáveis medidas como desvios de suas médias, $\mathbf{x}_k = \mathbf{z}_k - \bar{z}_k \mathbf{1}$.

Seja $\mathbf{X}_{(k)}$ a matriz com todas as variáveis exceto $\mathbf{x}_k$. Então, o k -ésimo elemento de $(\mathbf{X}^\top \mathbf{X})^{-1}$ é dado por

$$\begin{aligned} (\mathbf{x}_k^\top \mathbf{M}_{(k)} \mathbf{x}_k)^{-1} &= \left[\mathbf{x}_k^\top \mathbf{x}_k - \mathbf{x}_k^\top \mathbf{X}_{(k)} (\mathbf{X}_{(k)}^\top \mathbf{X}_{(k)})^{-1} \mathbf{X}_{(k)}^\top \mathbf{x}_k \right]^{-1} \\ &= \left[\mathbf{x}_k^\top \mathbf{x}_k \left(1 - \frac{\mathbf{x}_k^\top \mathbf{X}_{(k)} (\mathbf{X}_{(k)}^\top \mathbf{X}_{(k)})^{-1} \mathbf{X}_{(k)}^\top \mathbf{x}_k}{\mathbf{x}_k^\top \mathbf{x}_k} \right) \right]^{-1} \\ &= \frac{1}{(1 - R_k^2) s_{kk}}, \end{aligned}$$

sendo R_k^2 o coeficiente de determinação da regressão de x_k sobre todas as outras variáveis do modelo e $s_{kk} = \sum_{i=1}^n (x_{ik} - \bar{x}_k)^2$.

A variância do estimador de mínimos quadrados do k -ésimo coeficiente é σ^2 vezes a razão acima,

$$\text{Var}(b_k | \mathbf{X}) = \frac{\sigma^2}{(1 - R_k^2) \sum_{i=1}^n (x_{ik} - \bar{x}_k)^2}.$$

Com as outras coisas iguais,

- quanto maior a correlação de x_k com as outras variáveis, maior será a variância de b_k , devido a multicolinearidade;
- quanto maior a variância de x_k , menor será a variância de b_k ; e
- quanto melhor o ajuste geral da regressão, menor será a variância devido a σ^2 .

Dados reais nunca serão ortogonais ($R_k^2 = 0$). Então, a multicolinearidade estará sempre presente.

Quando a multicolinearidade é um problema?

Podemos utilizar o **número de condição** de $\mathbf{X}^T \mathbf{X}$: a raiz quadrada da razão da maior raiz característica de $\mathbf{X}^T \mathbf{X}$ (sendo as colunas padronizadas para tamanho unitário) sobre a menor. Valores maiores que 20 sugerem uma indicação de problema.

Possíveis caminhos são

- aumentar o número de observações;
- retirar as variáveis mais problemáticas do modelo; e
- utilizar componentes principais (pode haver dificuldade de interpretação).

Exemplo - dados de Longley

Tabela 10: Correlação entre as variáveis.

	Emprego	Preço	PIB	Militar	Ano
Emprego	1,0000	0,9709	0,9836	0,4573	0,9713
Preço	0,9709	1,0000	0,9916	0,4647	0,9911
PIB	0,9836	0,9916	1,0000	0,4464	0,9953
Militar	0,4573	0,4647	0,4464	1,0000	0,4172
Ano	0,9713	0,9911	0,9953	0,4172	1,0000

Ver arquivo exemplo_12.r

Tabela 11: Ajuste do modelo com 15 observações.

	Estimativa	Erro padrão	Estat. t	Valor p
Constante	1459415,1	714182,9	2,04348	0,06825
Ano	-721,756	369,985	-1,95077	0,07965
Preço	-181,123	135,525	-1,33646	0,21101
PIB	0,09107	0,02026	4,49478	0,00115
Militar	-0,07494	0,26113	-0,28698	0,77999

Tabela 12: Ajuste do modelo com 16 observações.

	Estimativa	Erro padrão	Estat. t	Valor p
Constante	1169087,5	835902,4	1,39859	0,18949
Ano	-576,464	433,487	-1,32983	0,21049
Preço	-19,7681	138,893	-0,14233	0,88940
PIB	0,06439	0,01995	3,22746	0,00805
Militar	-0,01015	0,30857	-0,03288	0,97436

Testes de hipóteses

Estamos interessados em implicações testáveis. Por exemplo, o modelo apresenta algumas restrições testáveis.

Suponha o seguinte modelo de investimento (I_t):

$$\mathcal{M}_1 : \ln I_t = \beta_1 + \beta_2 \iota_t + \beta_3 \Delta p_t + \beta_4 \ln y_t + \beta_5 t + \varepsilon_t, \quad (1)$$

sendo

- ι_t : taxa de juros nominal;
- Δp_t : taxa de inflação; e
- $\ln y_t$: logaritmo da produção real.

Entretanto, uma teoria alternativa afirma que “investidores se preocupam com a taxa de juros real”.

$$\ln I_t = \beta_1 + \beta_2 (\iota_t - \Delta p_t) + \beta_3 \Delta p_t + \beta_4 \ln y_t + \beta_5 t + \varepsilon_t.$$

Não temos como testar a hipótese, pois ι_t e Δp_t estão em ambos os modelos.

Nota: $\iota_t - \Delta p_t$ é conhecida como a equação de Fisher.

Considere agora a alternativa em que “investidores se preocupam *somente* com a taxa de juros real

$$\mathcal{M}_2 : \ln I_t = \beta_1 + \beta_2(\iota_t - \Delta p_t) + \beta_4 \ln y_t + \beta_5 t + \varepsilon_t. \quad (2)$$

Agora, em relação ao modelo $\mathcal{M}_1$, temos a restrição $\beta_2 + \beta_3 = 0$.

Os modelos em (1) e (2) são encaixados.

Se os investidores não se preocuparem com a inflação (modelo $\mathcal{M}_3$), então $\beta_3 = 0$. Neste caso, o modelo $\mathcal{M}_3$ também seria encaixado com $\mathcal{M}_1$.

Então, suponha os seguintes modelos:

Modelo $\mathcal{M}_3$: Investidores se preocupam somente com a taxa de juros nominal $(\beta_1, \beta_2, 0, \beta_4, \beta_5)$.

Modelo $\mathcal{M}_4$: Investidores se preocupam somente com a inflação $(\beta_1, 0, \beta_3, \beta_4, \beta_5)$.

Estes dois modelos ($\mathcal{M}_3$ e $\mathcal{M}_4$) não são encaixados. Eles têm o mesmo número de parâmetros, mas diferentes covariáveis.

A seguir trataremos de modelos encaixados.

O modelo de regressão linear

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

Restrições lineares:

$$\begin{array}{rcl} r_{11}\beta_1 + r_{12}\beta_2 + \cdots + r_{1p}\beta_p & = & q_1 \\ r_{21}\beta_1 + r_{22}\beta_2 + \cdots + r_{2p}\beta_p & = & q_2 \\ & \vdots & \vdots \\ & \vdots & \vdots \\ r_{J1}\beta_1 + r_{J2}\beta_2 + \cdots + r_{Jp}\beta_p & = & q_J. \end{array}$$

Podemos organizar em notação matricial $\mathbf{R}\boldsymbol{\beta} = \mathbf{q}$.

A matriz $\mathbf{R}$ tem J linhas e J restrições, então $J < p$. Estamos eliminando a possibilidade de $J = p$ para evitar solução degenerada. As linhas de $\mathbf{R}$ devem ser linearmente independentes.

A restrição $\mathbf{R}\beta = \mathbf{q}$ impõe J restrições em p parâmetros livres. Então, existem $p - J$ parâmetros livres.

Abordagem para testes de hipóteses

Utilizaremos a hipótese de normalidade **HP.6**.

Queremos testar:

$$H_0 : \mathbf{R}\boldsymbol{\beta} - \mathbf{q} = \mathbf{0}$$

$$H_1 : \mathbf{R}\boldsymbol{\beta} - \mathbf{q} \neq \mathbf{0}.$$

Exemplos:

1. $\beta_j = 0$, $\mathbf{R} = [0 \ \dots \ 0 \ 1 \ 0 \ \dots \ 0]$, $\mathbf{q} = 0$;
2. $\beta_k = \beta_j$, $\mathbf{R} = [0 \ \dots \ 0 \ 1 \ 0 \ \dots \ -1 \ 0 \ \dots \ 0]$, $\mathbf{q} = 0$;
3. $\beta_2 + \beta_3 + \beta_4 = 1$, $\mathbf{R} = [0 \ 1 \ 1 \ 1 \ 0 \ \dots \ 0]$, $\mathbf{q} = 1$;
4. $\beta_2 + \beta_3 = 1$, $\beta_4 + \beta_6 = 0$, $\beta_5 + \beta_6 = 0$,

$$\mathbf{R} = \begin{bmatrix} 0 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 1 & 1 \end{bmatrix} \quad \text{e} \quad \mathbf{q} = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix};$$

5. $\beta_2 = \beta_3 = \dots = \beta_p = 0$, $\mathbf{R} = [\mathbf{0} : \mathbf{I}_{p-1}]$, $\mathbf{q} = \mathbf{0}$.

Dado $\mathbf{b}$, estamos interessados no vetor de discrepâncias $\mathbf{m} = \mathbf{R}\mathbf{b} - \mathbf{q}$.

Em termos estatísticos, $\mathbf{m}$ é significativamente diferente de $\mathbf{0}$?

Temos que

$$\mathbf{b} \sim \mathcal{N}(\boldsymbol{\beta}, \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}) \Rightarrow \mathbf{m} \sim \mathcal{N}(\mathbf{R}\boldsymbol{\beta} - \mathbf{q}, \sigma^2 \mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top).$$

Sob H_0 , $\mathbf{m} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top)$ e $\mathbf{R}\boldsymbol{\beta} = \mathbf{q}$ tal que
 $(1/\sigma)[\mathbf{R}\mathbf{b} - \mathbf{q}] = (1/\sigma)\mathbf{R}(\mathbf{b} - \boldsymbol{\beta}) = \mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top (\boldsymbol{\varepsilon}/\sigma) = \mathbf{D}(\boldsymbol{\varepsilon}/\sigma).$

Podemos testar a H_0 com o critério de Wald. Condicional a $\mathbf{X}$, temos

$$\begin{aligned} W &= \mathbf{m}^\top [\text{Var}(\mathbf{m}|\mathbf{X})]^{-1} \mathbf{m} \\ &= \frac{(\mathbf{R}\mathbf{b} - \mathbf{q})^\top [\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top]^{-1} (\mathbf{R}\mathbf{b} - \mathbf{q})}{\sigma^2} \\ &= \frac{1}{J} \left(\frac{\boldsymbol{\varepsilon}}{\sigma} \right)^\top \mathbf{D}^\top [\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top]^{-1} \mathbf{D} \left(\frac{\boldsymbol{\varepsilon}}{\sigma} \right), \end{aligned}$$

uma forma quadrática idempotente em $(\boldsymbol{\varepsilon}/\sigma)$ com

$$\text{tr} \left(\mathbf{D}^\top [\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top]^{-1} \mathbf{D} \right) = \text{tr} \left([\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top]^{-1} \mathbf{D} \mathbf{D}^\top \right) = \text{tr}(\mathbf{I}_J) = J.$$

Logo, $W \sim \chi_J^2$.

Mas a estatística W não pode ser utilizada porque desconhecemos σ^2 .

Podemos utilizar s^2 (regressão) para estimar σ^2 .

Resultado 7

Dado $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, se $\mathbf{x}^\top \mathbf{A} \mathbf{x}$ e $\mathbf{x}^\top \mathbf{B} \mathbf{x}$ são duas formas quadráticas idempotentes em $\mathbf{x}$, então $\mathbf{x}^\top \mathbf{A} \mathbf{x}$ e $\mathbf{x}^\top \mathbf{B} \mathbf{x}$ são independentes se $\mathbf{AB} = \mathbf{0}$.

Note que dado $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, temos que $\mathbf{x}^\top \mathbf{A} \mathbf{x}$ e $\mathbf{x}^\top \mathbf{B} \mathbf{x}$ têm distribuição qui-quadrado.

Sabemos que $\mathbf{b}$ e s^2 são independentes. Isto implica que W e s^2 são independentes. Logo,

$$\begin{aligned} F &= \frac{W/J}{(n-p)s^2/[\sigma^2(n-p)]} \\ &= \frac{(\mathbf{Rb} - \mathbf{q})^\top [\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top]^{-1} (\mathbf{Rb} - \mathbf{q})}{Js^2} \sim \mathcal{F}_{J, n-p}. \end{aligned}$$

Teste de hipóteses com uma restrição

Agora, considere uma abordagem alternativa para testar a hipótese

$$H_0 : r_1\beta_1 + r_2\beta_2 + \cdots + r_p\beta_p = \mathbf{r}^\top \boldsymbol{\beta} = q.$$

Sabemos que $\hat{q} = \mathbf{r}^\top \mathbf{b}$ é o estimador de $q = \mathbf{r}^\top \boldsymbol{\beta}$ e que

$$\hat{q} = \mathbf{r}^\top \mathbf{b} \sim \mathcal{N}(\mathbf{r}^\top \boldsymbol{\beta}, \sigma^2 \mathbf{r}^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{r}).$$

Logo,

$$t = \frac{\hat{q} - q}{\text{ep}(\hat{q})} \sim t_{n-p}, \quad \text{sendo} \quad \text{ep}(\hat{q}) = s \sqrt{\mathbf{r}^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{r}}.$$

Note que $t^2 \sim \mathcal{F}_{1, n-p}$. Portanto, os testes são equivalentes.

Exemplo 12 - investimento restrito

Analizaremos dados trimestrais dos EUA entre jan/1950 e dez/2000. O modelo é dado por

$$\ln I_t = \beta_1 + \beta_2 \iota_t + \beta_3 \Delta p_t + \beta_4 \ln y_t + \beta_5 t + \varepsilon_t.$$

Primeiro queremos testar $H_0 : \beta_2 + \beta_3 = 0$ (somente taxa de juros real).

- I_t é o investimento real (Realinvs);
- ι_t é a taxa de juros (Tbilrate);
- Δp_t é a inflação medida por variação no logaritmo do índice de preço ao consumidor (CPI_U); e
- $\ln y_t$ é o logaritmo do PIB (Realgdp).

Uma observação é "perdida" ao calcular Δp_t .

Ver arquivo exemplo_13.r

Tabela 13: Resumo do ajuste do modelo irrestrito.

	Estimativa	Erro padrão	Estat. t	Valor p
Constante	-9,134092	1,366459	-6,684	$< 10^{-10}$
Taxa de juros	-0,008598	0,003196	-2,691	0,00774
Inflação	0,003306	0,002337	1,415	0,15872
ln PIB	1,930156	0,183272	10,532	$< 10^{-15}$
Tempo	-0,005659	0,001488	-3,803	0,00019

Além disso, $\hat{\sigma} = 0,08618$ e $R^2 = 0,9798$.

Ainda, $F = 2.395$ com 4 e 198 graus de liberdade, e valor $p < 10^{-15}$.

Com os resultados do modelo irrestrito, temos

- $\hat{q} = b_2 + b_3 = -0,008598 + 0,003306 = -0,00529,$
- $ep(\hat{q}) = \sqrt{ep(b_2)^2 + ep(b_3)^2 + 2cov(b_2, b_3)} =$
 $\sqrt{0,00320^2 + 0,00234^2 + 2 * (-3,717e - 06)} = 0,002878, e$
- $t = \hat{q}/ep(\hat{q}) = -1,838.$
- Como $|t| < t_{0,975;198} = 1,972$, a hipótese nula não é rejeitada.
- Podemos estimar o modelo supondo $\beta_2 = -\beta_3$. (Próxima página.)

Ver arquivo exemplo_13.r

Tabela 14: Resumo do ajuste do modelo restrito.

	Estimativa	Erro padrão	Estat. t	Valor p
Constante	-7,90716	1,20063	-6,59	$< 10^{-10}$
Taxa de juros real	-0,00443	0,00227	-1,95	0,0526
ln PIB	1,76406	0,16056	10,99	$< 10^{-15}$
Tempo	-0,00440	0,00133	-3,31	0,0011

Além disso, $\hat{\sigma} = 0,0867$ e $R^2 = 0,979$.

Ainda, $F = 3.150$ com 3 e 198 graus de liberdade, e valor $p < 10^{-15}$.

Agora, queremos testar (conjuntamente) a seguinte hipótese nula:

- $\beta_2 + \beta_3 = 0$ (taxa de juros real),
- $\beta_4 = 1$ (a propensão marginal a investir é 1), e
- $\beta_5 = 0$ (não existe tendência temporal).

Ainda utilizando o ajuste no modelo irrestrito, temos

$$\mathbf{R} = \begin{bmatrix} 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}, \quad \mathbf{q} = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}, \quad \mathbf{m} = \mathbf{R}\mathbf{b} - \mathbf{q} = \begin{bmatrix} -0,005292 \\ 0,930156 \\ -0,005659 \end{bmatrix} \quad \text{e}$$

$$\widehat{\text{Var}(\mathbf{m}|\mathbf{X})} = \begin{bmatrix} 8,238e-06 & -0,0002586 & 1,956e-06 \\ -2,586e-04 & 0,0335887 & -2,718e-04 \\ 1,956e-06 & -0,0002718 & 2,214e-06 \end{bmatrix}.$$

$$\text{Assim, } F = \frac{\mathbf{m}^\top [\widehat{\text{Var}(\mathbf{m}|\mathbf{X})}]^{-1} \mathbf{m}}{J} = 109,8, \text{ sendo } J = 3.$$

Como $F > F_{3; 198; 0,95} = 2,65$, a hipótese nula é rejeitada (valor $p \approx 0$).

Predição para investimento

Queremos prever o primeiro trimestre de 2001 ao nível $\alpha = 0,05$.

- Taxa de juros: 4,48 (90-day T-Bill);
- Inflação: 528,0 (CPI_U);
- PIB: 9316,8 (realGDP); e
- Tempo: 204.

Ver arquivo exemplo_13.r

O intervalo de predição é dado por:

$$IC_{95\%}(\hat{y}_{*,204}) = (7,159; 7,503).$$

A taxa anual de investimento real no primeiro trimestre de 2001 foi 1.721.

O logaritmo é 7,4507. Assim, a predição contém este valor.

Estimadores de mínimos quadrados restritos

- Sabemos que $\mathbf{b}$ é escolhido para minimizar $\mathbf{e}^\top \mathbf{e}$.
- Sendo $R^2 = 1 - \frac{\mathbf{e}^\top \mathbf{e}}{\mathbf{y}^\top \tilde{\mathbf{M}} \mathbf{y}}$ e $\mathbf{y}^\top \tilde{\mathbf{M}} \mathbf{y}$ é uma constante que não envolve $\mathbf{b}$, então $\mathbf{b}$ é escolhido para maximizar R^2 .
- O estimador de mínimos quadrados restrito é obtido pela solução de

$$\min_{S(\mathbf{b})} = (\mathbf{y} - \mathbf{X}\mathbf{b})^\top (\mathbf{y} - \mathbf{X}\mathbf{b}) \quad \text{sujeito a} \quad \mathbf{R}\mathbf{b} = \mathbf{q}.$$

- Uma função Lagrangeana para este problema pode ser escrita como

$$L(\mathbf{b}_*, \lambda_*) = (\mathbf{y} - \mathbf{X}\mathbf{b}_*)^\top (\mathbf{y} - \mathbf{X}\mathbf{b}_*) + 2\lambda_*^\top (\mathbf{R}\mathbf{b}_* - \mathbf{q}).$$

- Como λ é irrestrito, utilizamos 2λ como restrição.
- A solução de $\mathbf{b}_*$ e λ_* satisfazem as condições necessárias:

$$\begin{aligned} \frac{\partial L(\mathbf{b}_*, \lambda_*)}{\partial \mathbf{b}_*} &= -2\mathbf{X}^\top (\mathbf{y} - \mathbf{X}\mathbf{b}_*) + 2\mathbf{R}^\top \lambda = \mathbf{0} \\ \frac{\partial L(\mathbf{b}_*, \lambda_*)}{\partial \lambda_*} &= 2(\mathbf{R}\mathbf{b}_* - \mathbf{q}) = \mathbf{0}. \end{aligned}$$

$$\begin{bmatrix} \mathbf{X}^\top \mathbf{X} & \mathbf{R}^\top \\ \mathbf{R} & \mathbf{0} \end{bmatrix} \begin{bmatrix} \mathbf{b}_* \\ \lambda_* \end{bmatrix} = \begin{bmatrix} \mathbf{X}^\top \mathbf{y} \\ \mathbf{q} \end{bmatrix} \quad \text{ou} \quad \mathbf{A} \mathbf{d}_* = \mathbf{v}.$$

Se a matriz $\mathbf{A}$ for não singular, então $\mathbf{d}_* = \mathbf{A}^{-1} \mathbf{v}$.

Se $\mathbf{X}^\top \mathbf{X}$ for não singular, então uma solução explícita para $\mathbf{b}_*$ e λ_* pode ser obtida usando a fórmula da inversa particionada:

$$\mathbf{b}_* = \mathbf{b} - (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top \left[\mathbf{R} (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top \right]^{-1} (\mathbf{R} \mathbf{b} - \mathbf{q}) = \mathbf{b} - \mathbf{C} \mathbf{m} \quad \text{e}$$

$$\lambda_* = \left[\mathbf{R} (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top \right]^{-1} (\mathbf{R} \mathbf{b} - \mathbf{q}).$$

Além disso,

$$\text{Var}(\mathbf{b}_* | \mathbf{X}) = \sigma^2 (\mathbf{X}^\top \mathbf{X})^{-1} - \sigma^2 (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top \left[\mathbf{R} (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top \right]^{-1} \mathbf{R} (\mathbf{X}^\top \mathbf{X})^{-1}.$$

(A prova está na próxima página.)

Logo, $\text{Var}(\mathbf{b}_* | \mathbf{X})$ é a variância de $\mathbf{b}$ reduzida pelas restrições.

Prova:

$$\begin{aligned}\text{Var}(\mathbf{b}_*|\mathbf{X}) &= \text{Var}(\mathbf{b}|\mathbf{X}) + \text{Var}(\mathbf{C}\mathbf{m}|\mathbf{X}) - 2\text{Cov}(\mathbf{b}, \mathbf{C}\mathbf{m}|\mathbf{X}) \\&= \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1} + \mathbf{C}\text{Var}(\mathbf{R}\mathbf{b} - \mathbf{q}|\mathbf{X})\mathbf{C}^\top - 2\text{Cov}(\mathbf{b}, \mathbf{R}\mathbf{b} - \mathbf{q}|\mathbf{X})\mathbf{C}^\top \\&= \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1} + \mathbf{C}\mathbf{R}\sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \mathbf{C}^\top - 2\text{Cov}(\mathbf{b}, \mathbf{b}|\mathbf{X})\mathbf{R}^\top \mathbf{C}^\top \\&= \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1} + \sigma^2 \mathbf{C}\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \mathbf{C}^\top - 2\sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \mathbf{C}^\top.\end{aligned}$$

Note que

$$\begin{aligned}\mathbf{C}\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top &= (\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \left[\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \right]^{-1} \mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \\&= (\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top.\end{aligned}$$

Daí, temos que

$$\begin{aligned}\text{Var}(\mathbf{b}_*|\mathbf{X}) &= \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1} + \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \mathbf{C}^\top - 2\sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \mathbf{C}^\top \\&= \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1} - \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \mathbf{C}^\top \\&= \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1} - \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \left[\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1}\mathbf{R}^\top \right]^{-1} \mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1}.\end{aligned}$$

A solução restrita de mínimos quadrados não pode ser melhor que a irrestrita:

$$\begin{aligned} \mathbf{e}_\star &= \mathbf{y} - \mathbf{X}\mathbf{b}_\star = \mathbf{y} - \mathbf{X}\mathbf{b} + \mathbf{X}\mathbf{b} - \mathbf{X}\mathbf{b}_\star = \mathbf{e} - \mathbf{X}(\mathbf{b}_\star - \mathbf{b}) \Rightarrow \\ \mathbf{e}_\star^\top \mathbf{e}_\star &= \mathbf{e}^\top \mathbf{e} + (\mathbf{b}_\star - \mathbf{b})^\top \mathbf{X}^\top \mathbf{X} (\mathbf{b}_\star - \mathbf{b}) \geq \mathbf{e}^\top \mathbf{e} \Rightarrow \\ R_\star^2 &\leq R^2. \end{aligned}$$

A perda de ajuste é dada por

$$\mathbf{e}_\star^\top \mathbf{e}_\star - \mathbf{e}^\top \mathbf{e} = (\mathbf{R}\mathbf{b} - \mathbf{q})^\top \left[\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top \right]^{-1} (\mathbf{R}\mathbf{b} - \mathbf{q}).$$

Esta é a expressão no numerador do teste F:

$$\begin{aligned} F &= \frac{(\mathbf{R}\mathbf{b} - \mathbf{q})^\top \left[\mathbf{R}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top \right]^{-1} (\mathbf{R}\mathbf{b} - \mathbf{q})}{J s^2} = \frac{(\mathbf{e}_\star^\top \mathbf{e}_\star - \mathbf{e}^\top \mathbf{e})/J}{\mathbf{e}^\top \mathbf{e}/(n-p)} \\ &= \frac{(R^2 - R_\star^2)/J}{(1 - R^2)/(n-p)} \sim \mathcal{F}_{J, n-p}. \end{aligned}$$

Se a restrição for $\beta_2 = \dots = \beta_p = 0$ (exceto intercepto), então temos $p - 1$ restrições e o **teste F** usual de significância da regressão.

Exemplo - função de produção

Função de produção de Cobb-Douglas para a indústria de metal.

Regressão de mínimos quadrados no logaritmo da produção (valor adicionado) sobre uma constante, logaritmo do trabalho e produção de capital.

Uma generalização do modelo de Cobb-Douglas é o modelo **translog**:

$$\begin{aligned}\ln y &= \beta_1 + \beta_2 \ln \text{TR} + \beta_3 \ln \text{CP} \\ &+ \beta_4 \left[\frac{1}{2} (\ln \text{TR})^2 \right] + \beta_5 \left[\frac{1}{2} (\ln \text{CP})^2 \right] + \beta_6 \ln \text{TR} \times \ln \text{CP} + \varepsilon,\end{aligned}$$

sendo TR a variável trabalho e CP a variável capital.

Veremos depois que este modelo relaxa a hipótese de elasticidade unitária de substituição. O modelo de Cobb-Douglas é obtido com $\beta_4 = \beta_5 = \beta_6 = 0$.

Sem **HP.6 (normalidade)** os testes podem ser aproximados para amostras grandes.

Ver arquivo exemplo_14.r

Tabela 15: Resumo do ajuste do modelo translog.

	Estimativa	Erro padrão	Estat. t	Valor p
Constante	0,94420	2,91075	0,324	0,7489
$\ln TR$	3,61364	1,54807	2,334	0,0296
$\ln CP$	-1,89311	1,01626	-1,863	0,0765
$(\ln TR)^2/2$	-0,96405	0,70738	-1,363	0,1874
$(\ln CP)^2/2$	0,08529	0,29261	0,291	0,7735
$(\ln TR)(\ln CP)$	0,31239	0,43893	0,712	0,4845

Além disso, $\hat{\sigma} = 0,1799$ e $R^2 = 0,9549$.

Ainda, $F = 88,85$ com 5 e 21 graus de liberdade, e valor $p < 10^{-12}$.

Note o valor da estatística do teste $\mathcal{F}$ e seu valor p . Os demais resultados e outros testes estão descritos no arquivo R.

Restrições não lineares

Trabalharemos com uma restrição e o teste do tipo $H_0 : c(\beta) = q$.

A estatística de teste é dada por

$$Z = \frac{c(\mathbf{b}) - q}{\text{ep}(c(\mathbf{b}))} \approx \mathcal{N}(0, 1), \quad \text{sendo} \quad c(\mathbf{b}) \approx c(\beta) + \left(\frac{\partial c(\beta)}{\partial \beta} \right)^\top (\mathbf{b} - \beta).$$

O resultado é baseado em consistência e não em não tendenciosidade visto que, em geral, $E(g(X)) \neq g(E(X))$. Precisamos estimar o erro padrão do estimador.

$$\text{Var}(c(\mathbf{b})) \approx \left(\frac{\partial c(\beta)}{\partial \beta} \right)^\top \text{Var}(\mathbf{b}|\mathbf{X}) \left(\frac{\partial c(\beta)}{\partial \beta} \right),$$

usando $\mathbf{h}(\mathbf{b})$ para estimar $\mathbf{h}(\beta) = \frac{\partial c(\beta)}{\partial \beta}$.

Assim,

$$\text{ep}(c(\mathbf{b})) = s \sqrt{\mathbf{h}(\mathbf{b})^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{h}(\mathbf{b})}.$$

Exemplo - propensão marginal de longo prazo a consumir

Dados trimestrais de consumo agregados dos EUA (c_t) e renda pessoal (y_t) de 1950 a 2000. O modelo é dado por:

$$\ln c_t = \alpha + \beta \ln y_t + \gamma \ln c_{t-1} + \varepsilon_t.$$

A propensão marginal a consumir (PMC) a curto prazo (elasticidade, pois as variáveis estão no logaritmo) é β .

A PMC de longo prazo é $\delta = \frac{\beta}{1 - \gamma}$.

Considere o teste $H_0 : \delta = 1$. Temos,

$$h_1(\alpha, \beta, \gamma) = \frac{\partial \delta}{\partial \beta} = \frac{1}{1 - \gamma} \quad \text{e} \quad h_2(\alpha, \beta, \gamma) = \frac{\partial \delta}{\partial \gamma} = \frac{\beta}{(1 - \gamma)^2}.$$

Ver arquivo exemplo_15.r

Tabela 16: Resumo do ajuste do modelo de consumo.

	Estimativa	Erro padrão	Estat. t	Valor p
Constante	0,003142	0,010553	0,298	0,76624
$\ln(y_t)$	0,074958	0,028727	2,609	0,00976
$\ln(c_{t-1})$	0,924625	0,028594	32,337	$< 10^{-15}$

Além disso, $\hat{\sigma} = 0,008742$ e $R^2 = 0,9997$.

Ainda, $F = 3,476 \times 10^5$ com 2 e 200 graus de liberdade, e valor $p < 10^{-15}$.

Para o teste $H_0 : \delta = 1$, temos que

$$\hat{\delta} = 0,9945, \quad \text{ep}(\hat{\delta}) = 0,01636, \quad z_{obs} = 0,3386 \quad \text{e} \quad z_{0,975} = 1,96.$$

Como $|z_{obs}| < 1,96$, então não rejeitamos a hipótese nula que $\delta = 1$.

Critérios de seleção de modelos

Definição

A função de verossimilhança de θ é a função que associa o valor $f_y(\mathbf{y}|\theta)$ para cada θ é definida da seguinte maneira

$$\begin{aligned} L(.; \mathbf{y}) : \Theta &\rightarrow \mathbb{R}^+ \\ \theta &\rightarrow L(\theta; \mathbf{y}) = f_y(\mathbf{y}|\theta). \end{aligned}$$

Definição

O critério de informação de Akaike (AIC) e o critério de informação bayesiana (BIC) para o modelo $\mathcal{M}_j$ são dados por

$$\begin{aligned} AIC(\mathcal{M}_j) &= -2 \ln L(\hat{\theta}_j; \mathbf{y}) + 2p_j \quad \text{e} \\ BIC(\mathcal{M}_j) &= -2 \ln L(\hat{\theta}_j; \mathbf{y}) + p_j \ln(n), \end{aligned}$$

sendo p_j o número de parâmetros e $\hat{\theta}_j$ a estimativa de máxima verossimilhança do modelo $\mathcal{M}_j$.

Quanto menor o AIC (ou BIC), melhor o modelo ajustado.

Seleção bayesiana de modelos

Temos M modelos, $m = 1, 2, \dots, M$, $f_m(\mathbf{y}|\mathbf{Z}, \boldsymbol{\theta}_m)$, sendo $f_m(\cdot)$ a densidade conjunta, $\mathbf{y}$ e $\mathbf{Z}$ os dados, e $\boldsymbol{\theta}_m$ o vetor de parâmetros do m -ésimo modelo.

Suponha que m^* seja o modelo correto, desconhecido pelo analista, mas que seja um dos M modelos.

Seja $g(m) = \pi_m$ a probabilidade a priori do modelo m . Assim, a probabilidade a posteriori do m -ésimo modelo é dada por

$$\omega_m = g(m|\mathbf{y}, \mathbf{Z}) = \frac{f_m(\mathbf{y}|\mathbf{Z}, m)\pi_m}{\sum_{r=1}^M f_r(\mathbf{y}|\mathbf{Z}, r)\pi_r},$$

sendo $f_m(\mathbf{y}|\mathbf{Z}, m)$ a função de verossimilhança marginal do modelo m ,

$$f_m(\mathbf{y}|\mathbf{Z}, m) = \int f_m(\mathbf{y}|\mathbf{Z}, \boldsymbol{\theta}_m, m) f_m(\boldsymbol{\theta}_m) d\boldsymbol{\theta}_m.$$

Seja $\hat{\theta}_m$ uma estimativa bayesiana, em geral a média a posteriori, dos parâmetros do modelo m . Então, uma possível aproximação é dada por

$$f_m(\mathbf{y}|\mathbf{Z}, m) \simeq \hat{f}_m(\mathbf{y}|\mathbf{Z}, m) = f_m(\mathbf{y}|\mathbf{Z}, \hat{\theta}_m, m) f_m(\hat{\theta}_m).$$

Daí, a densidade de previsão média do modelo bayesiano pode ser aproximada por

$$\bar{f} = \sum_{r=1}^M f_r(\mathbf{y}|\mathbf{Z}, \hat{\theta}_r, r) f_r(\hat{\theta}_r).$$

Assim,

$$\hat{\omega}_m = \frac{\hat{f}_m(\mathbf{y}|\mathbf{Z}, m)}{\bar{f}}, \quad \text{para } m = 1, 2, \dots, M.$$

Voltaremos a este ponto quando tratarmos dos modelos bayesianos.

Forma funcional e mudança de estrutura

Analisaremos maneiras de mudar a estrutura da variável dependente - no contexto de regressão - utilizando covariáveis.

Dados dicotômicos

- Os dados dicotômicos (binário, *dummy*) são muito utilizados em análise de regressão.
- Em geral, utiliza-se estas variáveis em conjunto com outras contínuas.

Exemplo - rendimentos

- $T = 428$ participantes do mercado de trabalho formal;
- Uma equação de (semilog) ganhos:

$$\ln(\text{rendimentos}) = \beta_1 + \beta_2 \text{idade} + \beta_3 \text{idade}^2 + \beta_4 \text{educação} + \beta_5 \text{filhos} + \varepsilon;$$

- rendimentos: salário por hora vezes horas trabalhadas;
- educação: anos de escola; e
- filhos: variável dicotômica indicando (1) a presença de filhos menores de 18 anos.

Ver arquivo exemplo_16.r

Tabela 17: Resumo do ajuste do modelo de rendimentos.

Parâmetro	Estimativa	Erro padrão	Estat. t	Valor p
β_1	3,24010	1,76743	1,833	0,06747
β_2	0,20056	0,08386	2,392	0,01721
β_3	-0,00231	0,00099	-2,345	0,01947
β_3	0,06748	0,02525	2,672	0,00782
β_5	-0,35122	0,14753	-2,380	0,01773

Além disso, $\hat{\sigma} = 1,19$ e $R^2 = 0,041$.

O valor $-0,3511952$ é considerado como um efeito extremamente grande (semilog).

Note que $\epsilon = \frac{\partial y/y}{\partial x/x} = \frac{\partial \ln y}{\partial \text{filhos}} \times \text{filhos} = \beta_k \text{filhos}$ (é um modelo semilog).

$1 - \exp(-0,3511952) = 0,2961537 \simeq 1/3$ (mulheres com filhos recebem quase 1/3 menos).

Possíveis aplicações de variáveis dicotômicas

- Estudo do efeito de tratamentos em algum tipo de resposta.
- Efeito de diploma universitário na renda ao longo da vida.
- Diferenças de sexo em relação ao salário na iniciativa privada.
- Modelo envolvendo uma variável dicotômica:

$$y_i = \mathbf{x}_i^\top \boldsymbol{\beta} + \delta d_i + \varepsilon_i.$$

- É prática comum incluir variáveis dicotômicas na regressão para modelar algo somente em uma observação.

Uma variável dicotômica que toma valor 1 (um) somente para uma observação tem o efeito de excluir esta observação do cálculo dos coeficientes de mínimos quadrados e a variância do estimador (mas não o R^2).

- Diversas categorias (por exemplo, efeitos sazonais):

$$C_t = \beta_1 + \beta_2 x_t + \delta_1 D_{t1} + \delta_2 D_{t2} + \delta_3 D_{t3} + \varepsilon_t,$$

sendo x_t os ganhos.

- Uma formulação alternativa é dada por

$$C_t = \beta x_t + \delta_1 D_{t1} + \delta_2 D_{t2} + \delta_3 D_{t3} + \delta_4 D_{t4} + \varepsilon_t.$$

- Diversos grupos de variáveis dicotômicas.

$$y_{it} = \alpha + \beta x_{it} + \delta_i + \theta_t + \varepsilon_{it}.$$

- Precisamos retirar uma dicotômica de cada grupo.

Exemplo - companhias aéreas

O estudo trata da eficiência na produção de serviços de companhias aéreas.

Temos observações de 6 firmas durante 15 anos.

O modelo é dado por

$$\begin{aligned} \ln C_{i,t} &= \beta_1 + \beta_2 \ln Q_{i,t} + \beta_3 (\ln Q_{i,t})^2 + \beta_4 \ln P_{i,t} + \beta_5 \text{Cargas}_{i,t} \\ &+ \sum_{t=1}^{14} \theta_t D_{i,t} + \sum_{i=1}^5 \delta_i F_{i,t} + \varepsilon_{i,t} \end{aligned}$$

Ver arquivo exemplo_17.r

Em várias aplicações, variáveis dicotômicas são utilizadas para fatores qualitativos somente.

Suponha o interesse no seguinte modelo:

$$\text{renda} = \beta_1 + \beta_2 \text{idade} + \text{efeito de educação} + \varepsilon.$$

Ensino médio (EM), graduação (GR), mestrado (MS) e doutorado (DS).

Uma maneira **insatisfatória** de modelar é considerar a variável E igual a 0 para o primeiro grupo, 1 para o segundo, 2 para o terceiro e 3 para o quarto.

$$\text{renda} = \beta_1 + \beta_2 \text{idade} + \beta_3 E + \varepsilon.$$

Um modelo mais flexível seria um com variáveis dicotômicas:

$$\text{renda} = \beta_1 + \beta_2 \text{idade} + \delta_1 \text{GR} + \delta_2 \text{MS} + \delta_3 \text{DS} + \varepsilon.$$

$$\text{Ensino Médio: } E(\text{renda}|\text{idade}, \text{EM}) = \beta_1 + \beta_2 \text{idade},$$

$$\text{Graduação: } E(\text{renda}|\text{idade}, \text{GR}) = \beta_1 + \beta_2 \text{idade} + \delta_1,$$

$$\text{Mestrado: } E(\text{renda}|\text{idade}, \text{MS}) = \beta_1 + \beta_2 \text{idade} + \delta_2,$$

$$\text{Doutorado: } E(\text{renda}|\text{idade}, \text{DS}) = \beta_1 + \beta_2 \text{idade} + \delta_3.$$

Regressão *spline* (função definida segmentarmente por polinômios)

A função que desejamos estimar é

$$\begin{aligned} E(\text{renda}|\text{idade}) &= \alpha^0 + \beta^0 \text{idade}, & \text{se } \text{idade} < 18, \\ &\alpha^1 + \beta^1 \text{idade}, & \text{se } 18 \leq \text{idade} < 22, \\ &\alpha^2 + \beta^2 \text{idade}, & \text{se } \text{idade} \geq 22. \end{aligned}$$

Os valores 18 e 22 são chamados de nós. Seja $d_1 = 1$, se $\text{idade} \geq t_1^*$ e $d_2 = 1$, se $\text{idade} \geq t_2^*$, sendo $t_1^* = 18$ e $t_2^* = 22$. Temos a combinação das equações dada por

$$\text{renda} = \beta_1 + \beta_2 \text{idade} + \gamma_1 d_1 + \delta_1 d_1 \text{idade} + \gamma_2 d_2 + \delta_2 d_2 \text{idade} + \varepsilon.$$

Os coeficientes angulares dos três segmentos são β_2 , $\beta_2 + \delta_1$ e $\beta_2 + \delta_1 + \delta_2$.

Para fazer com que a função seja contínua, precisamos impor que

$$\begin{aligned} \beta_1 + \beta_2 t_1^* &= (\beta_1 + \gamma_1) + (\beta_2 + \delta_1) t_1^* & \text{e} \\ (\beta_1 + \gamma_1) + (\beta_2 + \delta_1) t_2^* &= (\beta_1 + \gamma_1 + \gamma_2) + (\beta_2 + \delta_1 + \delta_2) t_2^*. \end{aligned}$$

Isto implica restrições lineares nos coeficientes:

$$\gamma_1 + \delta_1 t_1^* = 0 \quad \text{and} \quad \gamma_2 + \delta_2 t_2^* = 0.$$

Assim,

$$\text{renda} = \beta_1 + \beta_2 \text{idade} + \delta_1 d_1 (\text{idade} - t_1^*) + \delta_2 d_2 (\text{idade} - t_2^*) + \varepsilon.$$

Os estimadores de mínimos quadrados restritos são obtidos pela regressão múltipla:

$$\begin{aligned} x_1 &= \text{idade}, \\ x_2 &= \text{idade} - 18, \quad \text{se } \text{idade} \geq 18 \quad \text{e } 0 \quad \text{caso contrário}, \\ x_3 &= \text{idade} - 22, \quad \text{se } \text{idade} \geq 22 \quad \text{e } 0 \quad \text{caso contrário}. \end{aligned}$$

Podemos testar a hipótese nula $H_0 : \delta_1 = \delta_2 = 0$ que implica mesmo coeficiente angular.

Não linearidade nas variáveis

Seja $\mathbf{z} = \{z_1, z_2, \dots, z_L\}$ um conjunto de L variáveis independentes.

Seja $f_1, f_2, \dots, f_p$ um conjunto de p funções linearmente independentes de $\mathbf{z}$.

Seja $g(y)$ uma função observável de y .

$$\begin{aligned}g(y) &= \beta_1 f_1(\mathbf{z}) + \beta_2 f_2(\mathbf{z}) + \dots + \beta_p f_p(\mathbf{z}) + \varepsilon \\&= \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_p x_p + \varepsilon \\&= \mathbf{x}^\top \boldsymbol{\beta} + \varepsilon.\end{aligned}$$

Um modelo comum é o log-linear

$$\ln y = \ln \alpha + \sum_k \beta_k \ln z_k + \varepsilon = \beta_1 + \sum_k \beta_k x_k + \varepsilon.$$

Neste modelo, os coeficientes angulares são elasticidades:

$$\left(\frac{\partial y}{\partial x_k} \right) \left(\frac{x_k}{y} \right) = \frac{\partial \ln y}{\partial \ln x_k} = \beta_k.$$

No modelo log-linear,

- Mudanças são medidas em termos proporcionais ou porcentagem.
- β_k mede a mudança percentual em y associada com a mudança de 1 por cento em x_k .

Termos de interação

Um modelo relacionando distância de frenagem D com velocidade S e humidade da pista W ,

$$D = \beta_1 + \beta_2 S + \beta_3 W + \beta_4 SW + \varepsilon.$$

Neste modelo,

$$\frac{\partial E(D|S, W)}{\partial S} = \beta_2 + \beta_4 W.$$

Isto implica que o efeito marginal da velocidade na distância de frenagem cresce quando a pista fica mais molhada.

$$\text{Var} \left(\frac{\partial \hat{E}(D|S, W)}{\partial S} \right) = \text{Var}(\hat{\beta}_2) + W^2 \text{Var}(\hat{\beta}_4) + 2W \text{Cov}(\hat{\beta}_2, \hat{\beta}_4).$$

Podemos tentar identificar não linearidades:

- Análise de resíduos;
- Regressão linear por partes; e
- Regressão polinomial.

Definição de linearidade intrínseca

No modelo de regressão linear clássico, se os p parâmetros $\beta_1, \beta_2, \dots, \beta_p$ podem ser escritos como p funções biunívocas, possivelmente não lineares, de um conjunto de p parâmetros $\theta_1, \theta_2, \dots, \theta_p$, então o modelo é intrinsecamente linear em θ .

Exemplo - regressão linear intrínseca

Considere a estimação dos parâmetros do modelo

$$f(y|x, \rho, \beta) = \frac{(\beta + x)^{-\rho}}{\Gamma(\rho)} y^{\rho-1} \exp\{-y/(\beta + x)\}$$

por máxima verossimilhança.

Ver arquivo exemplo_18.r

- Neste modelo, $E(y|x) = \beta\rho + \rho x$.
- Um modelo de regressão linear intrínseco seria $E(y|x) = \beta_1 + \beta_2 x$.
- Podemos estimar β_1 e β_2 por mínimos quadrados.
- Podemos estimar também β por b_1/b_2 .
- A estimação do erro padrão de β se dá através do método delta,

$$\widehat{\text{Var}}(\widehat{\beta}) = \left[\frac{\partial \widehat{\beta}}{\partial b_1} \right]^2 \widehat{\text{Var}}(b_1) + \left[\frac{\partial \widehat{\beta}}{\partial b_2} \right]^2 \widehat{\text{Var}}(b_2) + 2 \left[\frac{\partial \widehat{\beta}}{\partial b_1} \right] \left[\frac{\partial \widehat{\beta}}{\partial b_2} \right] \widehat{\text{Cov}}(b_1, b_2).$$

Exemplo - função de produção de elasticidade constante de substituição

O modelo é dado por

$$\ln y = \ln \gamma - \frac{\nu}{\rho} \ln(\delta K^{-\rho} + (1 - \delta)L^{-\rho}) + \varepsilon$$

A aproximação de Taylor em torno do ponto $\rho = 0$ é

$$\begin{aligned}\ln y &= \ln \gamma + \nu \delta \ln K + \nu(1 - \delta) \ln L + \rho \nu \delta(1 - \delta) \left\{ -\frac{1}{2} [\ln K - \ln L]^2 \right\} + \varepsilon^* \\ &= \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + \beta_4 x_4 + \varepsilon^*,\end{aligned}$$

sendo $x_1 = 1$, $x_2 = \ln K$, $x_3 = \ln L$ e $x_4 = -(1/2) \ln^2(K/L)$.

As transformações são dadas por

$$\begin{array}{ll}\beta_1 &= \ln \gamma & \gamma &= \exp(\beta_1) \\ \beta_2 &= \nu \delta & \delta &= \beta_2 / (\beta_2 + \beta_3) \\ \beta_3 &= \nu(1 - \delta) & \nu &= \beta_2 + \beta_3 \\ \beta_4 &= \rho \nu \delta(1 - \delta) & \rho &= \beta_4 (\beta_2 + \beta_3) / (\beta_2 \beta_3).\end{array}$$

Estimativas de β_1 , β_2 , β_3 e β_4 podem ser obtidas por mínimos quadrados.

Para estimar a matriz de covariância de $\theta = (\gamma, \delta, \nu, \rho)$ utilizamos o método delta.

$$\mathbf{C} = \frac{\partial \theta}{\partial \beta^\top} = \begin{bmatrix} \exp\{\beta_1\} & 0 & 0 & 0 \\ 0 & \frac{\beta_3}{(\beta_2 + \beta_3)^2} & -\frac{\beta_2}{(\beta_2 + \beta_3)^2} & 0 \\ 0 & 1 & 1 & 0 \\ 0 & -\frac{\beta_3\beta_4}{\beta_2^2\beta_3} & -\frac{\beta_2\beta_4}{\beta_2\beta_3^2} & \frac{\beta_2 + \beta_3}{\beta_2\beta_3} \end{bmatrix}.$$

A matriz de covariâncias de θ é, então, estimada por

$$\hat{\mathbf{C}} \left[s^2 (\mathbf{X}^\top \mathbf{X})^{-1} \right] \hat{\mathbf{C}}^\top.$$

Vale ressaltar que nem todos os modelos são intrinsecamente lineares.

Mudança estrutural

- Suponha o modelo de regressão múltipla $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$.
- Queremos testar a hipótese de que $\beta_1 = \beta_2$ para dois conjuntos diferentes da mesma variável.
- Podemos pensar no modelo

$$\begin{bmatrix} y_1 \\ y_2 \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1 & \mathbf{0} \\ \mathbf{0} & \mathbf{X}_2 \end{bmatrix} \begin{bmatrix} \beta_1 \\ \beta_2 \end{bmatrix} + \begin{bmatrix} \varepsilon_1 \\ \varepsilon_2 \end{bmatrix}$$

- A solução é

$$\mathbf{b} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = \begin{bmatrix} \mathbf{X}_1^\top \mathbf{X}_1 & \mathbf{0} \\ \mathbf{0} & \mathbf{X}_2^\top \mathbf{X}_2 \end{bmatrix} \begin{bmatrix} \mathbf{X}_1 \mathbf{y}_1 \\ \mathbf{X}_2 \mathbf{y}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{b}_1 \\ \mathbf{b}_2 \end{bmatrix}$$

- A soma total de quadrados é $\mathbf{e}^\top \mathbf{e} = \mathbf{e}_1^\top \mathbf{e}_1 + \mathbf{e}_2^\top \mathbf{e}_2$.
- Temos, $H_0 : \mathbf{R}\boldsymbol{\beta} = \mathbf{q}$, sendo $\mathbf{R} = [\mathbf{I} \quad -\mathbf{I}]$ e $\mathbf{q} = \mathbf{0}$.
- Prossegue-se com o teste F .
- O número de graus de liberdade é J e $n_1 + n_2 - 2p$. J é o número de colunas em $\mathbf{X}_2$.

Omissão de variáveis relevantes

Suponha que o modelo correto seja $\mathbf{y} = \mathbf{X}_1\beta_1 + \mathbf{X}_2\beta_2 + \varepsilon$.

Se regredimos $\mathbf{y}$ sobre $\mathbf{X}_1$, excluindo $\mathbf{X}_2$, temos

$$\mathbf{b}_1 = (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{y} = \beta_1 + (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{X}_2 \beta_2 + (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \varepsilon$$

$$E(\mathbf{b}_1|\mathbf{X}) = \beta_1 + \mathbf{P}_{1,2}\beta_2 \quad \text{e} \quad \text{Var}(\mathbf{b}_1|\mathbf{X}) = \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1},$$

sendo $\mathbf{P}_{1,2} = (\mathbf{X}_1^\top \mathbf{X}_1)^{-1} \mathbf{X}_1^\top \mathbf{X}_2$. Cada coluna de $\mathbf{P}_{1,2}$ é o resultado da regressão das colunas de $\mathbf{X}_2$ sobre $\mathbf{X}_1$.

Com $\mathbf{X}_2$ incluído na regressão, teríamos

$$\text{Var}(\mathbf{b}_{1,2}|\mathbf{X}) = \sigma^2(\mathbf{X}_1^\top \mathbf{M}_2 \mathbf{X}_1)^{-1} = \sigma^2[\mathbf{X}_1^\top \mathbf{X}_1 - \mathbf{X}_1^\top \mathbf{X}_2(\mathbf{X}_2^\top \mathbf{X}_2)^{-1} \mathbf{X}_2^\top \mathbf{X}_1]^{-1},$$

pois $\mathbf{M}_2 = \mathbf{I} - \mathbf{X}_2(\mathbf{X}_2^\top \mathbf{X}_2)^{-1} \mathbf{X}_2^\top$.

Podemos comparar as inversas das matrizes (precisão)

$$[\text{Var}(\mathbf{b}_1|\mathbf{X})]^{-1} - [\text{Var}(\mathbf{b}_{1,2}|\mathbf{X})]^{-1} = (1/\sigma^2) \mathbf{X}_1^\top \mathbf{X}_2(\mathbf{X}_2^\top \mathbf{X}_2)^{-1} \mathbf{X}_2^\top \mathbf{X}_1,$$

que é não negativa definida.

Concluimos que apesar de $\mathbf{b}_1$ ser tendencioso, sua variância nunca é maior do que a de $\mathbf{b}_{1,2}$.

Modelo de regressão generalizado e heterocedasticidade

Trataremos de modelos que não satisfazem a hipótese **HP.4** de homocedasticidade e correlação nula.

O modelo de regressão linear generalizado é dado por

$$\begin{aligned}\mathbf{y} &= \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon} \\ E(\boldsymbol{\varepsilon}|\mathbf{X}) &= \mathbf{0} \\ E(\boldsymbol{\varepsilon}\boldsymbol{\varepsilon}^\top|\mathbf{X}) &= \sigma^2\boldsymbol{\Omega} = \boldsymbol{\Sigma},\end{aligned}$$

em que $\boldsymbol{\Sigma}$ é uma matriz positiva definida.

No modelo (somente) heterocedástico:

$$\sigma^2\boldsymbol{\Omega} = \sigma^2 \begin{bmatrix} \omega_1 & 0 & \cdots & 0 \\ 0 & \omega_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \omega_n \end{bmatrix} = \begin{bmatrix} \sigma_1^2 & 0 & \cdots & 0 \\ 0 & \sigma_2^2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \sigma_n^2 \end{bmatrix}.$$

No modelo (somente) correlacionado:

$$\sigma^2 \mathbf{\Omega} = \sigma^2 \begin{bmatrix} 1 & \rho_1 & \cdots & \rho_{n-1} \\ \rho_1 & 1 & \cdots & \rho_{n-2} \\ \vdots & \vdots & \ddots & \vdots \\ \rho_{n-1} & \rho_{n-2} & \cdots & 1 \end{bmatrix}.$$

Em geral, aparece em modelos de séries temporais (modelo AR(1)).

Conjuntos de dados de painéis, consistindo de observações por seções cruzadas em vários pontos no tempo, podem exibir heterocedasticidade e correlação não nula.

Estimação por mínimos quadrados

Para o modelo com distúrbios esféricos,

$$E(\varepsilon|\mathbf{X}) = \mathbf{0} \quad \text{e} \quad E(\varepsilon\varepsilon^\top|\mathbf{X}) = \sigma^2\mathbf{I},$$

o estimador de mínimos quadrados ordinários (OLS) é

$$\mathbf{b} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = \boldsymbol{\beta} + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \varepsilon.$$

- o melhor estimador linear não tendencioso (BLUE);
- consistente;
- assintoticamente normal; e
- se os erros forem normais, então também é eficiente assintoticamente.

Teorema

Se os regressores e os distúrbios forem não correlacionados, então o estimador de mínimos quadrados para o modelo de regressão linear generalizado é não tendencioso. Condicional a $\mathbf{X}$, a variância amostral do estimador de mínimos quadrados é

$$\begin{aligned}\text{Var}(\mathbf{b}|\mathbf{X}) &= \text{E}((\mathbf{b} - \beta)(\mathbf{b} - \beta)^\top | \mathbf{X}) \\ &= \text{E}((\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \varepsilon \varepsilon^\top \mathbf{X} (\mathbf{X}^\top \mathbf{X})^{-1} | \mathbf{X}) \\ &= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \sigma^2 \boldsymbol{\Omega} \mathbf{X} (\mathbf{X}^\top \mathbf{X})^{-1} \\ &= \frac{\sigma^2}{n} \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right)^{-1} \left(\frac{1}{n} \mathbf{X}^\top \boldsymbol{\Omega} \mathbf{X} \right) \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right)^{-1}\end{aligned}$$

Se ε for normalmente distribuído, então

$$\mathbf{b} \sim \mathcal{N}(\beta, \sigma^2 (\mathbf{X}^\top \mathbf{X})^{-1} (\mathbf{X}^\top \boldsymbol{\Omega} \mathbf{X}) (\mathbf{X}^\top \mathbf{X})^{-1}).$$

- Não podemos mais estimar $\text{Var}(\mathbf{b}|\mathbf{X})$ com $s^2 (\mathbf{X}^\top \mathbf{X})^{-1}$; e
- Os testes t e F já não são mais apropriados.

Propriedades assintóticas de mínimos quadrados

Teorema

Se $\mathbf{Q} = \text{plim}(\mathbf{X}^\top \mathbf{X}/n)$ e $\text{plim}(\mathbf{X}^\top \mathbf{\Omega} \mathbf{X}/n)$ são matrizes positivas definidas e finitas, então $\mathbf{b}$ é consistente para β . Sob estas hipóteses, $\text{plim} \mathbf{b} = \beta$.

A condição no teorema acima depende de $\mathbf{X}$ e de $\mathbf{\Omega}$.

Teorema

Se os regressores forem suficientemente bem comportados e os termos fora da diagonal da matriz $\mathbf{\Omega}$ diminuirão suficientemente rápido, então o estimador de mínimos quadrados é assintoticamente normal com média β e matriz de covariâncias

$$\text{Var}(\mathbf{b}) = \frac{\sigma^2}{n} \mathbf{Q}^{-1} \text{plim} \left(\frac{1}{n} \mathbf{X}^\top \mathbf{\Omega} \mathbf{X} \right) \mathbf{Q}^{-1}.$$

Estimação robusta das matrizes de covariâncias assintóticas

Se $\sigma^2 \mathbf{\Omega}$ fosse conhecido, então o estimador da matriz de covariância assintótica de $\mathbf{b}$ seria

$$\mathbf{v}_{OLS} = \frac{1}{n} \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right)^{-1} \left(\frac{1}{n} \mathbf{X}^\top \left[\sigma^2 \mathbf{\Omega} \right] \mathbf{X} \right) \left(\frac{1}{n} \mathbf{X}^\top \mathbf{X} \right)^{-1}.$$

Nosso problema é $\sigma^2 \mathbf{\Omega}$. Assumiremos que $\text{tr}(\mathbf{\Omega}) = n$ e $\mathbf{\Sigma} = \sigma^2 \mathbf{\Omega}$.

Em princípio, parece que temos $n(n+1)/2$ parâmetros para estimar em $\mathbf{\Sigma}$.

Na verdade, precisamos estimar $p(p+1)/2$ parâmetros da matriz

$$\text{plim} \mathbf{Q}_* = \text{plim} \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^n \sigma_{ij} \mathbf{x}_i \mathbf{x}_j^\top.$$

Voltaremos a este ponto mais à frente.

Estimação eficiente por mínimos quadrados generalizados

Temos que Ω é uma matriz positiva definida, logo

$$\Omega = \mathbf{C}\mathbf{\Lambda}\mathbf{C}^\top,$$

sendo as colunas de $\mathbf{C}$ os vetores característicos de Ω e $\mathbf{\Lambda}$ é uma matriz diagonal com as raízes características de Ω .

Seja $\mathbf{T} = \mathbf{C}\mathbf{\Lambda}^{1/2}$, então $\Omega = \mathbf{T}\mathbf{T}^\top$.

Seja $\mathbf{P}^\top = \mathbf{C}\mathbf{\Lambda}^{-1/2}$, então $\Omega^{-1} = \mathbf{P}^\top \mathbf{P}$.

Segue-se que

$$\mathbf{P}\mathbf{y} = \mathbf{P}\mathbf{X}\boldsymbol{\beta} + \mathbf{P}\boldsymbol{\varepsilon} \quad \text{ou} \quad \mathbf{y}_* = \mathbf{X}_*\boldsymbol{\beta} + \boldsymbol{\varepsilon}_* \Rightarrow \text{E}(\boldsymbol{\varepsilon}_*\boldsymbol{\varepsilon}_*^\top | \mathbf{X}_*) = \mathbf{P}\sigma^2\Omega\mathbf{P}^\top = \sigma^2\mathbf{I}.$$

Então, o modelo de regressão linear clássico se aplica a este modelo transformado.

Por hipótese, Ω é conhecido, então $\mathbf{y}_*$ e $\mathbf{X}_*$ são os dados observados. Logo,

$$\hat{\beta} = (\mathbf{X}_*^\top \mathbf{X}_*)^{-1} \mathbf{X}_*^\top \mathbf{y}_* = (\mathbf{X}^\top \mathbf{P}^\top \mathbf{P} \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{P}^\top \mathbf{P} \mathbf{y} = (\mathbf{X}^\top \Omega^{-1} \mathbf{X})^{-1} \mathbf{X}^\top \Omega^{-1} \mathbf{y}$$

é um estimador eficiente de β . Este é o estimador de mínimos quadrados generalizados (GLS).

Teorema

1. Se $E(\varepsilon_* | \mathbf{X}_*) = \mathbf{0} \Leftrightarrow E(\mathbf{P}\varepsilon | \mathbf{P}\mathbf{X}) = \mathbf{0} \Leftrightarrow E(\varepsilon | \mathbf{X}) = \mathbf{0}$, então $E(\hat{\beta} | \mathbf{X}_*) = \beta$.
2. O estimador GLS é consistente se $\text{plim}(1/n) \mathbf{X}_*^\top \mathbf{X}_* = \mathbf{Q}_*$, sendo $\mathbf{Q}_*$ uma matriz positiva definida. Logo, $\text{plim}[(1/n) \mathbf{X}^\top \Omega^{-1} \mathbf{X}]^{-1} = \mathbf{Q}_*^{-1}$.
3. O estimador GLS é normalmente distribuído com média β e a variância amostral

$$\text{Var}(\hat{\beta} | \mathbf{X}_*) = \sigma^2 (\mathbf{X}_*^\top \mathbf{X}_*)^{-1} = \sigma^2 (\mathbf{X}^\top \Omega \mathbf{X})^{-1}.$$

4. O estimador GLS $\hat{\beta}$ é o estimador linear não tendencioso de variância mínima no modelo de regressão generalizado.

Para testar J restrições lineares, $\mathbf{R}\beta = \mathbf{q}$, a estatística apropriada é

$$F_{J,n-p} = \frac{(\mathbf{R}\hat{\beta} - \mathbf{q})^\top [\mathbf{R}\hat{\sigma}^2(\mathbf{X}_*^\top \mathbf{X}_*)^{-1} \mathbf{R}^\top]^{-1} (\mathbf{R}\hat{\beta} - \mathbf{q})}{J} = \frac{(\hat{\varepsilon}_c^\top \hat{\varepsilon}_c - \hat{\varepsilon}^\top \hat{\varepsilon})/J}{\hat{\sigma}^2},$$

sendo $\hat{\varepsilon} = \mathbf{y}_* - \mathbf{X}_* \hat{\beta}$ e $\hat{\sigma}^2 = \frac{\hat{\varepsilon}^\top \hat{\varepsilon}}{n-p} = \frac{(\mathbf{y} - \mathbf{X}\hat{\beta})^\top \Omega^{-1} (\mathbf{y} - \mathbf{X}\hat{\beta})}{n-p}$.

Os resíduos do GLS restrito, $\hat{\varepsilon}_c = \mathbf{y}_* - \mathbf{X}_* \hat{\beta}_c$, são baseados em

$$\hat{\beta}_c = \hat{\beta} - [\mathbf{X}^\top \Omega^{-1} \mathbf{X}]^{-1} [\mathbf{R}(\mathbf{X}^\top \Omega^{-1} \mathbf{X})^{-1} \mathbf{R}^\top]^{-1} (\mathbf{R}\hat{\beta} - \mathbf{q}).$$

Não existe nenhuma contrapartida precisa para o R^2 no modelo de regressão generalizado.

Se Ω for desconhecido, então não é possível utilizar o estimador GLS.

Precisamos impor restrições em Ω para diminuir o número efetivo de parâmetros a estimar.

Heterocedasticidade

Distúrbios da regressão cujas variâncias não são constantes para diferentes observações são heterocedásticos.

Em modelos de regressão heterocedásticos,

$$\text{Var}(\varepsilon_i|\mathbf{X}) = \sigma_i^2 = \sigma^2 \omega_i, \quad i = 1, 2, \dots, n \quad \text{ou}$$

$$\text{E}(\varepsilon \varepsilon^\top) = \sigma^2 \mathbf{\Omega} = \sigma^2 \begin{bmatrix} \omega_1 & 0 & \cdots & 0 \\ 0 & \omega_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \omega_n \end{bmatrix} = \begin{bmatrix} \sigma_1^2 & 0 & \cdots & 0 \\ 0 & \sigma_2^2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \sigma_n^2 \end{bmatrix},$$

sendo $\text{tr}(\mathbf{\Omega}) = \sum_{i=1}^n \omega_i = n$ para evitar problemas de escala.

Ineficiência de mínimos quadrados

- $\mathbf{b}$, o estimador OLS, é ineficiente relativo ao estimador GLS.
- Quanto maior a dispersão de ω_i para diferentes observações, maior a eficiência do GLS sobre OLS.
- Suponha que a forma da heterocedasticidade seja desconhecida.
- O estimador GLS não pode ser usado.
- Se usarmos $\mathbf{b}$, então a matriz de covariância $\sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}$ é inapropriada.
- A matriz $\sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}(\mathbf{X}^\top \mathbf{\Omega} \mathbf{X})(\mathbf{X}^\top \mathbf{X})^{-1}$ é apropriada.
- Temos, usualmente,

$$s^2 = \frac{\mathbf{e}^\top \mathbf{e}}{n-p} = \frac{\boldsymbol{\varepsilon}^\top \mathbf{M} \boldsymbol{\varepsilon}}{n-p} = \frac{\boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon}}{n-p} - \frac{\boldsymbol{\varepsilon}^\top \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon}}{n-p},$$

pois $\mathbf{M} = \mathbf{I} - \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$. Isto implica que

$$\mathbb{E} \left(\frac{\boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon}}{n-p} \middle| \mathbf{X} \right) = \frac{\text{tr}(\mathbb{E}(\boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^\top | \mathbf{X}))}{n-p} = \frac{n\sigma^2}{n-p} \quad \text{e}$$

$$\begin{aligned}
E\left(\frac{\boldsymbol{\varepsilon}^\top \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon}}{n-p} \middle| \mathbf{X}\right) &= \frac{\text{tr}(E[(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^\top \mathbf{X}])}{n-p} \\
&= \frac{\text{tr}(\sigma^2 (\mathbf{X}^\top \mathbf{X}/n)^{-1} (\mathbf{X}^\top \boldsymbol{\Omega} \mathbf{X}/n))}{n-p} \\
&= \frac{\sigma^2}{n-p} \text{tr}\left(\left(\frac{\mathbf{X}^\top \mathbf{X}}{n}\right)^{-1} \mathbf{Q}_n^*\right),
\end{aligned}$$

sendo $\mathbf{Q}_n^* = \frac{\mathbf{X}^\top \boldsymbol{\Omega} \mathbf{X}}{n} = \frac{1}{n} \sum_{i=1}^n \omega_i \mathbf{x}_i \mathbf{x}_i^\top$.

Se $\mathbf{b}$ é consistente, então $\lim_{n \rightarrow \infty} E(s^2) = \sigma^2$.

Se a heterocedasticidade não é correlacionada com as variáveis no modelo, então pelo menos em amostras grandes, os cálculos de mínimos quadrados, apesar de não serem a maneira ótima de utilizar os dados, não levarão a “erros grosseiros”.

Observações:

- É verdade que o OLS é ineficiente.
- Se por hipótese da análise - que a heterocedasticidade não é relacionada as variáveis do modelo - é incorreta, então os erros padrões convencionais podem estar longe dos valores apropriados.

Estamos interessados em um estimador para $\mathbf{Q}_* = \frac{1}{n} \sum_{i=1}^n \sigma_i^2 \mathbf{x}_i \mathbf{x}_i^\top$.

Pode ser mostrado sob condições bem gerais que o estimador

$$\mathbf{S}_0 = \frac{1}{n} \sum_{i=1}^n e_i^2 \mathbf{x}_i \mathbf{x}_i^\top$$

tem $\text{plim} \mathbf{S}_0 = \text{plim} \mathbf{Q}_*$.

Assim, o estimador consistente de heterocedasticidade de **White** é

$$\widehat{\text{Var}}(\mathbf{b}|\mathbf{X}) = n(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{S}_0 (\mathbf{X}^\top \mathbf{X})^{-1}.$$

Um dos problemas dos quadrados dos resíduos de OLS é que eles tendem a subestimar os quadrados dos distúrbios verdadeiros.

Davidson e McKinnon propuseram

1. escalar o resultado final por $n/(n - p)$; e
2. usar e_i^2/m_{ii} ao invés de e_i^2 , sendo $m_{ii} = 1 - \mathbf{x}_i^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{x}_i$.

Exemplo - estimadores de White, e de Davidson e McKinnon

Ver arquivo exemplo_19.r

Modelaremos os gastos mensais com cartão de crédito (G):

$$G = \beta_1 + \beta_2 \text{idade} + \beta_3 \text{casa} + \beta_4 \text{renda} + \beta_5 \text{renda}^2 + \varepsilon.$$

A variável “casa” é uma binária indicando casa própria ou não.

Testando a heterocedasticidade

O teste geral de White é dado por $H_0 : \sigma_i^2 = \sigma^2$ para todo i , e $H_1 : \text{Não } H_0$.

A dificuldade é em estimar um modelo com n parâmetros.

A matriz de covariância correta para o estimador de mínimos quadrados é

$$\text{Var}(\mathbf{b}|\mathbf{X}) = \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}(\mathbf{X}^\top \mathbf{\Omega} \mathbf{X})(\mathbf{X}^\top \mathbf{X})^{-1}$$

que pode ser estimada como anteriormente.

O estimador convencional é $\mathbf{V} = s^2(\mathbf{X}^\top \mathbf{X})^{-1}$.

Se não existir heterocedasticidade, então $\mathbf{V}$ dará um estimador consistente de $\text{Var}(\mathbf{b}|\mathbf{X})$.

Uma versão simples e operacional do teste de White é feita ao obter nR^2 na regressão de e_i^2 sobre uma constante, todas as variáveis em $\mathbf{x}_i$ e todos os quadrados e produtos cruzados de $\mathbf{x}_i$.

A estatística tem distribuição assintótica qui-quadrada com $r - 1$ graus de liberdade, sendo r o número de regressores na equação, incluindo a constante.

O teste de White não é construtivo. Se H_0 é rejeitada, então o resultado do teste não dá nenhuma indicação do que fazer a seguir.

O teste de **Breusch-Pagan** é baseado nos multiplicadores de Lagrange e tem $H_0 : \sigma_i^2 = \sigma^2 f(\alpha_0 + \alpha^\top \mathbf{z}_i)$, sendo $\mathbf{z}_i$ o vetor de variáveis independentes.

O modelo é homocedástico se $\alpha = \mathbf{0}$.

O teste pode ser conduzido com uma regressão simples.

A estatística do teste é

$$LM = \frac{1}{2} \mathbf{g}^\top \mathbf{Z} (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{g},$$

sendo $\mathbf{g}$ o vetor de observações $g_i = e_i^2 / (\mathbf{e}^\top \mathbf{e} / n) - 1$.

Sob H_0 , $LM \approx \chi_r^2$, sendo r o número de variáveis em $\mathbf{z}$.

Este teste é sensível a hipótese de normalidade.

Koenker e Basset sugeriram utilizar um estimador robusto para a variância de ε_i^2 ,

$$V = \frac{1}{n} \sum_{i=1}^n \left[e_i^2 - \frac{\mathbf{e}^\top \mathbf{e}}{n} \right]^2.$$

Assim,

$$LM = \frac{1}{V} (\mathbf{u} - \bar{u}\mathbf{1})^\top \mathbf{Z} (\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top (\mathbf{u} - \bar{u}\mathbf{1}),$$

sendo $\mathbf{u} = (e_1^2, e_2^2, \dots, e_n^2)$ e $\bar{u} = \mathbf{e}^\top \mathbf{e} / n$.

Ver arquivo [exemplo_19.r](#)

Mínimos quadrados ponderados

Tendo testado e encontrado evidência de heterocedasticidade, o próximo passo é levar isto em conta na estimação.

O estimador GLS é

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

Suponha que $\text{Var}(\varepsilon_i | \mathbf{X}) = \sigma_i^2 = \sigma^2 \omega_i$.

Então, o i -ésimo elemento da diagonal da matriz $\mathbf{\Omega}^{-1}$ é $1/\omega_i$.

O estimador GLS é obtido ao regredir

$$\mathbf{Py} = \begin{bmatrix} y_1 / \sqrt{\omega_1} \\ y_2 / \sqrt{\omega_2} \\ \vdots \\ y_n / \sqrt{\omega_n} \end{bmatrix} \quad \text{sobre} \quad \mathbf{PX} = \begin{bmatrix} \mathbf{x}_1^\top / \sqrt{\omega_1} \\ \mathbf{x}_2^\top / \sqrt{\omega_2} \\ \vdots \\ \mathbf{x}_n^\top / \sqrt{\omega_n} \end{bmatrix}.$$

Aplicando-se o OLS ao modelo transformado, nós obtemos o estimador de mínimos quadrados ponderados (WLS),

$$\hat{\beta} = \left[\sum_{i=1}^n w_i \mathbf{x}_i \mathbf{x}_i^\top \right]^{-1} \left[\sum_{i=1}^n w_i \mathbf{x}_i y_i \right],$$

sendo $w_i = 1/\omega_i$.

Uma especificação comum é que a variância seja proporcional a um dos regressores ou de seu quadrado, por exemplo se $\sigma_i^2 = \sigma^2 x_{ik}^2$, então o modelo de regressão transformado para GLS é

$$\frac{y}{x_k} = \beta_k + \beta_1 \frac{x_1}{x_k} + \beta_2 \frac{x_2}{x_k} + \cdots + \frac{\varepsilon}{x_k}.$$

O estimador WLS é consistente seja qual for os pesos utilizados, mas os pesos não podem ser correlacionados com os distúrbios.

Entretanto, a escolha dos pesos deve ser feita com cuidado para não “afetar em demasia” os erros padrões do estimador GLS.

Modelos para dados de painel

Dados que combinam séries temporais e seções cruzadas são comuns em economia.

Temos um modelo de regressão da forma:

$$\begin{aligned}y_{it} &= \mathbf{x}_{it}^{\top} \boldsymbol{\beta} + \mathbf{z}_i^{\top} \boldsymbol{\alpha} + \varepsilon_{it} \\ &= \mathbf{x}_{it}^{\top} \boldsymbol{\beta} + c_i + \varepsilon_{it}.\end{aligned}$$

- $\mathbf{x}_{it}$ inclui p regressores, excluindo o termo constante.
- A heterogeneidade ou o efeito do indivíduo é $\mathbf{z}_i^{\top} \boldsymbol{\alpha}$ com $\mathbf{z}_i$ contendo o termo constante e um conjunto de variáveis individuais ou em grupo (raça, sexo, localização, etc.).
- Exogeneidade estrita: $E(\varepsilon_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots) = 0$.
- Independência na média: $E(c_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots) = \boldsymbol{\alpha}$.

Estruturas do modelo

1. **Regressão de coleção (pooled)**: se $\mathbf{z}_i$ contiver somente um termo constante, então OLS resulta em estimativas consistentes e eficientes de α e β .
2. **Efeitos fixos**: se $\mathbf{z}_i$ não é observada, mas correlacionada com $\mathbf{x}_{it}$, então o estimador de mínimos quadrados de β é tendencioso e inconsistente como consequência da omissão de uma variável.

O modelo $y_{it} = \mathbf{x}_{it}^\top \beta + \alpha_i + \varepsilon_{it}$ tem α_i como um termo constante para cada grupo i .

Efeito fixo está relacionado com a correlação de c_i e $\mathbf{x}_{it}$.

3. **Efeitos aleatórios**: se a heterogeneidade individual não observável, entretanto formulada, for por hipótese não correlacionada com as variáveis incluídas, então o modelo pode ser formulado como:

$$\begin{aligned} y_{it} &= \mathbf{x}_{it}^\top \beta + E(\mathbf{z}_i^\top \alpha) + \{\mathbf{z}_i^\top \alpha - E(\mathbf{z}_i^\top \alpha)\} + \varepsilon_{it} \\ &= \mathbf{x}_{it}^\top \beta + \alpha + u_i + \varepsilon_{it}. \end{aligned}$$

4. **Parâmetros aleatórios**: $y_{it} = \mathbf{x}_{it}^\top (\beta + \mathbf{h}_i) + (\alpha + u_i) + \varepsilon_{it}$.

Modelo de regressão de coleção (*pooled*)

Temos o seguinte modelo e hipóteses:

$$y_{it} = \alpha + \mathbf{x}_{it}^{\top} \boldsymbol{\beta} + \varepsilon_{it}, \quad i = 1, \dots, n, \quad t = 1, \dots, T_i$$

$$E(\varepsilon_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT_i}) = 0$$

$$\text{Var}(\varepsilon_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT_i}) = \sigma_{\varepsilon}^2$$

$$\text{Cov}(\varepsilon_{it}, \varepsilon_{js} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT_i}) = 0 \text{ se } i \neq j \text{ ou } t \neq s.$$

Nesta forma, se as outras hipóteses do modelo de regressão linear clássico são satisfeitas, então podemos aplicar os resultados anteriores de mínimos quadrados ordinários.

Estimação por mínimos quadrados

- As hipóteses para a aplicação de mínimos quadrados ordinários são raramente satisfeitas.
- No modelo de efeitos aleatórios, $y_{it} = c_i + \mathbf{x}_{it}^\top \boldsymbol{\beta} + \varepsilon_{it}$, sendo $E(c_i | \mathbf{X}_i) = \alpha$, temos

$$\begin{aligned}y_{it} &= \alpha + \mathbf{x}_{it}^\top \boldsymbol{\beta} + \varepsilon_{it} + (c_i - E(c_i | \mathbf{X}_i)) \\&= \alpha + \mathbf{x}_{it}^\top \boldsymbol{\beta} + \varepsilon_{it} + u_i \\&= \alpha + \mathbf{x}_{it}^\top \boldsymbol{\beta} + \omega_{it}.\end{aligned}$$

Deste modo, $E(\omega_{it}\omega_{is}) = \sigma_u^2$ quando $t \neq s$ (correlacionado).

- O estimador de mínimos quadrados ordinários no modelo de regressão generalizado pode ser consistente.
- Entretanto, o estimador convencional da variância assintótica pode subestimar a verdadeira.

Estimação robusta da matriz de covariância

- Suponha o modelo (agrupando observações)

$$\mathbf{y}_i = \mathbf{X}_i\boldsymbol{\beta} + \boldsymbol{\omega}_i.$$

- Pode existir heterocedasticidade entre os indivíduos.
- Entretanto, o maior problema de dados de painel é a correlação entre as observações ou a autocorrelação.
- Suponha que $\text{Var}(\boldsymbol{\omega}_i|\mathbf{X}_i) = \boldsymbol{\Omega}_i$.
- O estimador de mínimos quadrados ordinários de $\boldsymbol{\beta}$ é

$$\begin{aligned}\mathbf{b} &= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} = \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1} \sum_{i=1}^n \mathbf{x}_i^\top \mathbf{y}_i \\ &= \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1} \sum_{i=1}^n \mathbf{x}_i^\top (\mathbf{x}_i \boldsymbol{\beta} + \boldsymbol{\omega}_i) \\ &= \boldsymbol{\beta} + \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1} \sum_{i=1}^n \mathbf{x}_i^\top \boldsymbol{\omega}_i.\end{aligned}$$

A variância assintótica

$$\begin{aligned}\text{Var}(\mathbf{b}) &= \text{plim} \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1} \text{plim} \left[\sum_{i=1}^n \mathbf{x}_i^\top \boldsymbol{\omega}_i \boldsymbol{\omega}_i^\top \mathbf{x}_i \right] \text{plim} \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1} \\ &= \text{plim} \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1} \text{plim} \left[\sum_{i=1}^n \mathbf{x}_i^\top \boldsymbol{\Omega}_i \mathbf{x}_i \right] \text{plim} \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1}\end{aligned}$$

A matriz no centro deve ser estimada, então

$$\widehat{\text{Var}}(\mathbf{b}) = \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1} \left[\sum_{i=1}^n \mathbf{x}_i^\top \widehat{\boldsymbol{\omega}}_i \widehat{\boldsymbol{\omega}}_i^\top \mathbf{x}_i \right] \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i \right]^{-1}.$$

Este estimador da variância assintótica leva mais em conta a correlação entre as observações cruzadas e não a heterocedasticidade.

Exemplo - salários

Trataremos da modelagem dos salários de chefes de família entre 1976 e 1982. Para isto, são utilizadas diversas covariáveis.

Ver arquivo [exemplo_20.r](#)

O modelo de efeitos fixos

O modelo de efeitos fixos é dado por

$$y_{it} = \mathbf{x}_{it}^{\top} \boldsymbol{\beta} + \alpha_i + \varepsilon_{it}$$

A formulação de efeitos fixos implica que a diferença entre grupos pode ser capturada em diferenças no termo constante.

Cada α_i é tratado como um parâmetro a ser estimado.

Problema: qualquer variável invariante no tempo terá o efeito do termo constante específico do indivíduo.

Os coeficientes nas variáveis invariantes no tempo não podem ser estimados.

Esta falta de identificação é o preço que se paga por robustez da especificação de correlação não medida entre os efeitos comuns e as variáveis exógenas.

Estimação por mínimos quadrados

Vamos supor que os dados sejam balanceados e que $\mathbf{y}_i$ e $\mathbf{X}_i$ os dados agrupados de cada grupo i .

$$\mathbf{y}_i = \mathbf{X}_i\boldsymbol{\beta} + \mathbf{1}\alpha_i + \boldsymbol{\varepsilon}_i.$$

Temos,

$$\begin{bmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \\ \vdots \\ \mathbf{y}_n \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1 \\ \mathbf{X}_2 \\ \vdots \\ \mathbf{X}_n \end{bmatrix} \boldsymbol{\beta} + \begin{bmatrix} \mathbf{1} & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \mathbf{1} & \cdots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \cdots & \mathbf{1} \end{bmatrix} \begin{bmatrix} \alpha_1 \\ \alpha_2 \\ \vdots \\ \alpha_n \end{bmatrix} + \begin{bmatrix} \boldsymbol{\varepsilon}_1 \\ \boldsymbol{\varepsilon}_2 \\ \vdots \\ \boldsymbol{\varepsilon}_n \end{bmatrix}$$

ou

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{D}\boldsymbol{\alpha} + \boldsymbol{\varepsilon} = [\mathbf{X} \quad \mathbf{d}_1 \quad \mathbf{d}_2 \quad \cdots \quad \mathbf{d}_n] \begin{bmatrix} \boldsymbol{\beta} \\ \boldsymbol{\alpha} \end{bmatrix} + \boldsymbol{\varepsilon},$$

sendo $\mathbf{D} = [\mathbf{d}_1 \quad \mathbf{d}_2 \quad \cdots \quad \mathbf{d}_n]$ e $\mathbf{d}_i$ uma variável indicadora da i -ésima unidade.

Este é um modelo de regressão clássico. Portanto, não é necessário nenhum resultado novo para analisá-lo.

O modelo de efeitos aleatórios

O modelo de efeitos fixos permite que o efeito individual não observado seja correlacionado com as variáveis incluídas.

Uma reformulação do modelo é

$$y_{it} = \mathbf{x}_{it}^{\top} \boldsymbol{\beta} + (\alpha + u_i) + \varepsilon_{it}$$

com p regressores, incluindo uma constante. Temos,

$$u_i = \mathbf{z}_i^{\top} \boldsymbol{\alpha} - E(\mathbf{z}_i^{\top} \boldsymbol{\alpha}) \quad \text{e} \quad E(\mathbf{z}_i^{\top} \boldsymbol{\alpha}) = \alpha$$

Vamos supor exogeneidade estrita:

$$\begin{aligned} E(\varepsilon_{it} | \mathbf{X}) &= E(u_i | \mathbf{X}) = 0, \\ E(\varepsilon_{it}^2 | \mathbf{X}) &= \sigma_{\varepsilon}^2, \quad E(u_i^2 | \mathbf{X}) = \sigma_u^2, \\ E(\varepsilon_{it} u_j | \mathbf{X}) &= 0, \quad \text{para todo } i, t \text{ e } j, \\ E(\varepsilon_{it} \varepsilon_{js} | \mathbf{X}) &= 0, \quad \text{se } t \neq s \text{ e } i \neq j, \\ E(u_i u_j | \mathbf{X}) &= 0, \quad \text{se } i \neq j. \end{aligned}$$

Agruparemos em blocos de T observações, $\mathbf{y}_i$, $\mathbf{X}_i$, $u_i\mathbf{1}$ e ε_i .

$$\begin{aligned}\eta_{it} &= \varepsilon_{it} + u_i, & \boldsymbol{\eta}_i &= (\eta_{i1}, \eta_{i2}, \dots, \eta_{iT})^\top \\ \mathbb{E}(\eta_{ij}^2 | \mathbf{X}) &= \sigma_\varepsilon^2 + \sigma_u^2, & \mathbb{E}(\eta_{it}\eta_{is} | \mathbf{X}) &= \sigma_u^2 \quad t \neq s, \\ \mathbb{E}(\eta_{it}\eta_{js} | \mathbf{X}) &= \sigma_u^2 \quad t \neq s & \text{para todo } t \text{ e } s \text{ se } i \neq j.\end{aligned}$$

Assim,

$$\mathbb{E}(\boldsymbol{\eta}_i \boldsymbol{\eta}_i^\top | \mathbf{X}) = \boldsymbol{\Sigma} = \begin{bmatrix} \sigma_\varepsilon^2 + \sigma_u^2 & \sigma_u^2 & \cdots & \sigma_u^2 \\ \sigma_u^2 & \sigma_\varepsilon^2 + \sigma_u^2 & \cdots & \sigma_u^2 \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_u^2 & \sigma_u^2 & \cdots & \sigma_\varepsilon^2 + \sigma_u^2 \end{bmatrix} = \sigma_\varepsilon^2 \mathbf{I}_T + \sigma_u^2 \mathbf{1}\mathbf{1}^\top.$$

As observações i e j são independentes, logo a matriz de covariância para as nT observações é $\boldsymbol{\Omega} = \mathbf{I}_n \otimes \boldsymbol{\Sigma}$.

Mínimos quadrados generalizados

$$\hat{\beta} = (\mathbf{X}^\top \mathbf{\Omega}^{-1} \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{\Omega}^{-1} \mathbf{y} = \left[\sum_{i=1}^n \mathbf{x}_i^\top \mathbf{\Sigma}^{-1} \mathbf{x}_i \right]^{-1} \sum_{i=1}^n \mathbf{x}_i^\top \mathbf{\Sigma}^{-1} \mathbf{y}_i.$$

Precisamos estimar a matriz $\mathbf{\Sigma}$. Uma abordagem para estimar as componentes de variância é a seguinte:

- Estimaremos $\sigma_\eta^2 = \sigma_\varepsilon^2 + \sigma_u^2$ conjuntamente através da variância do modelo de regressão de coleção (*pooled*).
- Estimaremos σ_ε^2 através da variância do modelo de regressão com efeitos fixos.
- Finalmente, estimaremos σ_u^2 pela relação $\sigma_u^2 = \sigma_\eta^2 - \sigma_\varepsilon^2$.

Após estimar a matriz de variância $\mathbf{\Sigma}$, estimamos β pela equação acima e

$$\widehat{\text{Var}}(\hat{\beta}|\mathbf{X}) = \left[\sum_{i=1}^n \mathbf{x}_i^\top \hat{\mathbf{\Sigma}}^{-1} \mathbf{x}_i \right]^{-1}.$$

Exemplo - salários (revisitado)

Trataremos da modelagem dos salários de chefes de família entre 1976 e 1982. Para isto, são utilizadas diversas covariáveis.

Ver arquivo exemplo_21.r

Regressão não linear e mínimos quadrados não lineares

A forma geral do modelo de regressão não linear é

$$y_i = h(\mathbf{x}_i, \boldsymbol{\beta}) + \varepsilon_i.$$

Existem modelos que podem ser linearizados como

$$y = e^{\beta_0} x_1^{\beta_1} x_2^{\beta_2} e^{\varepsilon}.$$

Entretanto, estamos interessados em modelos em que este tipo de transformação não é possível.

Exemplo

Modelo de função de produção da elasticidade constante de substituição:

$$\ln y = \ln \gamma - \frac{\nu}{\rho} \ln (\delta K^{-\rho} + (1 - \delta)L^{-\rho}) + \varepsilon.$$

Hipóteses do modelo de regressão não linear

- HP.1** Forma funcional: $E(y_i|\mathbf{x}_i) = h(\mathbf{x}_i, \beta)$, $i = 1, 2, \dots, n$, sendo $h(\mathbf{x}_i, \beta)$ uma função diferenciável continuamente de β .
- HP.2** Identificação: o vetor de parâmetros do modelo é identificado (estimável) se existir um parâmetro não nulo $\beta^0 \neq \beta$ tal que $h(\mathbf{x}_i, \beta^0) = h(\mathbf{x}_i, \beta)$ para todo $\mathbf{x}_i$.
- HP.3** Média zero dos distúrbios: pela Hipótese 1 temos $y_i = h(\mathbf{x}_i, \beta) + \varepsilon_i$ sendo $E(\varepsilon_i|h(\mathbf{x}_i, \beta)) = 0$.
- HP.4** Homocedasticidade e autocorrelação nula:
 $E(\varepsilon_i^2|h(\mathbf{x}_j, \beta), j = 1, \dots, n) = \sigma^2 < \infty$ e
 $E(\varepsilon_i \varepsilon_j|h(\mathbf{x}_j, \beta), j = 1, \dots, n) = 0$ para todo $i \neq j$.
- HP.5** Processo de geração de dados: $\mathbf{x}_i$ é estritamente exógena ao processo ε_i e tem os dois primeiros momentos bem definidos e convergentes.
- HP.6** Modelo de probabilidade: existe uma distribuição de probabilidade bem definida para ε_i .

Condição de primeira ordem para um modelo não linear

Exemplo

Seja o modelo

$$y_i = \beta_1 + \beta_2 \exp\{\beta_3 x_i\} + \varepsilon_i.$$

Então,

$$S(\mathbf{b}) = \sum_{i=1}^n (y_i - b_1 - b_2 \exp\{b_3 x_i\})^2,$$

tal que

$$\frac{\partial S(\mathbf{b})}{\partial b_1} = - \sum_{i=1}^n (y_i - b_1 - b_2 \exp\{b_3 x_i\}) = 0$$

$$\frac{\partial S(\mathbf{b})}{\partial b_2} = - \sum_{i=1}^n (y_i - b_1 - b_2 \exp\{b_3 x_i\}) \exp\{b_3 x_i\} = 0$$

$$\frac{\partial S(\mathbf{b})}{\partial b_3} = - \sum_{i=1}^n (y_i - b_1 - b_2 \exp\{b_3 x_i\}) b_2 x_i \exp\{b_3 x_i\} = 0,$$

não tem solução analítica.

Definição

Um modelo de regressão não linear é um na qual as condições de primeira ordem de estimação por mínimos quadrados dos parâmetros são funções não lineares dos parâmetros.

Regressão linearizada

Seja o modelo de regressão linear é

$$y = h(\mathbf{x}, \boldsymbol{\beta}) + \varepsilon.$$

A expansão de Taylor em torno de $\boldsymbol{\beta}^0$:

$$\begin{aligned} h(\mathbf{x}, \boldsymbol{\beta}) &\simeq h(\mathbf{x}, \boldsymbol{\beta}^0) + \sum_{k=1}^K \frac{\partial h(\mathbf{x}, \boldsymbol{\beta}^0)}{\partial \beta_k^0} (\beta_k - \beta_k^0) \\ &= \left[h(\mathbf{x}, \boldsymbol{\beta}^0) - \sum_{k=1}^K \beta_k^0 \left(\frac{\partial h(\mathbf{x}, \boldsymbol{\beta}^0)}{\partial \beta_k^0} \right) \right] + \sum_{k=1}^K \beta_k \left(\frac{\partial h(\mathbf{x}, \boldsymbol{\beta}^0)}{\partial \beta_k^0} \right). \end{aligned}$$

Seja $x_k^0 = \frac{\partial h(\mathbf{x}, \boldsymbol{\beta}^0)}{\partial \beta_k^0}$. Dado $\boldsymbol{\beta}^0$, temos que x_k^0 é função dos dados.

$$\begin{aligned}
h(\mathbf{x}, \boldsymbol{\beta}) &\simeq \left[h^0 - \sum_{k=1}^K x_k^0 \beta_k^0 \right] + \sum_{k=1}^K x_k^0 \beta_k = h^0 - \mathbf{x}^{0\top} \boldsymbol{\beta}^0 + \mathbf{x}^{0\top} \boldsymbol{\beta} \Rightarrow \\
y &\simeq h^0 - \mathbf{x}^{0\top} \boldsymbol{\beta}^0 + \mathbf{x}^{0\top} \boldsymbol{\beta} + \varepsilon \Rightarrow \\
y^0 &= y - h^0 + \mathbf{x}^{0\top} \boldsymbol{\beta}^0 = \mathbf{x}^{0\top} \boldsymbol{\beta} + \varepsilon^0.
\end{aligned}$$

ε^0 contem o distúrbio verdadeiro mais o erro de aproximação de Taylor de primeira ordem.

Conhecendo $\boldsymbol{\beta}^0$, podemos calcular y^0 e x^0 e estimar os parâmetros $\boldsymbol{\beta}$ por mínimos quadrados.

Exemplo - regressão linearizada

$$y_i = \beta_1 + \beta_2 \exp\{\beta_3 x_i\} + \varepsilon_i.$$

$$x_1^0 = \frac{\partial h(\cdot)}{\partial \beta_1} = 1, \quad x_2^0 = \frac{\partial h(\cdot)}{\partial \beta_2} = \exp\{\beta_3^0 x\}, \quad \text{e} \quad x_3^0 = \frac{\partial h(\cdot)}{\partial \beta_3} = \beta_2^0 x \exp\{\beta_3^0 x\}.$$

$$y^0 = y - h(x, \beta_1^0, \beta_2^0, \beta_3^0) + \beta_1^0 x_1^0 + \beta_2^0 x_2^0 + \beta_3^0 x_3^0.$$

Sob certas condições, o estimador de mínimos quadrados não linear é consistente.

Além disto, $\mathbf{b} \approx \mathcal{N}\left(\boldsymbol{\beta}, \frac{\sigma^2}{n}[\mathbf{Q}^0]^{-1}\right)$ sendo $\mathbf{Q}^0 = \text{plim} \frac{1}{n} \mathbf{X}^{0\top} \mathbf{X}^0$ e

$$\widehat{\text{Var}}(\mathbf{b}) = \widehat{\sigma}^2 (\mathbf{X}^{0\top} \mathbf{X}^0)^{-1}.$$

Solução iterativa

$$\begin{aligned} b_{t+1} &= \left[\sum_{i=1}^n \mathbf{x}_i^0 \mathbf{x}_i^{0\top} \right]^{-1} \left[\sum_{i=1}^n \mathbf{x}_i^0 (y_i - h_i^0 + \mathbf{x}_i^{0\top} b_t) \right] \\ &= b_t + \left[\sum_{i=1}^n \mathbf{x}_i^0 \mathbf{x}_i^{0\top} \right]^{-1} \left[\sum_{i=1}^n \mathbf{x}_i^0 (y_i - h_i^0) \right] \\ &= b_t + (\mathbf{X}^{0\top} \mathbf{X}^0)^{-1} \mathbf{X}^{0\top} \mathbf{e}^0 \\ &= b_t + \boldsymbol{\Delta}_t. \end{aligned}$$

Um estimador consistente de σ^2 é

$$\widehat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (y_i - h(\mathbf{x}_i, \mathbf{b}))^2.$$

Estimação por variáveis instrumentais

- Temos trabalhado com a hipótese de que $\mathbf{x}_i$ e ε_i são não correlacionados no modelo de regressão linear.
- Existem aplicações em economia onde isto não é possível.
- Sem esta hipótese, nenhuma das provas de consistência ou não tendenciosidade do estimador de mínimos quadrados permanecem válidos.

Exemplo - modelos de dados de painel dinâmico.

$$y_{it} = \mathbf{x}_{it}^{\top} \boldsymbol{\beta} + \gamma y_{i,t-1} + \varepsilon_{it} + u_i.$$

Temos $y_{i,t-1}$ é correlacionado com o distúrbio $\varepsilon_{it} + u_i$.

- Podemos utilizar o método de variáveis instrumentais.

Hipóteses do modelo

Das hipóteses anteriores, manteremos o resultado

$$\text{plim}(\mathbf{X}^\top \mathbf{X}/n) = \mathbf{Q}_{xx} \quad \text{e} \quad \mathbf{HP.3I} : E(\varepsilon_i | \mathbf{x}_i) = \eta_i \quad \Rightarrow \quad E(\mathbf{x}_i \varepsilon_i) = \boldsymbol{\gamma}.$$

Sob certas condições de regularidade, $\text{plim}(1/n)\mathbf{X}^\top \boldsymbol{\varepsilon} = \boldsymbol{\gamma}$.

Assim, os regressores $\mathbf{X}$ não são mais exógenos.

Agora, suporemos que existe um conjunto de variáveis $\mathbf{Z}$ tal que

1 - Exogeneidade: elas são não correlacionadas com os distúrbios; e

2 - Relevância: elas são correlacionadas com as variáveis $\mathbf{X}$.

Estas (variáveis $\mathbf{Z}$) são chamadas de variáveis instrumentais.

Hipóteses:

HP.7I: $[\mathbf{x}_i, \mathbf{z}_i, \varepsilon_i]$, $i = 1, \dots, n$, formam uma sequência iid de variáveis aleatórias.

HP.8Ia: $E(x_{ik}^2) = \mathbf{Q}_{xx, kk} < \infty$, uma constante finita, $k = 1, \dots, K$.

HP.8Ib: $E(z_{i\ell}^2) = \mathbf{Q}_{zz, \ell\ell} < \infty$, uma constante finita, $\ell = 1, \dots, L$.

HP.8Ic: $E(z_{i\ell} x_{ik}) = \mathbf{Q}_{zx, \ell k} < \infty$, uma constante finita, $\ell = 1, \dots, L$ e $k = 1, \dots, K$.

HP.9I: $E(\varepsilon_i | \mathbf{z}_i) = 0$.

Além disso,

- $\text{plim}(1/n) \mathbf{Z}^\top \mathbf{Z} = \mathbf{Q}_{zz}$,
- $\text{plim}(1/n) \mathbf{Z}^\top \mathbf{X} = \mathbf{Q}_{zx}$ (relevância), e
- $\text{plim}(1/n) \mathbf{Z}^\top \boldsymbol{\varepsilon} = \mathbf{0}$ (exogeneidade).

Estimação

Mínimos quadrados ordinários

$$E(\mathbf{b}|\mathbf{X}) = \beta + (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\eta} \neq \beta \quad (\text{tendencioso})$$

$$\text{plim} \mathbf{b} = \beta + \text{plim} \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)^{-1} \text{plim} \left(\frac{\mathbf{X}^\top \boldsymbol{\varepsilon}}{n} \right) = \beta + \mathbf{Q}_{xx}^{-1} \boldsymbol{\gamma} \neq \beta \quad (\text{inconsistente}).$$

Estimador por variáveis instrumentais

Temos, $E(\mathbf{z}_i \varepsilon_i) = 0$, então

$$\text{plim} \left(\frac{\mathbf{Z}^\top \boldsymbol{\varepsilon}}{n} \right) = \text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{y}}{n} \right) - \text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{X} \beta}{n} \right) = \mathbf{0}.$$

Portanto,

$$\text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{y}}{n} \right) = \left[\text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{X}}{n} \right) \right] \beta + \text{plim} \left(\frac{\mathbf{Z}^\top \boldsymbol{\varepsilon}}{n} \right) = \left[\text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{X}}{n} \right) \right] \beta.$$

Suponha que $\mathbf{Z}$ tenha o mesmo número de variáveis do que $\mathbf{X}$.

O posto de $\mathbf{Z}^\top \mathbf{X}$ é p , e $\mathbf{Z}^\top \mathbf{X}$ é uma matriz quadrada. Logo,

$$\left[\text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{X}}{n} \right) \right]^{-1} \text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{y}}{n} \right) = \beta.$$

que nos leva ao **estimador por variáveis instrumentais**

$$\mathbf{b}_{IV} = (\mathbf{Z}^\top \mathbf{X})^{-1} \mathbf{Z}^\top \mathbf{y}.$$

Já provamos que $\mathbf{b}_{IV}$ é consistente.

A distribuição assintótica de $\mathbf{b}_{IV}$ é dada por

$$\mathbf{b}_{IV} \approx \mathcal{N}(\beta, \sigma^2 \mathbf{Q}_{zx}^{-1} \mathbf{Q}_{zz} \mathbf{Q}_{xz}^{-1})$$

sendo $\mathbf{Q}_{zx} = \text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{X}}{n} \right)$ e $\mathbf{Q}_{zz} = \text{plim} \left(\frac{\mathbf{Z}^\top \mathbf{Z}}{n} \right)$.

Estimamos σ^2 por

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \mathbf{x}_i^\top \mathbf{b}_{IV})^2$$

e temos

$$\begin{aligned} \widehat{\text{Var}}(\mathbf{b}_{IV}) &= \frac{1}{n} \left[\left(\frac{\hat{\boldsymbol{\varepsilon}}^\top \hat{\boldsymbol{\varepsilon}}}{n} \right) \left(\frac{\mathbf{Z}^\top \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{Z}^\top \mathbf{Z}}{n} \right) \left(\frac{\mathbf{X}^\top \mathbf{Z}}{n} \right)^{-1} \right] \\ &= \hat{\sigma}^2 (\mathbf{Z}^\top \mathbf{X})^{-1} (\mathbf{Z}^\top \mathbf{Z}) (\mathbf{X}^\top \mathbf{Z})^{-1}, \end{aligned}$$

sendo $\hat{\boldsymbol{\varepsilon}} = [\mathbf{I} - \mathbf{X}(\mathbf{Z}^\top \mathbf{X})^{-1} \mathbf{Z}^\top] \mathbf{y} = \mathbf{y} - \mathbf{X} \mathbf{b}_{IV}$.

Se $\mathbf{Z}$ contiver mais variáveis que $\mathbf{X}$, então podemos ter a projeção das colunas de $\mathbf{X}$ no espaço das colunas de $\mathbf{Z}$:

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

Neste caso, temos as variáveis instrumentais $\hat{\mathbf{X}}$ e

$$\mathbf{b}_{IV} = (\hat{\mathbf{X}}^\top \mathbf{X})^{-1} \hat{\mathbf{X}}^\top \mathbf{y} = [\mathbf{X}^\top \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{X}]^{-1} \mathbf{X}^\top \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{y}.$$

Temos também

$$\widehat{\text{Var}}(\mathbf{b}_{IV}) = \hat{\sigma}^2 [\mathbf{X}^\top \mathbf{Z}(\mathbf{Z}^\top \mathbf{Z})^{-1} \mathbf{Z}^\top \mathbf{X}]^{-1}.$$

Podemos reescrever

$$\mathbf{b}_{IV} = (\hat{\mathbf{X}}^\top \mathbf{X})^{-1} \hat{\mathbf{X}}^\top \mathbf{y} = (\mathbf{X}^\top (\mathbf{I} - \mathbf{M}_z) \mathbf{X})^{-1} \mathbf{X}^\top (\mathbf{I} - \mathbf{M}_z) \mathbf{y} = (\hat{\mathbf{X}}^\top \hat{\mathbf{X}})^{-1} \hat{\mathbf{X}}^\top \mathbf{y}.$$

Isto é chamado de estimador de mínimos quadrados em dois passos (primeiro $\hat{\mathbf{X}}$, depois $\mathbf{b}_{IV}$).

O estimador de σ^2 não deve ser baseado em $\hat{\mathbf{X}}$.

O estimador

$$s_{IV}^2 = \frac{(\mathbf{y} - \hat{\mathbf{X}}\mathbf{b}_{IV})^\top (\mathbf{y} - \hat{\mathbf{X}}\mathbf{b}_{IV})}{n}$$

é inconsistente para σ^2 , com ou sem a correção para os graus de liberdade.

Exemplo

No modelo de equilíbrio de mercado, podemos ter a equação

$$\text{Wks}_{it} = \gamma_1 + \gamma_2 \ln \text{Wage}_{it} + \gamma_3 \text{Ed}_i + \gamma_4 \text{Union}_{it} + \gamma_5 \text{Fem}_i + u_{it}.$$

Se o número de semanas trabalhadas e o salário aceito são determinados conjuntamente, então $\ln \text{Wage}$ e u_{it} nesta equação são correlacionados.

Consideramos dois conjuntos de variáveis instrumentais:

$$\mathbf{Z}_1 = [1, \text{Ind}_{it}, \text{Ed}_i, \text{Union}_{it}, \text{Fem}_i]$$

e

$$\mathbf{Z}_2 = [1, \text{Ind}_{it}, \text{Ed}_i, \text{Union}_{it}, \text{Fem}_i, \text{SMSA}_{it}].$$

Tabela 18: Equação do trabalho estimada.

Variável	OLS		IV com Z_1		IV com Z_2	
	Estimado	E.P.	Estimado	E.P.	Estimado	E.P.
Constant	44,7665	1,2153	18,8986	13,0590	30,7044	4,9997
$\ln Wage$	0,7326	0,1972	5,1828	2,2454	3,1518	0,8572
Education	-0,1532	0,0321	-0,4600	0,1578	-0,3200	0,0661
Union	-1,9960	0,1701	-2,3602	0,2567	-2,1940	0,1860
Female	-1,3498	0,2642	0,6957	1,0650	-0,2378	0,4679

- Notemos o valor baixo para o coeficiente do $\ln Wage$ no OLS.
- As duas estimativas por variáveis instrumentais são maiores.
- Para as variáveis Z_2 , utilizamos a estimação em dois passos discutidos anteriormente.

Estimação por máxima verossilhança

A função de verossimilhança é dada por

$$L(\boldsymbol{\theta}|\mathbf{y}) = f(y_1, \dots, y_n|\boldsymbol{\theta}) = \prod_{i=1}^n f(y_i|\boldsymbol{\theta})$$

para $y_i, i = 1, \dots, n$, iid condicional em $\boldsymbol{\theta}$.

O logaritmo da função de verossimilhança é dada por

$$\ell(\boldsymbol{\theta}|\mathbf{y}) = \ln L(\boldsymbol{\theta}|\mathbf{y}) = \sum_{i=1}^n \ln f(y_i|\boldsymbol{\theta}).$$

Suponha que no modelo clássico de regressão os distúrbios sejam normais. Então,

$$\ell(\boldsymbol{\theta}|\mathbf{y}, \mathbf{X}) = \sum_{i=1}^n \ln f(y_i|\mathbf{x}_i, \boldsymbol{\theta}) = -\frac{1}{2} \sum_{i=1}^n \left[\ln \sigma^2 + \ln(2\pi) + \frac{1}{\sigma^2} (y_i - \mathbf{x}_i^\top \boldsymbol{\beta})^2 \right]$$

sendo $\mathbf{X}$ a matriz $n \times p$ dos regressores com a i -ésima linha igual a $\mathbf{x}_i$.

Estamos interessados em obter estimativas dos parâmetros, θ .

Definição

O vetor de parâmetros θ é **identificável** (estimável) se para qualquer outro vetor de parâmetros, $\gamma \neq \theta$, para algum dado $\mathbf{y}$, $L(\gamma|\mathbf{y}) \neq L(\theta|\mathbf{y})$.

O princípio da máxima verossimilhança

Suponha uma amostra de 10 observações de uma distribuição Poisson: 5, 0, 1, 1, 0, 3, 2, 3, 4 e 1.

$$L(\theta|\mathbf{y}) = \frac{e^{-100\theta} \theta^{20}}{207.360}.$$

Qual valor mais provável de θ que produziu esta amostra?

A função tem um único máximo em $\theta = 2$, que é a estimativa de máxima verossimilhança.

$$\begin{aligned}
\ell(\theta|\mathbf{y}) &= -n\theta + n\bar{y} \ln \theta - \sum_{i=1}^n \ln(y_i!) \\
\frac{\partial \ell(\theta|\mathbf{y})}{\partial \theta} &= -n + \frac{n\bar{y}}{\theta} = 0, \quad \Rightarrow \quad \hat{\theta}_{MV} = \bar{y}_n \\
\frac{\partial^2 \ell(\theta|\mathbf{y})}{\partial \theta^2} &= -\frac{n\bar{y}}{\theta^2} < 0, \quad \Rightarrow \quad \bar{y}_n \text{ é ponto de máximo.}
\end{aligned}$$

Chamamos de **equação de verossimilhança** a expressão:

$$\frac{\partial \ell(\theta|\mathbf{dados})}{\partial \theta} = \frac{\partial \ln L(\theta|\mathbf{dados})}{\partial \theta} = \mathbf{0}.$$

Denotaremos:

- $\hat{\theta}$ como o estimador de máxima verossimilhança;
- θ_0 o valor verdadeiro do vetor de parâmetros; e
- θ um outro valor possível do vetor parâmetros.

Teorema

Sob regularidade, o EMV tem as seguintes propriedades assintóticas:

M1: Consistência: $\text{plim } \hat{\theta} = \theta_0$.

M2: Normalidade: $\hat{\theta} \approx \mathcal{N}(\theta_0, [I(\theta_0)]^{-1})$ sendo

$$I(\theta_0) = E_0 \left(-\frac{\partial^2 \ell(\theta|\mathbf{y})}{\partial \theta \partial \theta^\top} \right).$$

M3: Eficiência: $\hat{\theta}$ é assintoticamente eficiente e atinge o limite inferior de Cramér-Rao.

M4: Invariância: o EMV de $\gamma = \mathbf{c}(\theta_0)$ é $\mathbf{c}(\hat{\theta})$ se $\mathbf{c}(\theta_0)$ é uma função contínua e continuamente diferenciável (transformações biunívocas).

Se $\mathbf{g}_i = \frac{\partial \ln f(y_i|\theta)}{\partial \theta}$ (a função escore) e $\mathbf{H}_i = \frac{\partial \ln f(y_i|\theta)}{\partial \theta \partial \theta^\top}$ (a matriz de informação), então

D1: $E_0(\mathbf{g}_i(\theta)) = \mathbf{0}$; e

D2: $\text{Var}_0(\mathbf{g}_i(\theta)) = E_0(-\mathbf{H}_i(\theta))$.

Teorema

$E_0 \left(\frac{1}{n} \ell(\theta_0) \right) > E_0 \left(\frac{1}{n} \ell(\theta) \right)$ para qualquer $\theta \neq \theta_0$ (incluindo $\hat{\theta}$).

O valor esperado da log-verossimilhança é maximizado no valor verdadeiro dos parâmetros.

Teorema (Cramér-Rao)

Sob condições de regularidade, a variância assintótica de um estimador, consistente e assintoticamente normal, do vetor de parâmetros θ_0 será sempre no mínimo maior ou igual a

$$[I(\theta_0)]^{-1} = \left[E_0 \left(-\frac{\partial^2 \ell(\theta_0)}{\partial \theta_0 \partial \theta_0^\top} \right) \right]^{-1} = \left[E_0 \left(\left(\frac{\partial \ell(\theta_0)}{\partial \theta_0} \right) \left(\frac{\partial \ell(\theta_0)}{\partial \theta_0} \right)^\top \right) \right]^{-1}.$$

Teste da razão de verossimilhança

Seja $\hat{\theta}_U$ o EMV sob o modelo irrestrito e $\hat{\theta}_R$ sob o modelo restrito.

Sejam também $\hat{L}_U$ e $\hat{L}_R$ as respectivas funções de verossimilhança avaliadas nestes pontos.

Então, a razão de verossimilhança é

$$\hat{\lambda} = \frac{\hat{L}_R}{\hat{L}_U}, \quad \text{com } 0 \leq \lambda \leq 1.$$

Teorema

Sob regularidade e sob H_0 , a distribuição para amostras grandes de $-2 \ln \hat{\lambda}$ é qui-quadrada, com os graus de liberdade igual ao número de restrições impostas.

A hipótese nula é rejeitada se este valor ultrapassa o valor crítico apropriado da distribuição qui-quadrada.

Teste de Wald

Se $\mathbf{x} \sim \mathcal{N}_J(\boldsymbol{\mu}, \boldsymbol{\Sigma})$, então $(\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu}) \sim \chi_J^2$.

Suponha que $\hat{\boldsymbol{\theta}}$ seja a estimativa do vetor de parâmetros sem restrições e que $H_0 : \mathbf{c}(\boldsymbol{\theta}) = \mathbf{q}$.

Se as restrições são válidas, então $\hat{\boldsymbol{\theta}}$ deve satisfazê-las. Caso contrário, $\mathbf{c}(\hat{\boldsymbol{\theta}}) - \mathbf{q}$ deve ser longe de $\mathbf{0}$.

Teorema

A estatística de Wald é

$$W = (\mathbf{c}(\hat{\boldsymbol{\theta}}) - \mathbf{q})^\top \left[\text{Var}(\mathbf{c}(\hat{\boldsymbol{\theta}}) - \mathbf{q}) \right]^{-1} (\mathbf{c}(\hat{\boldsymbol{\theta}}) - \mathbf{q}).$$

Sob H_0 , em amostras grandes, W tem uma distribuição qui-quadrada com os graus de liberdade igual ao número de restrições, isto é, o número de equações em $\mathbf{c}(\hat{\boldsymbol{\theta}}) - \mathbf{q} = \mathbf{0}$.

Ademais, $\widehat{\text{Var}}(\mathbf{c}(\hat{\boldsymbol{\theta}}) - \mathbf{q}) = \widehat{\mathbf{C}} \widehat{\text{Var}}(\hat{\boldsymbol{\theta}}) \widehat{\mathbf{C}}^\top$ sendo $\widehat{\mathbf{C}} = \left[\frac{\partial \mathbf{c}(\hat{\boldsymbol{\theta}})}{\partial \hat{\boldsymbol{\theta}}} \right]$.

Teste do Multiplicador de Lagrange

Suponha que maximizamos a log-verossimilhança sob a restrição $\mathbf{c}(\boldsymbol{\theta}) - \mathbf{q} = \mathbf{0}$. Seja $\boldsymbol{\lambda}$ o vetor de multiplicadores de Lagrange e definamos a função lagrangeana

$$\ln L^*(\boldsymbol{\theta}) = \ln L(\boldsymbol{\theta}) + \boldsymbol{\lambda}^\top (\mathbf{c}(\boldsymbol{\theta}) - \mathbf{q}).$$

A solução é dada por

$$\begin{aligned}\frac{\partial \ln L^*}{\partial \boldsymbol{\theta}} &= \frac{\partial \ln L(\boldsymbol{\theta})}{\partial \boldsymbol{\theta}} + \mathbf{c}^\top \boldsymbol{\lambda} \\ \frac{\partial \ln L^*}{\partial \boldsymbol{\lambda}} &= \mathbf{c}(\boldsymbol{\theta}) - \mathbf{q} = \mathbf{0}.\end{aligned}$$

Podemos testar $H_0 : \boldsymbol{\lambda} = \mathbf{0}$.

Teorema

A estatística de teste do multiplicador de Lagrange é

$$LM = \left(\frac{\partial \ell(\hat{\boldsymbol{\theta}}_R)}{\partial \hat{\boldsymbol{\theta}}_R} \right)^\top \left[\mathbf{I}(\hat{\boldsymbol{\theta}}_R) \right]^{-1} \left(\frac{\partial \ell(\hat{\boldsymbol{\theta}}_R)}{\partial \hat{\boldsymbol{\theta}}_R} \right).$$

Sob H_0 , LM tem distribuição limite como qui-quadrada com graus de liberdade igual ao número de restrições. Todos os termos são calculados no estimador restrito.

Comparando modelos

Duas medidas comuns para modelos não encaixados baseados na mesma lógica são

- Critério de Informação de Akaike: $AIC = -2 \ln L(\hat{\boldsymbol{\theta}}) + 2p$; e
- Critério de Informação Bayesiana/Schwarz: $BIC = -2 \ln L(\hat{\boldsymbol{\theta}}) + p \ln n$,

sendo p o número de parâmetros do modelo. Menor AIC (ou BIC), melhor o modelo.

Modelo de regressão linear normal

$$y_i = \mathbf{x}_i^\top \boldsymbol{\beta} + \varepsilon_i.$$

$$\begin{aligned} L(\boldsymbol{\beta}, \sigma^2) &= (2\pi\sigma^2)^{-n/2} \exp \left\{ -\frac{1}{2\sigma^2} \boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon} \right\} \\ &= (2\pi\sigma^2)^{-n/2} \exp \left\{ -\frac{1}{2\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \right\}. \end{aligned}$$

$$\ell(\boldsymbol{\beta}, \sigma^2) = -\frac{n}{2} \ln(2\pi) - \frac{n}{2} \ln \sigma^2 - \frac{(\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})}{2\sigma^2}.$$

$$\begin{bmatrix} \frac{\partial \ell(\boldsymbol{\beta}, \sigma^2)}{\partial \boldsymbol{\beta}} \\ \frac{\partial \ell(\boldsymbol{\beta}, \sigma^2)}{\partial \sigma^2} \end{bmatrix} = \begin{bmatrix} \frac{\mathbf{X}^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})}{\sigma^2} \\ -\frac{n}{2\sigma^2} + \frac{(\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})}{2\sigma^4} \end{bmatrix} = \begin{bmatrix} \mathbf{0} \\ 0 \end{bmatrix} \Rightarrow$$

$$\hat{\boldsymbol{\beta}}_{MV} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} \quad \text{e} \quad \hat{\sigma}_{MV}^2 = \frac{\mathbf{e}^\top \mathbf{e}}{n}.$$

Temos,

$$\begin{bmatrix} \frac{\partial^2 \ell(\boldsymbol{\beta}, \sigma^2)}{\partial \boldsymbol{\beta} \partial \boldsymbol{\beta}^\top} & \frac{\partial^2 \ell(\boldsymbol{\beta}, \sigma^2)}{\partial \boldsymbol{\beta} \partial \sigma^2} \\ \frac{\partial^2 \ell(\boldsymbol{\beta}, \sigma^2)}{\partial \sigma^2 \partial \boldsymbol{\beta}} & \frac{\partial^2 \ell(\boldsymbol{\beta}, \sigma^2)}{\partial (\sigma^2)^2} \end{bmatrix} = \begin{bmatrix} -\frac{\mathbf{X}^\top \mathbf{X}}{\sigma^2} & -\frac{\mathbf{X}^\top \boldsymbol{\varepsilon}}{\sigma^4} \\ -\frac{\boldsymbol{\varepsilon}^\top \mathbf{X}}{\sigma^4} & \frac{n}{2\sigma^4} - \frac{\boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon}}{\sigma^6} \end{bmatrix} \Rightarrow$$

$$\left[I(\boldsymbol{\beta}, \sigma^2) \right]^{-1} = \begin{bmatrix} \sigma^2 (\mathbf{X}^\top \mathbf{X})^{-1} & \mathbf{0} \\ \mathbf{0} & \frac{2\sigma^4}{n} \end{bmatrix}.$$

Temos, $\sqrt{n}(\hat{\sigma}_{MV}^2 - \sigma^2) \xrightarrow{d} \mathcal{N}(0, 2\sigma^4)$.

Se $H_0 : \mathbf{R}\boldsymbol{\beta} - \mathbf{q} = \mathbf{0}$, o teste da razão F ,

$$F = \frac{(\mathbf{R}\mathbf{b} - \mathbf{q})^\top [\mathbf{R}\mathbf{s}^2(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{R}^\top](\mathbf{R}\mathbf{b} - \mathbf{q})}{J} \sim \mathcal{F}_{J, n-p},$$

para qualquer tamanho de amostra se os distúrbios forem normalmente distribuídos.

Os outros testes e suas estatísticas de teste continuam não tendo distribuição exata conhecida para amostras finitas ou pequenas.

Modelo de regressão generalizado

$$y_i = \mathbf{x}_i^\top \boldsymbol{\beta} + \varepsilon_i, \quad i = 1, 2, \dots, n$$

$$E(\varepsilon|\mathbf{X}) = \mathbf{0}$$

$$E(\varepsilon\varepsilon^\top|\mathbf{X}) = \sigma^2\boldsymbol{\Omega}.$$

Por hipótese, teremos $\boldsymbol{\Omega}$ como uma matriz constante conhecida.

$$\ell(\boldsymbol{\beta}, \sigma^2) = -\frac{n}{2} \ln(2\pi) - \frac{n}{2} \ln \sigma^2 - \frac{1}{2} \ln |\boldsymbol{\Omega}| - \frac{1}{2\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top \boldsymbol{\Omega}^{-1} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})$$

$$\frac{\partial \ell}{\partial \boldsymbol{\beta}} = \frac{1}{\sigma^2} \mathbf{X}^\top \boldsymbol{\Omega}^{-1} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) = \frac{1}{\sigma^2} \mathbf{X}_*^\top (\mathbf{y}_* - \mathbf{X}_*\boldsymbol{\beta}) = \mathbf{0} \quad \boldsymbol{\Omega}^{-1} = \mathbf{P}^\top \mathbf{P}$$

$$\frac{\partial \ell}{\partial \sigma^2} = -\frac{n}{2\sigma^2} + \frac{1}{\sigma^4} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top \boldsymbol{\Omega}^{-1} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \quad \mathbf{X}_* = \mathbf{P}\mathbf{X}$$

$$= -\frac{n}{2\sigma^2} + \frac{1}{\sigma^4} (\mathbf{y}_* - \mathbf{X}_*\boldsymbol{\beta})^\top (\mathbf{y}_* - \mathbf{X}_*\boldsymbol{\beta}) = 0. \quad \mathbf{y}_* = \mathbf{P}\mathbf{y}$$

$$\hat{\boldsymbol{\beta}}_{MV} = (\mathbf{X}_*^\top \mathbf{X}_*)^{-1} \mathbf{X}_*^\top \mathbf{y} = (\mathbf{X}^\top \boldsymbol{\Omega}^{-1} \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\Omega}^{-1} \mathbf{y} \quad \text{e}$$

$$\begin{aligned} \hat{\sigma}_{MV}^2 &= \frac{1}{n} (\mathbf{y}_* - \mathbf{X}_* \hat{\boldsymbol{\beta}}_{MV})^\top (\mathbf{y}_* - \mathbf{X}_* \hat{\boldsymbol{\beta}}_{MV}) \\ &= \frac{1}{n} (\mathbf{y} - \mathbf{X} \hat{\boldsymbol{\beta}}_{MV})^\top \boldsymbol{\Omega}^{-1} (\mathbf{y} - \mathbf{X} \hat{\boldsymbol{\beta}}_{MV}). \end{aligned}$$

Temos que $\hat{\sigma}_{MV}^2$ é tendencioso para σ^2 . Para obter um estimador não tendencioso precisamos multiplicá-lo pelo fator $n/(n-p)$.

Se Ω for desconhecido, então precisamos estimar $(\beta, \sigma^2, \Omega)$ simultaneamente.

Mas Ω tem $n(n+1)/2 - 1$ parâmetros. Precisamos impor restrições.

Em geral, a estimação conjunta de todos os parâmetros será complicada.

Modelo de heterocedasticidade multiplicativa

Considere um modelo de regressão com variância dada por

$$\sigma_i^2 = \sigma^2 \exp\{\mathbf{w}_i^\top \boldsymbol{\alpha}\} = \exp\{\mathbf{z}_i^\top \boldsymbol{\gamma}\},$$

sendo $\mathbf{z}_i^\top = (1, \mathbf{w}_i^\top)$ e $\boldsymbol{\gamma}^\top = (\ln \sigma^2, \boldsymbol{\alpha}^\top)$.

Neste caso, tomamos $\Omega = \text{diag}(\exp\{\mathbf{z}_1^\top \boldsymbol{\gamma}\}, \exp\{\mathbf{z}_2^\top \boldsymbol{\gamma}\}, \dots, \exp\{\mathbf{z}_n^\top \boldsymbol{\gamma}\})$.

$$\begin{aligned}
\ell(\boldsymbol{\beta}, \boldsymbol{\gamma}) &= -\frac{n}{2} \ln(2\pi) - \frac{1}{2} \sum_{i=1}^n \ln \sigma_i^2 - \frac{1}{2} \sum_{i=1}^n \frac{\varepsilon_i^2}{\sigma_i^2} \\
&= -\frac{n}{2} \ln(2\pi) - \frac{1}{2} \sum_{i=1}^n \ln \mathbf{z}_i^\top \boldsymbol{\gamma} - \frac{1}{2} \sum_{i=1}^n \frac{\varepsilon_i^2}{\exp\{\mathbf{z}_i^\top \boldsymbol{\gamma}\}}
\end{aligned}$$

$$\frac{\partial \ell(\boldsymbol{\beta}, \boldsymbol{\gamma})}{\partial \boldsymbol{\beta}} = \sum_{i=1}^n \mathbf{x}_i \frac{\varepsilon_i}{\exp\{\mathbf{z}_i^\top \boldsymbol{\gamma}\}} = \mathbf{X}^\top \boldsymbol{\Omega}^{-1} \boldsymbol{\varepsilon} = \mathbf{0}$$

$$\frac{\partial \ell(\boldsymbol{\beta}, \boldsymbol{\gamma})}{\partial \boldsymbol{\gamma}} = \frac{1}{2} \sum_{i=1}^n \mathbf{z}_i \left(\frac{\varepsilon_i^2}{\exp\{\mathbf{z}_i^\top \boldsymbol{\gamma}\}} - 1 \right) = \mathbf{0}.$$

$$\frac{\partial^2 \ell(\boldsymbol{\beta}, \boldsymbol{\gamma})}{\partial \boldsymbol{\beta} \partial \boldsymbol{\beta}^\top} = -\sum_{i=1}^n \frac{1}{\exp\{\mathbf{z}_i^\top \boldsymbol{\gamma}\}} \mathbf{x}_i \mathbf{x}_i^\top = -\mathbf{X}^\top \boldsymbol{\Omega}^{-1} \mathbf{X}$$

$$\frac{\partial^2 \ell(\boldsymbol{\beta}, \boldsymbol{\gamma})}{\partial \boldsymbol{\beta} \partial \boldsymbol{\gamma}^\top} = -\sum_{i=1}^n \frac{\varepsilon_i}{\exp\{\mathbf{z}_i^\top \boldsymbol{\gamma}\}} \mathbf{x}_i \mathbf{z}_i^\top$$

$$\frac{\partial^2 \ell(\boldsymbol{\beta}, \boldsymbol{\gamma})}{\partial \boldsymbol{\gamma} \partial \boldsymbol{\gamma}^\top} = -\frac{1}{2} \sum_{i=1}^n \frac{\varepsilon_i^2}{\exp\{\mathbf{z}_i^\top \boldsymbol{\gamma}\}} \mathbf{z}_i \mathbf{z}_i^\top.$$

Temos,

$$\begin{aligned} \mathbb{E} \left(-\frac{\partial^2 \ell(\beta, \gamma)}{\partial \beta \partial \gamma^\top} \right) &= \mathbf{0} \quad \text{porque} \quad \mathbb{E}(\varepsilon_i | \mathbf{x}_i, \mathbf{z}_i) = 0, \quad \text{e} \\ \mathbb{E} \left(-\frac{\partial^2 \ell(\beta, \gamma)}{\partial \gamma \partial \gamma^\top} \right) &= \frac{1}{2} \mathbf{Z}^\top \mathbf{Z} \quad \text{porque} \quad \mathbb{E} \left(\frac{\varepsilon_i^2}{\sigma_i^2} \middle| \mathbf{x}_i, \mathbf{z}_i \right) = 1. \end{aligned}$$

Seja $\delta = (\beta, \gamma)$. Então,

$$\mathbb{E} \left(-\frac{\partial^2 \ell(\delta)}{\partial \delta \partial \delta^\top} \right) = \begin{bmatrix} \mathbf{X}^\top \boldsymbol{\Omega}^{-1} \mathbf{X} & \mathbf{0} \\ \mathbf{0} & \frac{1}{2} \mathbf{Z}^\top \mathbf{Z} \end{bmatrix} = -\bar{\mathbf{H}}.$$

O método do escore é um algoritmo para encontrar uma solução:

$$\delta_{t+1} = \delta_t - \bar{\mathbf{H}}^{-1} \mathbf{g}_t$$

$$\text{sendo } \mathbf{g} = \left(-\frac{\partial \ell(\delta)}{\partial \beta^\top}, -\frac{\partial \ell(\delta)}{\partial \gamma^\top} \right)^\top.$$

Como $\bar{\mathbf{H}}$ é bloco diagonal, temos

$$\begin{aligned}\beta_{t+1} &= \beta_t + (\mathbf{X}^\top \Omega_t^{-1} \mathbf{X}) \mathbf{X}^\top \Omega_t^{-1} \varepsilon_t \\ &= \beta_t + (\mathbf{X}^\top \Omega_t^{-1} \mathbf{X}) \mathbf{X}^\top \Omega_t^{-1} (\mathbf{y} - \mathbf{X} \beta_t) \\ &= (\mathbf{X}^\top \Omega_t^{-1} \mathbf{X}) \mathbf{X}^\top \Omega_t^{-1} \mathbf{y} \\ \gamma_{t+1} &= \gamma_t + \left[2(\mathbf{Z}^\top \mathbf{Z})^{-1} \right] \left[\frac{1}{2} \sum_{i=1}^n \mathbf{z}_i \left(\frac{\varepsilon_{i,t}}{\exp\{\mathbf{z}_i^\top \gamma_t\}} - 1 \right) \right].\end{aligned}$$

Esboço do algoritmo:

1. Estime a variância dos distúrbios σ_i^2 com $\exp\{\mathbf{z}_i^\top \gamma\}$.
2. Calcule β_{t+1} .
3. Calcule γ_{t+1} .
4. Calcule $\mathbf{d}_{t+1} = \delta_{t+1} - \delta_t$. Se $\mathbf{d}_{t+1}$ for grande, retorne ao Passo 1.

Temos também que

$$\begin{aligned}\text{Var}(\hat{\beta}_{MV}) &= (\mathbf{X}^\top \Omega^{-1} \mathbf{X})^{-1} \\ \text{Var}(\hat{\gamma}_{MV}) &= 2(\mathbf{Z}^\top \mathbf{Z})^{-1}.\end{aligned}$$

Testar a hipótese de homocedasticidade, $H_0 : \alpha = \mathbf{0}$, usando o teste de Wald,

$$W = \begin{pmatrix} 0 & \hat{\alpha}^\top \end{pmatrix} \left[\frac{1}{2} \mathbf{Z}^\top \mathbf{Z} \right]^{-1} \begin{pmatrix} 0 \\ \hat{\alpha}^\top \end{pmatrix} = 2\hat{\alpha}^\top (\mathbf{Z}_1^\top \tilde{M} \mathbf{Z}_1) \hat{\alpha} \sim \chi_J^2.$$

Exemplo

$$\ln C_{i,t} = \beta_1 + \beta_2 \ln Q_{i,t} + \beta_3 [\ln Q_{i,t}]^2 + \beta_4 \ln Pf_{i,t} + \beta_5 \text{Loadfactor}_{i,t} + \varepsilon_{i,t},$$

sendo

- $C_{i,t}$ o custo total;
- $Q_{i,t}$ a produção; e
- $Pf_{i,t}$ o preço do combustível.

Os dados são compostos de 90 observações para seis firmas observadas em 15 anos.

Para completar o modelo, temos

$$\sigma_{i,t}^2 = \sigma^2 \exp\{\alpha \text{Loadfactor}_{i,t}\} = \exp\{\gamma_1 + \gamma_2 \text{Loadfactor}_{i,t}\}.$$

Tabela 19: Modelo de heterocedasticidade multiplicativa

β_1	β_2	β_3	β_4	γ_2
9,2774	0,91609	0,021643	0,40174	9,78076
(0,20977)	(0,032993)	(0,011017)	(0,016332)	(2,839)

Para testa $H_0 : \alpha = 0$, temos $W = (9,78076/2,839)^2 = 11,869 > 3,84$. Logo, rejeitamos a hipótese nula.

Modelo de fronteira estocástica

$y = f(\mathbf{x})$ produção dado os insumos $\mathbf{x}$.

Temos que $y \leq f(\mathbf{x}) \Rightarrow y = h(\mathbf{x}, \beta) + u$, u deve ser negativo.

Qualquer distúrbio diferente de zero deve ser interpretado como ineficiência.

Da função Cobb-Douglas, temos

$$\ln y = \beta_1 + \sum_k \beta_k \ln x_k - u, \quad u \geq 0.$$

Fronteira estocástica:

$$\begin{aligned} \ln y &= \beta_1 + \sum_k \beta_k \ln x_k - u + \nu, \quad u \geq 0, \quad \nu \sim \mathcal{N}(0, \sigma_\nu^2) \\ &= \beta_1 + \sum_k \beta_k \ln x_k + \varepsilon. \end{aligned}$$

A fronteira de qualquer firma particular é $h(\mathbf{x}, \beta) + \nu$.

Os dados estão em logaritmos, logo u é a medida da porcentagem na qual a observação particular falha em atingir a fronteira (razão de produção ideal).

Sejam $\varepsilon = \nu - u$, $\lambda = \sigma_u/\sigma_\nu$, $\sigma = (\sigma_u^2 + \sigma_\nu^2)^{1/2}$ e $\Phi(z)$ a f.d.a.

Para o modelo normal truncada,

$$\ln h(\varepsilon_i|\beta, \lambda, \sigma) = \left[-\ln \sigma + \frac{1}{2} \ln \frac{2}{\pi} - \frac{1}{2} \left(\frac{\varepsilon_i}{\sigma} \right)^2 + \ln \Phi \left(\frac{-\varepsilon_i \lambda}{\sigma} \right) \right],$$

$$E(\varepsilon) = -\sigma_u \left(\frac{2}{\pi} \right)^{1/2},$$

e para o modelo exponencial

$$\ln h(\varepsilon_i|\beta, \lambda, \sigma) = \left[\ln \theta + \frac{1}{2} \theta^2 \sigma_\nu^2 + \theta \varepsilon_i + \ln \Phi \left(\frac{-\varepsilon_i}{\sigma_\nu} - \theta \sigma \right) \right], \quad E(\varepsilon) = -\frac{1}{\theta}.$$

Fazemos $\beta_0 = \beta_1 + E(\varepsilon)$ e $\varepsilon^* = \varepsilon - E(\varepsilon)$, para obter

$$\ln y = \beta_0 + \sum_k \beta_k \ln x_k + \varepsilon^*.$$

Para o modelo normal truncada, temos

$$\ell(\boldsymbol{\beta}, \sigma, \lambda) = -n \ln \sigma + \frac{n}{2} \ln \frac{2}{\pi} - \frac{1}{2} \sum_{i=1}^n \left(\frac{\varepsilon_i}{\sigma} \right)^2 + \sum_{i=1}^n \ln \Phi \left(-\frac{\varepsilon_i \lambda}{\sigma} \right).$$

Seja $\theta = 1/\sigma$ e $\boldsymbol{\gamma} = (1/\sigma)\boldsymbol{\beta}$. Então,

$$\begin{aligned} \ell(\boldsymbol{\beta}, \sigma, \lambda) &= n \ln \theta + \frac{n}{2} \ln \frac{2}{\pi} - \frac{1}{2} \sum_{i=1}^n \left(\theta y_i - \mathbf{x}_i^\top \boldsymbol{\gamma} \right)^2 \\ &\quad + \sum_{i=1}^n \ln \Phi \left(-\lambda(\theta y_i - \mathbf{x}_i^\top \boldsymbol{\gamma}) \right). \end{aligned}$$

Definimos:

- $\alpha_i = \varepsilon_i / \sigma = \theta y_i - \mathbf{x}_i^\top \boldsymbol{\gamma}$
- $\delta(y_i, \mathbf{x}_i, \lambda, \theta, \boldsymbol{\gamma}) = \frac{\phi(-\lambda \alpha_i)}{\Phi(-\lambda \alpha_i)} = \delta_i$
- $\mathbf{z}_i = (\mathbf{x}_i^\top, -y_i)^\top$, com dimensão $(p+1) \times 1$
- $\mathbf{t}_i = (\mathbf{0}^\top, 1/\theta)^\top$, com dimensão $(p+1) \times 1$.

As equações de verossimilhança são

$$\begin{aligned}\frac{\partial \ell(\boldsymbol{\gamma}, \theta, \lambda)}{\partial (\boldsymbol{\gamma}^\top, \theta)^\top} &= \sum_{i=1}^n \mathbf{t}_i + \sum_{i=1}^n \alpha_i \mathbf{z}_i + \lambda \sum_{i=1}^n \delta_i \mathbf{z}_i = \mathbf{0} \\ \frac{\partial \ell(\boldsymbol{\gamma}, \theta, \lambda)}{\partial \lambda} &= - \sum_{i=1}^n \delta_i \alpha_i = 0.\end{aligned}$$

Temos também

$$\mathbf{H}(\boldsymbol{\gamma}, \sigma, \lambda) = \sum_{i=1}^n \left\{ \begin{bmatrix} (\lambda^2 \Delta_i - 1) \mathbf{z}_i \mathbf{z}_i^\top & (\delta_i - \lambda \alpha_i \Delta_i) \mathbf{z}_i \\ (\delta_i - \lambda \alpha_i \Delta_i) \mathbf{z}_i^\top & \alpha_i^2 \Delta_i \end{bmatrix} - \begin{bmatrix} \mathbf{t}_i \mathbf{t}_i^\top & \mathbf{0} \\ \mathbf{0}^\top & 0 \end{bmatrix} \right\}.$$

Segue-se que $\widehat{\text{Var}} \left((\hat{\boldsymbol{\gamma}}^\top, \hat{\theta}, \hat{\lambda})^\top \right) = \left[\mathbf{H}(\hat{\boldsymbol{\gamma}}^\top, \hat{\theta}, \hat{\lambda}) \right]^{-1},$

$$\hat{\sigma}_{MV} = 1/\hat{\theta}_{MV} \text{ e } \hat{\boldsymbol{\beta}}_{MV} = \frac{\hat{\boldsymbol{\gamma}}_{MV}}{\hat{\theta}_{MV}}.$$

Recorremos ao método delta para estimar as variâncias.

$$\mathbf{G} = \begin{bmatrix} \partial \hat{\beta} / \partial \hat{\gamma} & \partial \hat{\beta} / \partial \hat{\theta} & \partial \hat{\beta} / \partial \hat{\lambda} \\ \partial \hat{\sigma} / \partial \hat{\gamma} & \partial \hat{\sigma} / \partial \hat{\theta} & \partial \hat{\sigma} / \partial \hat{\lambda} \\ \partial \hat{\lambda} / \partial \hat{\gamma} & \partial \hat{\lambda} / \partial \hat{\theta} & \partial \hat{\lambda} / \partial \hat{\lambda} \end{bmatrix} = \begin{bmatrix} -(1/\hat{\theta})\mathbf{I} & -(1/\hat{\theta}^2)\hat{\gamma} & \mathbf{0} \\ \mathbf{0}^\top & -(1/\hat{\theta}^2) & 0 \\ \mathbf{0}^\top & 0 & 1 \end{bmatrix}.$$

$\widehat{\text{Var}} \left((\hat{\beta}, \hat{\sigma}, \hat{\lambda})^\top \right) = \mathbf{G} \left[\mathbf{H}(\hat{\gamma}^\top, \hat{\theta}, \hat{\lambda}) \right]^{-1} \mathbf{G}^\top$. Também podemos estimar

$$\hat{\sigma}_v^2 = \frac{\hat{\sigma}^2}{1 + \hat{\lambda}^2} \text{ e } \hat{\sigma}_u^2 = \frac{\hat{\sigma}^2 \hat{\lambda}^2}{1 + \hat{\lambda}^2}.$$

Em geral, estamos interessados nas eficiências u_i . Mas obtemos somente $\varepsilon = \mathbf{y} - \mathbf{x}^\top \beta$.

Podemos, entretanto, utilizar a seguinte aproximação:

$$\mathbf{E}(u|\varepsilon) = \frac{\sigma \lambda}{1 + \lambda^2} \left[\frac{\phi(z)}{1 - \Phi(z)} - z \right], \quad z = \frac{\varepsilon \lambda}{\sigma},$$

para o modelo normal truncada e

$$E(u|\varepsilon) = z + \sigma_\nu \frac{\phi(z/\sigma_\nu)}{\Phi(z/\sigma_\nu)}, \quad z = -\varepsilon - \theta\sigma_\nu^2,$$

para o modelo exponencial.

Podemos estimar estas quantidades com as estimativas de máxima verossimilhança.

Estamos interessados também na variância relativa. Para o modelo normal truncada,

$$\text{Var}(\varepsilon) = \text{Var}(u) + \text{Var}(\nu) = (1 - 2\pi)\sigma_u^2 + \sigma_\nu^2$$

e para o modelo exponencial

$$\text{Var}(\varepsilon) = \frac{1}{\theta^2} + \sigma_\nu^2.$$

Mais uma vez, podemos utilizar as estimativas de máxima verossimilhança para estimar estas quantidades.

O melhor neste caso seria aplicar a estimação bayesiana dos parâmetros do modelo, inclusive a estimação de todas eficiências u_i .

Estimação e inferência bayesiana

$$\begin{aligned} p(\theta|\mathbf{y}) &= \frac{p(\mathbf{y}|\theta)p(\theta)}{p(\mathbf{y})} = \frac{L(\theta|\mathbf{y})p(\theta)}{p(\mathbf{y})} \\ &\propto p(\mathbf{y}|\theta)p(\theta) = L(\theta|\mathbf{y})p(\theta), \end{aligned}$$

sendo

- $p(\theta|\mathbf{y})$ a distribuição a posteriori;
- $p(\theta)$ a distribuição a priori;
- $p(\mathbf{y}|\theta) = L(\theta|\mathbf{y})$ a função de verossimilhança; e
- $p(\mathbf{y})$ a distribuição marginal dos dados ou a verossimilhança marginal.

Temos,

$$p(\mathbf{y}) = \int p(\mathbf{y}|\theta)p(\theta)d\theta.$$

Além disso, $p(\mathbf{y}) < \infty$ para existência da distribuição a posteriori.

Análise bayesiana do modelo de regressão clássico

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}, \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}).$$

Logo,

$$L(\boldsymbol{\beta}, \sigma^2) = (2\pi\sigma^2)^{-n/2} \exp \left\{ -\frac{1}{2\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \right\}.$$

Seja $d = n - p$ os graus de liberdade e

$$\mathbf{y} - \mathbf{X}\boldsymbol{\beta} = \mathbf{y} - \mathbf{X}\mathbf{b} - \mathbf{X}(\boldsymbol{\beta} - \mathbf{b}) = \mathbf{e} - \mathbf{X}(\boldsymbol{\beta} - \mathbf{b}).$$

Então,

$$-\frac{1}{2\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) = \left(-\frac{ds^2}{2} \right) \sigma^{-2} + \frac{1}{2} (\boldsymbol{\beta} - \mathbf{b})^\top [\sigma^{-2} \mathbf{X}^\top \mathbf{X}] (\boldsymbol{\beta} - \mathbf{b}).$$

Logo,

$$\begin{aligned} L(\beta, \sigma^2 | \mathbf{y}, \mathbf{X}) &= (2\pi)^{-d/2} (\sigma^2)^{-d/2} \exp \left\{ -\frac{d\mathbf{s}^2}{2\sigma^2} \right\} \\ &\times (2\pi)^{-p/2} (\sigma^2)^{-p/2} \exp \left\{ -\frac{1}{2} (\beta - \mathbf{b})^\top [\sigma^2 (\mathbf{X}^\top \mathbf{X})^{-1}]^{-1} (\beta - \mathbf{b}) \right\} \\ &\propto \frac{(\nu \mathbf{s}^2)^{\nu-1}}{\Gamma(\nu-1)} (\sigma^2)^{-\nu} \exp \left\{ -\frac{\nu \mathbf{s}^2}{\sigma^2} \right\} \\ &\times (2\pi)^{-p/2} |\sigma^2 \mathbf{\Delta}|^{-1/2} \exp \left\{ -\frac{1}{2} (\beta - \mathbf{b})^\top [\sigma^2 \mathbf{\Delta}]^{-1} (\beta - \mathbf{b}) \right\}, \end{aligned}$$

sendo $n/2 = d/2 + p/2$, $\nu = d/2$ e $\mathbf{\Delta} = (\mathbf{X}^\top \mathbf{X})^{-1}$.

A verossimilhança é proporcional ao produto de uma inversa gama (σ^2) com parâmetros $\delta = \nu - 1$ e $\lambda = \nu \mathbf{s}^2$, e uma normal p dimensional ($\beta | \sigma^2$) com média $\mathbf{b}$ e matriz de covariância $\sigma^2 \mathbf{\Delta}$.

Podemos considerar uma priori não informativa,

$$p(\beta, \sigma^2) \propto \frac{1}{\sigma^2}.$$

$$\begin{aligned}
p(\boldsymbol{\beta}, \sigma^2 | \mathbf{y}, \mathbf{X}) &\propto L(\boldsymbol{\beta}, \sigma^2 | \mathbf{y}, \mathbf{X}) p(\boldsymbol{\beta}, \sigma^2) \\
&\propto \frac{(\nu \mathbf{s}^2)^\nu}{\Gamma(\nu)} (\sigma^2)^{-(\nu+1)} \exp \left\{ -\frac{\nu \mathbf{s}^2}{\sigma^2} \right\} \\
&\quad \times (2\pi)^{-p/2} |\sigma^2 \boldsymbol{\Delta}|^{-1/2} \exp \left\{ -\frac{1}{2} (\boldsymbol{\beta} - \mathbf{b})^\top [\sigma^2 \boldsymbol{\Delta}]^{-1} (\boldsymbol{\beta} - \mathbf{b}) \right\}.
\end{aligned}$$

Agora seja $A(\boldsymbol{\beta}) = (\boldsymbol{\beta} - \mathbf{b})^\top \boldsymbol{\Delta}^{-1} (\boldsymbol{\beta} - \mathbf{b})/2$, então

$$\begin{aligned}
p(\boldsymbol{\beta} | \mathbf{y}, \mathbf{X}) &= \int p(\boldsymbol{\beta}, \sigma^2 | \mathbf{y}, \mathbf{X}) d\sigma^2 \propto \int_0^\infty (\sigma^2)^{-(n/2+1)} \exp \left\{ -\frac{A(\boldsymbol{\beta}) + \nu \mathbf{s}^2}{\sigma^2} \right\} d\sigma^2 \\
&= \frac{\Gamma(n/2)}{[A(\boldsymbol{\beta}) + \nu \mathbf{s}^2]^{n/2}} \propto \Gamma(n/2) \left[1 + \frac{1}{2\nu \mathbf{s}^2} A(\boldsymbol{\beta}) \right]^{-((n-p)+p)/2} \\
&= \Gamma(n/2) \left[1 + \frac{1}{2\nu} (\boldsymbol{\beta} - \mathbf{b})^\top [\mathbf{s}^2 (\mathbf{X}^\top \mathbf{X})^{-1}]^{-1} (\boldsymbol{\beta} - \mathbf{b}) \right]^{-((n-p)+p)/2}.
\end{aligned}$$

Assim,

$$(\boldsymbol{\beta} | \mathbf{y}, \mathbf{X}) \sim t_{n-p} \left(\mathbf{b}, \mathbf{s}^2 (\mathbf{X}^\top \mathbf{X})^{-1} \right) \quad (p \text{ dimensional}).$$

É fácil mostrar que

$$(\sigma^2 | \mathbf{y}, \mathbf{X}) \sim \mathcal{IG}(\nu, \nu s^2).$$

Logo,

$$\begin{aligned} E(\beta | \mathbf{y}, \mathbf{X}) &= \mathbf{b}, \quad \text{para } n - p > 1, \\ E(\sigma^2 | \mathbf{y}, \mathbf{X}) &= \frac{\nu}{\nu - 1} s^2, \quad \text{para } \nu > 1 \ (n - p > 2). \end{aligned}$$

Podemos considerar também uma priori informativa. Por exemplo,

$$(\beta | \sigma^2) \sim \mathcal{N}_p(\mathbf{b}_0, \sigma^2 \mathbf{B}_0) \quad \text{e} \quad \sigma^2 \sim \mathcal{IG}(m, m\sigma_0^2) \quad (\text{priori normal gama}).$$

$$\begin{aligned} p(\beta, \sigma^2 | \mathbf{y}, \mathbf{X}) &\propto (\sigma^2)^{-(\nu+m+1)} \exp \left\{ -\frac{1}{\sigma^2} (m\sigma_0^2 + \nu s^2) \right\} \\ &\times |\boldsymbol{\Sigma}|^{-1/2} \exp \left\{ -\frac{1}{2} (\beta - \mathbf{a})^\top \boldsymbol{\Sigma}^{-1} (\beta - \mathbf{a}) \right\}, \end{aligned}$$

sendo

$$\begin{aligned} \boldsymbol{\Sigma} &= \left[\sigma^{-2} \mathbf{B}_0^{-1} + \sigma^{-2} \mathbf{X}^\top \mathbf{X} \right]^{-1} = \sigma^2 \left[\mathbf{B}_0^{-1} + \mathbf{X}^\top \mathbf{X} \right]^{-1} \quad \text{e} \\ \mathbf{a} &= \boldsymbol{\Sigma} \left(\sigma^{-2} \mathbf{B}_0^{-1} \mathbf{b}_0 + \sigma^{-2} \mathbf{X}^\top \mathbf{X} \mathbf{b} \right). \end{aligned}$$

Assim,

$$\begin{aligned}(\boldsymbol{\beta}|\mathbf{y}, \mathbf{X}) &\sim t_{n-p}(\mathbf{a}, \boldsymbol{\Sigma}^*) \text{ e} \\ (\sigma^2|\mathbf{y}, \mathbf{X}) &\sim \mathcal{IG}(\nu + m, m\sigma_0^2 + \nu s^2).\end{aligned}$$

Exercício

Encontre a expressão exata de $\boldsymbol{\Sigma}^*$.

Em problemas mais complicados, teremos que recorrer a métodos numéricos, como a amostragem de Gibbs.

Para utilizar a amostragem de Gibbs no modelo de regressão com priori não informativa, temos

$$\begin{aligned}(\boldsymbol{\beta}|\sigma^2, \mathbf{y}, \mathbf{X}) &\sim \mathcal{N}_p(\mathbf{b}, \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1}) \text{ e} \\ (\sigma^2|\boldsymbol{\beta}, \mathbf{y}, \mathbf{X}) &\sim \mathcal{IG}(n/2, C(\boldsymbol{\beta})),\end{aligned}$$

sendo $C(\boldsymbol{\beta}) = \nu s^2 + (\boldsymbol{\beta} - \mathbf{b})^\top \mathbf{X}^\top \mathbf{X}(\boldsymbol{\beta} - \mathbf{b})$.

Modelo probit

Considere o modelo probit

$$\begin{aligned}y_i^* &= \mathbf{x}_i^\top \boldsymbol{\beta} + \varepsilon_i, \quad \varepsilon_i \sim \mathcal{N}(0, 1), \\y_i &= 1 \quad \text{if } y_i^* > 0, \quad \text{caso contrário } y_i = 0.\end{aligned}$$

Temos também que

$$Pr(y_i = 1 | \mathbf{x}_i) = \Phi(\mathbf{x}_i^\top \boldsymbol{\beta}) \quad \text{e} \quad Pr(y_i = 0 | \mathbf{x}_i) = 1 - \Phi(\mathbf{x}_i^\top \boldsymbol{\beta}).$$

A função de verossimilhança

$$L(\mathbf{y} | \mathbf{X}, \boldsymbol{\beta}) = \prod_{i=1}^n \left[\Phi(\mathbf{x}_i^\top \boldsymbol{\beta}) \right]^{y_i} \left[1 - \Phi(\mathbf{x}_i^\top \boldsymbol{\beta}) \right]^{1-y_i}$$

e teremos a prior como $p(\boldsymbol{\beta}) \propto 1$. Assim,

$$p(\boldsymbol{\beta} | \mathbf{y}, \mathbf{X}) = \frac{\prod_{i=1}^n \left[\Phi(\mathbf{x}_i^\top \boldsymbol{\beta}) \right]^{y_i} \left[1 - \Phi(\mathbf{x}_i^\top \boldsymbol{\beta}) \right]^{1-y_i}}{\int \prod_{i=1}^n \left[\Phi(\mathbf{x}_i^\top \boldsymbol{\beta}) \right]^{y_i} \left[1 - \Phi(\mathbf{x}_i^\top \boldsymbol{\beta}) \right]^{1-y_i} d\boldsymbol{\beta}}.$$

A média a posteriori é dada por

$$E(\beta|\mathbf{y}, \mathbf{X}) = \frac{\int \beta \prod_{i=1}^n [\Phi(\mathbf{x}_i^\top \beta)]^{y_i} [1 - \Phi(\mathbf{x}_i^\top \beta)]^{1-y_i} d\beta}{\int \prod_{i=1}^n [\Phi(\mathbf{x}_i^\top \beta)]^{y_i} [1 - \Phi(\mathbf{x}_i^\top \beta)]^{1-y_i} d\beta}.$$

O cálculo das integrais acima é complicado.

Consideraremos a técnica conhecida como variáveis auxiliares.

Trataremos y_i^* como quantidades não observadas a serem estimadas junto com β .

Primeiro, temos que $p(\beta|\mathbf{y}^*, \mathbf{y}, \mathbf{X}) = p(\beta|\mathbf{y}^*, \mathbf{X})$ e também

$$(\beta|\mathbf{y}^*, \mathbf{y}, \mathbf{X}) \sim \mathcal{N}(\mathbf{b}^*, (\mathbf{X}^\top \mathbf{X})^{-1}) \quad \text{sendo} \quad \mathbf{b}^* = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}^*.$$

Agora, sabemos que $(y_i^*|\beta, \mathbf{X}) \sim \mathcal{N}(\mathbf{x}_i^\top \beta, 1)$. Mas y_i tem informação sobre y_i^* . Logo,

$$\begin{aligned}(y_i^*|y_i = 1, \beta, \mathbf{x}_i) &\sim \mathcal{N}(\mathbf{x}_i^\top \beta, 1)I(0, +\infty) \\ (y_i^*|y_i = 0, \beta, \mathbf{x}_i) &\sim \mathcal{N}(\mathbf{x}_i^\top \beta, 1)I(-\infty, 0).\end{aligned}$$

Com estas equações podemos utilizar a amostragem de Gibbs.

Portanto, $E(\beta|\mathbf{y}, \mathbf{X})$ e $E(\mathbf{y}^*|\mathbf{y}, \mathbf{X})$ pode ser estimado com a amostra da posteriori.

Exemplo

Temos o modelo,

$$Pr(GRADE_i = 1 | \beta, \mathbf{x}_i) = \Phi(\beta_1 + \beta_2 GPA_i + \beta_3 TUCE_i + \beta_4 PSI_i),$$

sendo

- GPA_i a média dos pontos do estudante;
- $GRADE_i$ uma variável dicotômica indicando se a nota do estudante no curso de macroeconomia intermediário foi mais alta que nos cursos iniciais;
- PSI_i uma variável dicotômica indicando se o indivíduo participou do *Personalized System Instruction*; e
- $TUCE_i$ a nota do estudante em um teste de economia.

Tabela 20: Estimação do modelo probit para as notas.

Variável	EMV		Bayesiano	
	Estimativa	E.P.	Média	D.P.
Constante	-7,4523	2,5425	-8,6286	2,7995
GPA	1,6258	0,6939	1,8754	0,7668
TUCE	0,0517	0,0839	0,0628	0,0870
PSI	1,4263	0,5950	1,6072	0,6257

Estimação consistente: o método dos momentos

- Momentos populacionais e amostrais.
- Os momentos populacionais dependem dos parâmetros que estamos interessados em estimar.
- A amostra contém informações em forma de estatísticas como momentos.
- Se o modelo tiver p parâmetros, igualamos os p momentos populacionais aos p momentos amostrais correspondentes.
- Então, escrevemos os parâmetros como funções dos momentos amostrais (das estatísticas).
- Os momentos serão consistentes pela lei dos grandes números.
- Os momentos serão assintoticamente normais pelo Teorema de Lindeberg-Levy.
- Aplicação do teorema de Slutsky garante a propriedade dos estimadores dos parâmetros.

Seja $m_k(\cdot)$ uma função contínua e diferenciável que não depende do tamanho de amostra n e seja

$$\bar{m}_k = \frac{1}{n} \sum_{i=1}^n m_k(y_i), \quad k = 1, \dots, p.$$

Segue-se que

$$\text{plim } \bar{m}_k = E(m_k(y_i)) = \mu(\theta_1, \dots, \theta_p).$$

As p equações dos momentos são

$$\begin{aligned} \bar{m}_{n,1}(\theta) &= \bar{m}_1 - \mu_1(\theta_1, \dots, \theta_p) = 0 \\ \bar{m}_{n,2}(\theta) &= \bar{m}_2 - \mu_2(\theta_1, \dots, \theta_p) = 0 \\ &\vdots \quad \vdots \quad \quad \quad \vdots \quad \quad \vdots \\ \bar{m}_{n,p}(\theta) &= \bar{m}_p - \mu_p(\theta_1, \dots, \theta_p) = 0. \end{aligned}$$

- Na maioria dos casos os estimadores pelo método dos momentos não são eficientes.
- A exceção é na amostragem aleatória das distribuições na família exponencial.

Definição (família exponencial)

Uma distribuição pertence a família exponencial se tem o logaritmo da função de verossimilhança dada por

$$\ln L(\theta|\mathbf{dados}) = a(\mathbf{dados}) + b(\theta) + \sum_{k=1}^p c_k(\mathbf{dados}) s_k(\theta),$$

sendo $a(\cdot)$, $b(\cdot)$, $c_k(\cdot)$ e $s_k(\cdot)$ funções.

Lembremos que $c_k(\mathbf{dados})$ é uma estatística suficiente.

Em geral, consideramos os momentos como

$$\overline{m}_k = \frac{1}{n} \sum_{i=1}^n m_k(\mathbf{y}_i), \quad k = 1, \dots, p.$$

Um estimador apropriado da matriz de covariância assintótica de $\overline{\mathbf{m}} = (\overline{m}_1, \dots, \overline{m}_p)$ é

$$\frac{1}{n} F_{jk} = \frac{1}{n} \left\{ \frac{1}{n} \sum_{i=1}^n [m_j(\mathbf{y}_i) - \overline{m}_j][m_k(\mathbf{y}_i) - \overline{m}_k] \right\}, \quad j, k = 1, 2, \dots, p.$$

Agora, seja $\overline{\mathbf{G}}(\boldsymbol{\theta})$ a matriz $p \times p$ cuja k -ésima linha é o vetor de derivadas parciais

$$\overline{\mathbf{G}}_k^\top = \frac{\partial \overline{m}_k(\boldsymbol{\theta})}{\partial \boldsymbol{\theta}^\top}.$$

Da expansão de Taylor em torno do valor verdadeiro $\boldsymbol{\theta}_0$, temos

$$\mathbf{0} \simeq \overline{\mathbf{m}}(\boldsymbol{\theta}_0) + \overline{\mathbf{G}}^\top(\boldsymbol{\theta}_0)(\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}_0).$$

Portanto,

$$\sqrt{n}(\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}_0) \simeq -[\bar{\mathbf{G}}^\top(\boldsymbol{\theta}_0)]^{-1} \sqrt{n}[\bar{\mathbf{m}}(\boldsymbol{\theta}_0)].$$

Usando esta aproximação, pode-se provar que

$$\hat{\boldsymbol{\theta}} \approx \mathcal{N}(\boldsymbol{\theta}_0, \boldsymbol{\Delta}),$$

sendo $\boldsymbol{\Delta} = \frac{1}{2} \{ -[\boldsymbol{\Gamma}(\boldsymbol{\theta}_0)]^{-1} \} \boldsymbol{\Phi} \{ -[\boldsymbol{\Gamma}(\boldsymbol{\theta}_0)]^{-1} \}.$

Podemos estimar a variância assintótica $\boldsymbol{\Delta}$ por

$$\widehat{\text{Var}}(\hat{\boldsymbol{\theta}}) = \frac{1}{n} \left[\bar{\mathbf{G}}^\top(\hat{\boldsymbol{\theta}}) \mathbf{F}^{-1} \bar{\mathbf{G}}(\hat{\boldsymbol{\theta}}) \right]^{-1}.$$

GMM sob a condição de ortogonalidade

Considere o estimador OLS no modelo de regressão clássico. Por hipótese,

$$E(\mathbf{x}_i \varepsilon_i) = E(\mathbf{x}_i (y_i - \mathbf{x}_i^\top \boldsymbol{\beta})) = \mathbf{0}.$$

Na amostra, temos

$$\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \hat{\varepsilon}_i = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i (y_i - \mathbf{x}_i^\top \hat{\boldsymbol{\beta}}) = \mathbf{0}.$$

Podemos ver que o estimador OLS é um estimador do método dos momentos.

Para o estimador de variáveis instrumentais,

$$\text{plim} \left(\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i \varepsilon_i \right) = \text{plim} \left(\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i (y_i - \mathbf{x}_i^\top \beta) \right).$$

Na amostra, temos

$$\left(\frac{1}{n} \mathbf{X}^\top \mathbf{Z} \right) \left(\frac{1}{n} \mathbf{Z}^\top \mathbf{Z} \right)^{-1} \left(\frac{1}{n} \mathbf{X}^\top \hat{\boldsymbol{\varepsilon}} \right) = \frac{1}{n} \hat{\mathbf{X}}^\top \hat{\boldsymbol{\varepsilon}} = \frac{1}{n} \sum_{i=1}^n \hat{\mathbf{x}}_i \hat{\varepsilon}_i = \mathbf{0}.$$

Podemos ver que este estimador é um estimador do método dos momentos.

Estimadores de máxima verossimilhança são obtidos ao igualar a derivada da log-verossimilhança a zero.

A função de verossimilhança ponderada é

$$\frac{1}{n} \ln L(\boldsymbol{\theta}|\mathbf{y}, \mathbf{X}) = \frac{1}{n} \sum_{i=1}^n \ln f(y_i|\mathbf{x}_i, \boldsymbol{\theta}).$$

Sabemos que

$$E\left(\frac{\partial \ln f(y_i|\mathbf{x}_i, \boldsymbol{\theta})}{\partial \boldsymbol{\theta}}\right) = \mathbf{0}.$$

Na amostra, temos

$$\frac{1}{n} \frac{\partial \ln L(\hat{\boldsymbol{\theta}}|\mathbf{y}, \mathbf{X})}{\partial \hat{\boldsymbol{\theta}}} = \frac{1}{n} \sum_{i=1}^n \frac{\partial \ln f(y_i|\mathbf{x}_i, \hat{\boldsymbol{\theta}})}{\partial \hat{\boldsymbol{\theta}}} = \mathbf{0}.$$

Podemos também ver que este estimador é um estimador do método dos momentos.

Generalizando o método dos momentos

- Existem casos em que temos mais equações de momentos do que parâmetros. Logo, o sistema é sobre identificado.
- Para utilizar toda informação na amostra é necessário criar um critério para evitar os conflitos que podem existir em um sistema sobre identificado.
- Suponha que o modelo envolva p parâmetros, $\theta = (\theta_1, \dots, \theta_p)^\top$, e que a teoria fornece um conjunto de $L > p$ condições de momentos,

$$E(m_\ell(y_i, \mathbf{x}_i, \mathbf{z}_i, \theta)) = E(m_{i\ell}(\theta)) = 0,$$

sendo y_i , $\mathbf{x}_i$ e $\mathbf{z}_i$ variáveis que aparecem no modelo.

- Denotaremos os correspondentes momentos amostrais por

$$\bar{m}_\ell(\mathbf{y}, \mathbf{X}, \mathbf{Z}, \theta) = \frac{1}{n} \sum_{i=1}^n m_\ell(y_i, \mathbf{x}_i, \mathbf{z}_i, \theta) = \frac{1}{n} \sum_{i=1}^n m_{i\ell}(\theta).$$

- A menos que as equações sejam funcionalmente dependentes, o sistema de L equações em p parâmetros desconhecidos,

$$\overline{m}_\ell(\theta) = \frac{1}{n} \sum_{i=1}^n m_\ell(y_i, \mathbf{x}_i, \mathbf{z}_i, \theta) = 0, \quad \ell = 1, \dots, L,$$

não terá uma solução única.

- Uma das possibilidades é minimizar a soma de quadrados,

$$q = \sum_{\ell=1}^L \overline{m}_\ell^2 = \overline{\mathbf{m}}(\theta)^\top \overline{\mathbf{m}}(\theta).$$

- Pode ser mostrado que $\text{plim } \overline{\mathbf{m}}(\theta) = \mathbf{E}(\overline{\mathbf{m}}(\theta)) = \mathbf{0}$, isto é, o minimizador de q é um estimador consistente de θ .
- Podemos também utilizar mínimos quadrados ponderados,

$$q = \overline{\mathbf{m}}(\theta)^\top \mathbf{W}_n \overline{\mathbf{m}}(\theta),$$

sendo $\mathbf{W}_n$ qualquer matriz positiva definida que depende dos dados mas não dos parâmetros. Supomos também que $\text{plim } \mathbf{W}_n = \mathbf{W}$, uma matriz positiva definida.

- Podemos utilizar uma matriz $\mathbf{W}_n$ como diagonal com elementos que são os inversos das variâncias individuais dos momentos,

$$w_{\ell\ell} = \frac{1}{\text{Var}(\sqrt{n}\bar{m}_\ell)} = \frac{1}{\phi_{\ell\ell}}.$$

- O estimador de mínimos quadrados ponderados deve minimizar

$$q = \bar{\mathbf{m}}(\theta)^\top \boldsymbol{\Phi}^{-1} \bar{\mathbf{m}}(\theta).$$

- Em geral, os L elementos de $\bar{\mathbf{m}}$ podem ser correlacionados.
- Para utilizar mínimos quadrados generalizados, podemos definir a matriz,

$$\mathbf{W} = [\text{Var}(\sqrt{n}\bar{\mathbf{m}})]^{-1} = \boldsymbol{\Phi}^{-1}.$$

- Se $\mathbf{W}_n$ é uma matriz positiva definida e se $\text{plim } \bar{\mathbf{m}}(\theta) = \mathbf{0}$, então o estimador GMM é consistente.

- A matriz de covariância assintótica do estimador GMM é

$$\mathbf{V}_{GMM} = \frac{1}{n} [\mathbf{\Gamma}^\top \mathbf{W} \mathbf{\Gamma}]^{-1} = \frac{1}{n} [\mathbf{\Gamma}^\top \mathbf{\Phi}^{-1} \mathbf{\Gamma}]^{-1},$$

sendo $\mathbf{\Gamma}$ a matriz de derivadas com a j -ésima linha igual a

$$\mathbf{\Gamma}^j = \text{plim} \frac{\partial \bar{m}_j(\boldsymbol{\theta})}{\partial \boldsymbol{\theta}^\top},$$

e $\mathbf{\Phi} = \text{Var}(\sqrt{n}\bar{\mathbf{m}})$.

- Por hipótese, temos que

$$\sqrt{n}\bar{\mathbf{m}}(\boldsymbol{\theta}_0) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \mathbf{\Phi}).$$

Teorema

A distribuição assintótica para o estimador GMM é

$$\begin{aligned} \hat{\boldsymbol{\theta}}_{GMM} &\xrightarrow{p} \boldsymbol{\theta} \\ \hat{\boldsymbol{\theta}}_{GMM} &\approx \mathcal{N}(\boldsymbol{\theta}, \mathbf{V}_{GMM}). \end{aligned}$$

Estimação GMM em modelos econométricos

Consideraremos o modelo de regressão com variáveis instrumentais,

$$\begin{aligned}y_i &= \mathbf{x}_i^\top \boldsymbol{\beta} + \varepsilon_i \\ \mathbf{E}(\mathbf{z}_i \varepsilon_i) &= \mathbf{0},\end{aligned}$$

para p variáveis em $\mathbf{x}_i$ e um conjunto de $L \geq p$ variáveis instrumentais $\mathbf{z}_i$.

O resultado da regressão generalizada aparece se $\mathbf{z}_i = \mathbf{x}_i$ e a regressão clássica ao adicionar que $\boldsymbol{\Omega} = \mathbf{I}$.

Consideraremos três casos:

- Regressão clássica: $\text{Var}(\varepsilon_i | \mathbf{X}, \mathbf{Z}) = \sigma^2$,
- Heterocedasticidade: $\text{Var}(\varepsilon_i | \mathbf{X}, \mathbf{Z}) = \sigma_i^2$,
- Regressão generalizada: $\text{Cov}(\varepsilon_t, \varepsilon_s | \mathbf{X}, \mathbf{Z}) = \sigma^2 \omega_{ts}$.

A hipótese $E(\mathbf{z}_i \varepsilon_i) = 0$ implica a condição de ortogonalidade

$$\text{Cov}(\mathbf{z}_i, \varepsilon_i) = \mathbf{0} \quad \text{ou} \quad E(\mathbf{z}_i(y_i - \mathbf{x}_i^\top \beta)) = \mathbf{0}.$$

As equações de momentos da população são

$$E\left(\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i(y_i - \mathbf{x}_i^\top \beta)\right) = E(\overline{\mathbf{m}}(\beta)) = \mathbf{0}.$$

As equações de momentos da amostra são

$$\left[\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i(y_i - \mathbf{x}_i^\top \hat{\beta})\right] = \left[\frac{1}{n} \sum_{i=1}^n \mathbf{m}_i(\hat{\beta})\right] = \overline{\mathbf{m}}(\hat{\beta}) = \mathbf{0}.$$

Para os casos em que o modelo é identificável, podemos escrever

$$\overline{\mathbf{m}}(\hat{\beta}) = \left(\frac{1}{n} \mathbf{z}^\top \mathbf{y}\right) - \left(\frac{1}{n} \mathbf{z}^\top \mathbf{X}\right) \hat{\beta}.$$

Se o modelo for exatamente identificável ($L = p$), então pode ser mostrado que

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

Se o modelo for sobre identificável ($L > p$), então o sistema $\overline{\mathbf{m}}(\hat{\beta}) = \mathbf{0}$ não tem uma solução única.

Podemos escolher o critério de mínimos quadrados,

$$\min_{\hat{\beta}} q = \min_{\hat{\beta}} \overline{\mathbf{m}}(\hat{\beta})^\top \overline{\mathbf{m}}(\hat{\beta}).$$

Neste caso, as condições de primeira ordem são

$$\begin{aligned} \frac{\partial q}{\partial \beta} &= 2 \left(\frac{\partial \overline{\mathbf{m}}(\hat{\beta})}{\partial \beta} \right) \overline{\mathbf{m}}(\hat{\beta}) = 2 \overline{\mathbf{G}}(\hat{\beta})^\top \overline{\mathbf{m}}(\hat{\beta}) \\ &= 2 \left(\frac{1}{n} \mathbf{X}^\top \mathbf{Z} \right) \left(\frac{1}{n} \mathbf{Z}^\top \mathbf{y} - \frac{1}{n} \mathbf{Z}^\top \mathbf{X} \hat{\beta} \right) = \mathbf{0}, \end{aligned}$$

cujas solução é

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

Hipóteses sobre estimador GMM:

- Convergência dos momentos: $\text{plim } \bar{\mathbf{m}}(\hat{\beta}) = \mathbf{0}$.
- Identificação para estimação GMM: a matriz $L \times p$

$$\mathbf{\Gamma}(\beta) = E(\bar{\mathbf{G}}(\beta)) = \text{plim } \bar{\mathbf{G}}(\beta) = \text{plim } \frac{\partial \bar{\mathbf{m}}}{\partial \beta^\top} = \text{plim } \frac{1}{n} \sum_{i=1}^n \frac{\partial \bar{\mathbf{m}}_i}{\partial \beta^\top}$$

deve ter posto de linha igual a p .

- Os momentos amostrais convergem para uma distribuição normal.

É possível mostrar que sob certas condições o estimador GMM é consistente.

A matriz de covariância assintótica de $\hat{\beta}$ é dada por

$$\text{Var}(\hat{\beta}) = \frac{1}{n} [\mathbf{\Gamma}^\top \mathbf{\Gamma}]^{-1} \mathbf{\Gamma}^\top \text{Var}(\sqrt{n} \bar{\mathbf{m}}(\beta)) \mathbf{\Gamma} [\mathbf{\Gamma}^\top \mathbf{\Gamma}]^{-1}.$$

Para o modelo de regressão em questão, temos

$$\begin{aligned}\bar{\mathbf{m}} &= \frac{1}{n}(\mathbf{Z}^\top \mathbf{y} - \mathbf{Z}^\top \mathbf{Z} \boldsymbol{\beta}), \\ \bar{\mathbf{G}}(\boldsymbol{\beta}) &= \frac{1}{n} \mathbf{Z}^\top \mathbf{X}, \\ \boldsymbol{\Gamma}(\boldsymbol{\beta}) &= \mathbf{Q}_{\mathbf{Z}\mathbf{X}}.\end{aligned}$$

Precisamos, agora, determinar a forma de

$$\mathbf{V} = \text{Var}(\sqrt{n}\bar{\mathbf{m}}(\boldsymbol{\beta})).$$

Dado a forma de $\bar{\mathbf{m}}$ acima, temos

$$\mathbf{V} = \frac{1}{n} \text{Var} \left(\sum_{i=1}^n \mathbf{z}_i \varepsilon_i \right) = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^n \sigma^2 \omega_{ij} \mathbf{z}_i \mathbf{z}_j^\top = \sigma^2 \frac{\mathbf{Z}^\top \boldsymbol{\Omega} \mathbf{Z}}{n}$$

para o caso mais geral.

Agora, precisamos estimar a variância assintótica. Para isto, precisamos estimar $\mathbf{V}$.

O estimador de $\mathbf{V}$ é dado por

- Regressão clássica:

$$\hat{\mathbf{V}} = \frac{(\mathbf{e}^\top \mathbf{e}/n)}{n} \sum_{i=1}^n \mathbf{z}_i \mathbf{z}_i^\top = \frac{(\mathbf{e}^\top \mathbf{e}/n)}{n} \mathbf{Z}^\top \mathbf{Z}.$$

- Regressão heterocedástica:

$$\hat{\mathbf{V}} = \frac{1}{n} \sum_{i=1}^n e_i^2 \mathbf{z}_i \mathbf{z}_i^\top.$$

- Regressão generalizada:

$$\begin{aligned} \hat{\mathbf{V}} &= \frac{1}{n} \left[\sum_{t=1}^n e_t^2 \mathbf{z}_t \mathbf{z}_t^\top \right. \\ &\quad \left. + \sum_{\ell=1}^n \left(1 - \frac{\ell}{p+1} \right) \sum_{t=\ell+1}^n e_t e_{t-\ell} (\mathbf{z}_t \mathbf{z}_{t-\ell} + \mathbf{z}_{t-\ell} \mathbf{z}_t^\top) \right]. \end{aligned}$$

Entretanto, se os elementos de $\overline{\mathbf{m}}(\beta)$ forem não correlacionados, podemos utilizar

$$\hat{\mathbf{V}} = \frac{1}{n} \sum_{i=1}^n \mathbf{m}_i(\hat{\beta}) \mathbf{m}_i(\hat{\beta})^\top.$$

Assim, estimamos a variância assintótica por

$$\begin{aligned}\widehat{\text{Var}}(\hat{\beta}) &= \frac{1}{n} \left[\overline{\mathbf{G}}(\hat{\beta})^\top \overline{\mathbf{G}}(\hat{\beta}) \right]^{-1} \overline{\mathbf{G}}(\hat{\beta})^\top \widehat{\mathbf{V}} \overline{\mathbf{G}}(\hat{\beta}) \left[\overline{\mathbf{G}}(\hat{\beta})^\top \overline{\mathbf{G}}(\hat{\beta}) \right]^{-1} \\ &= n \left[(\mathbf{X}^\top \mathbf{Z})(\mathbf{Z}^\top \mathbf{X}) \right]^{-1} (\mathbf{X}^\top \mathbf{Z}) \widehat{\mathbf{V}} (\mathbf{Z}^\top \mathbf{X}) \left[(\mathbf{X}^\top \mathbf{Z})(\mathbf{Z}^\top \mathbf{X}) \right]^{-1}.\end{aligned}$$

Agora, queremos utilizar o critério de mínimos quadrados ponderados

$$\min_{\hat{\beta}} q = \min_{\hat{\beta}} \overline{\mathbf{m}}(\hat{\beta})^\top \mathbf{W} \overline{\mathbf{m}}(\hat{\beta})$$

sendo $\mathbf{W}$ qualquer matriz positiva definida de ponderação. Assim, temos

$$\widehat{\text{Var}}(\hat{\beta}) = \frac{1}{n} \left[\overline{\mathbf{G}}^\top \mathbf{W} \overline{\mathbf{G}} \right]^{-1} \overline{\mathbf{G}}^\top \mathbf{W} \widehat{\mathbf{V}} \mathbf{W} \overline{\mathbf{G}} \left[\overline{\mathbf{G}}^\top \mathbf{W} \overline{\mathbf{G}} \right]^{-1}.$$

No caso anterior, tínhamos $\mathbf{W} = \mathbf{I}$.

Se $L = p$, então $\mathbf{W}$ é irrelevante para a solução. Então, baseado na simplicidade, o valor ótimo de $\mathbf{W}$ é $\mathbf{I}$.

No caso de sobre identificação o valor ótimo de W é V^{-1} . O estimador resultante será o mais eficiente. Logo,

$$\hat{\beta}_{GMM} = [(X^T Z) \hat{V}^{-1} (Z^T X)]^{-1} (X^T Z) \hat{V}^{-1} (Z^T y)$$

e

$$\text{Var}(\hat{\beta}_{GMM}) = \frac{1}{n} [\bar{G}^T V^{-1} \bar{G}]^{-1} = [(X^T Z) V^{-1} (Z^T X)]^{-1}.$$

Concluindo, o estimador GMM é dado pela solução de

$$\min_{\beta} q = \min_{\beta} \bar{m}(\beta)^T [\text{Var}(\sqrt{n} \bar{m}(\beta))]^{-1} \bar{m}(\beta).$$

que sugere que a matriz de pesos é uma função de β - o vetor de parâmetros que estamos interessados em estimar.

A estimação GMM é, em geral, dada em dois passos. No Passo 1 estimamos V e no Passo 2 estimamos β utilizando a inversa de $\hat{V}$.

Passo 1: Utilize $\mathbf{W} = \mathbf{I}$ para obter um estimador consistente de β .
Então, estime $\mathbf{V}$ com

$$\hat{\mathbf{V}} = \frac{1}{n} \sum_{i=1}^n e_i^2 \mathbf{z}_i \mathbf{z}_i^{\top}$$

no caso de heterocedasticidade (o estimador de White) ou,
para o caso mais geral, o estimador de Newey-West.

Passo 2: Utilize $\mathbf{W} = \hat{\mathbf{V}}^{-1}$ para calcular o estimador GMM de β .