

EDUARDO HORTA

ECONOMETRIA

PPGE/UFRGS

Copyright © 2022 Eduardo Horta eduardo.horta@ufrgs.br

This is a manuscript — please do not share!

PUBLISHED BY PPGE/UFRGS

Licensed under the Apache License, Version 2.0 (the “License”); you may not use this file except in compliance with the License. You may obtain a copy of the License at <http://www.apache.org/licenses/LICENSE-2.0>. Unless required by applicable law or agreed to in writing, software distributed under the License is distributed on an “AS IS” BASIS, WITHOUT WARRANTIES OR CONDITIONS OF ANY KIND, either express or implied. See the License for the specific language governing permissions and limitations under the License.

First printing, October 2020

Sumário

Prelúdio 7

O modelo de regressão linear 13

*Esperança condicional e
modelos de regressão:
um interlúdio* 21

O método de mínimos quadrados 31

*Propriedades amostrais do estimador
de mínimos quadrados* 41

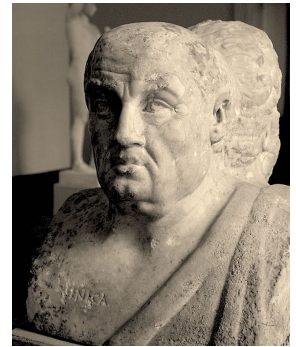
περὶ τῶν ἀσυμπτῶτων (*sobre assíntotas*) 49

Consistência e Gaussianidade assintóticas 61

Sob a tese... 67

Referências Bibliográficas 75

The time will come when diligent research over long periods will bring to light things which now lie hidden. A single lifetime, even though entirely devoted to the sky, would not be enough for the investigation of so vast a subject... And so this knowledge will be unfolded only through long successive ages. There will come a time when our descendants will be amazed that we did not know things that are so plain to them... Many discoveries are reserved for ages still to come, when memory of us will have been effaced.



Lucius Annaeus Seneca,
Naturales quaestiones.

Prelúdio

UM TEXTO SOBRE ECONOMETRIA — tanto mais um texto para alunos de pós-graduação — não deveria se esquivar de abordar questões epistemológicas: não resta dúvida de que esse é um assunto que desperta o interesse do estudante dotado de uma mente questionadora. Nosso tempo é exíguo, contudo, e me resta indicar aos curiosos um ponto de partida para a literatura relevante.¹ Um texto sobre Econometria tampouco deveria se esquivar de apresentar um panorama da pesquisa — teórica e aplicada — nessa disciplina. Com o risco de levar a pecha de pedante — afinal, para que mencionar aquilo que se deve fazer e não fazê-lo? — recorro novamente ao expediente de apenas indicar uma referência importante, qual seja, o *review* escrito por John Geweke, Joel Horowitz e Hashem Pesaran em 2006 com o título *Econometrics: A Bird's Eye View* e republicado em 2008 no *New Palgrave Dictionary of Economics*² com um título um tanto mais prosaico.

O leitor já pode deduzir, após o parágrafo acima, que esse texto apresentará a teoria Econométrica desde a perspectiva do formalismo matemático — é isso que o tempo nos permite, e é a partir do entendimento desse formalismo que se pode efetivamente compreender questões relativas à teoria do conhecimento e à metodologia científica, bem como proceder a aplicações que não se restrinjam a mero mimetismo (“ $p < 0.01 \rightarrow$ posso rejeitar a hipótese nula!”). Ainda assim, há pelo menos uma questão de caráter epistemológico sobre a qual eu eu gostaria de comentar antes de prosseguir, que é a seguinte: no jargão econométrico, o conceito de *Data Generating Process* (DGP) é amplamente empregado, em que pese o fato de ser uma noção imprecisa. Todavia, em que pese o fato de ser uma noção

¹ Trygve Haavelmo. The Probability Approach in Econometrics. *Econometrica*, 12:iii–115, 1944; Olav Bjerkholt. On the Founding of the Econometric Society. *Journal of the History of Economic Thought*, 39(2), 2017; David F. Hendry. Econometrics – Alchemy or Science? *Economica*, 47 (188):387–406, 1980; and Kevin D. Hoover. The Methodology of Econometrics. In Terence C. Mills and Kerry Patterson, editors, *Palgrave Handbook of Econometrics*, volume 1. Palgrave Macmillan UK, 2006

² John Geweke, Joel L. Horowitz, and Hashem Pesaran. *Econometrics*, pages 1–44. Palgrave Macmillan UK, London, 2017

imprecisa, devemos admitir que há muito didatismo em se considerar que a realidade/natureza/ Tύχη (Fortuna) gera os fenômenos que observamos utilizando um mecanismo capaz de produzir observáveis que *parecem* aleatórios (e.g., um programa de computador), mecanismo esse configurado com parâmetros por ela escolhidos (e, para nós, desconhecidos), e que nosso objetivo nesse jogo é oferecer um *educated guess* sobre qual o valor desses parâmetros, baseado nas observações (por Tύχη reveladas) a que temos acesso. Todavia, há pelo menos dois inconvenientes nessa ideia que merecem ser mencionados, ambos relacionados à dificuldade de se estabelecer uma ponte entre um discurso sobre fenômenos de um lado e um discurso delimitado por uma linguagem formal de outro:

1. **Quantidades numéricas observadas no mundo real não são variáveis aleatórias!** Variáveis aleatórias³ são objetos constrictos a uma linguagem formal/axiomática. Por exemplo, a ‘*taxa de juros de amanhã*’ é um valor numérico *in potentia*, sobre o qual temos incerteza, mas não faz sentido perguntar se é uma função mensurável definida em um espaço de probabilidade, etc. Essa distinção costumeiramente é deixada de lado, e dizemos despreocupadamente que a ‘*taxa de juros de amanhã*’ é uma variável aleatória. É importante ressaltar que é imaterial, no exemplo acima, o fato de estarmos considerando um evento futuro; poderíamos com igual valor didático tomar a ‘*taxa de juros de ontem*’ se, por alguma razão, desconhecêssemos essa quantidade.
2. Sentenças da forma “A probabilidade de E é p ”, onde E é um evento exprimível em linguagem natural e p é um número real entre 0 e 1, **tipicamente não fazem referência a um atributo do mundo real**. Por exemplo, considere o experimento que consiste em lançar uma moeda equilibrada duas vezes consecutivas, em lançamentos fisicamente independentes. Um modelo formal para esse experimento é o seguinte: consideramos o espaço amostral $\Omega = \{HH, HT, TH, TT\}$ (H para cara, T para coroa) munido da medida de probabilidade $\mathbb{P}$ definida via $\mathbb{P}\{\omega\} = 1/4$ para $\omega \in \Omega$. Nesse contexto, o evento (expresso em linguagem natural) “observar *cara* no primeiro lançamento” corresponde ao evento $\{HH, HT\}$, e a sentença “a probabilidade de observar-



Tύχη de Antioquia, cópia romana de um original em bronze feito por Eutíquides, *Galleria dei Candelabri*.

³ Lembre que uma variável aleatória é uma função $x: \Omega \rightarrow \mathbb{R}$, onde Ω é um conjunto não-vazio munido de uma σ -álgebra $\mathcal{A}$, tal que os conjuntos

$$\{x \leq \xi\} := \{\omega \in \Omega: x(\omega) \leq \xi\}$$

são **eventos** (isto é, elementos de $\mathcal{A}$), para qualquer $\xi \in \mathbb{R}$. Ao conjunto Ω chamamos de **espaço amostral** e consideramos que ele representa a coleção de todos os possíveis estados do mundo.

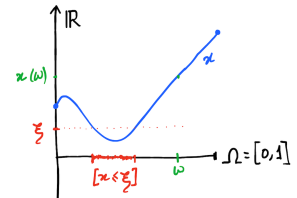


Figura 1: Exemplo de uma variável aleatória no espaço amostral $\Omega = [0, 1]$.

mos *cara* no primeiro lançamento é (igual a) $1/2$ ” corresponde à sentença formal

$$\mathbb{P}\{HH, HT\} = 1/2.$$

Note, porém que a função $\mathbb{P}$ acima foi *escolhida*, dentro do escopo de uma linguagem formal, para satisfazer nossa intuição de qual seria um modelo adequado para descrever o fenômeno sob consideração, de acordo com aquilo que o enunciado do problema deixa transparecer: “moeda equilibrada”, “lançamentos independentes”, etc. Note também que, se formos exigentes no nosso uso da linguagem, deveríamos traduzir a sentença “ $\mathbb{P}\{HH, HT\} = 1/2$ ” como “a $\mathbb{P}$ -probabilidade de observarmos *cara* no primeiro lançamento é (igual a) $1/2$ ”. Esse exemplo deixa transparecer — em que pese o fato de que o modelo escolhido satisfaz, sim, um critério (informal) de razoabilidade — que o valor numérico $1/2$ aparecendo na sentença acima não é “algo que está no mundo”. Sendo assim, o melhor é interpretar esse valor numérico como uma quantificação de nossa incerteza a respeito de um fenômeno que temos interesse em modelar, e essa quantificação depende de uma escolha (nossa!) de qual modelo formal vamos utilizar: em suma, sentenças de probabilidade (*probability statements*) são epistêmicas! Em suma, sentenças do tipo “o verdadeiro valor do parâmetro é p ” nada mais são, na maioria das vezes, do que uma abstração, referentes a objetos que existem apenas dentro do escopo de uma linguagem formal.⁴ Como disse o célebre George Box,⁵

“All models are wrong, but some are useful.”

É importante ter isso em mente quando falamos, por exemplo, sobre “o verdadeiro valor de β ”, etc.

Sobre notações e outras noções

Após algum tempo convivendo com notações matemáticas é inevitável dar-se conta de que, como a busca pelo Santo Graal, o desejo por uma notação abrangente e unificada não é mais que uma busca vã. Em Probabilidade Axiomática, a notação canônica (se é que existe tal coisa) recomenda que variáveis aleatórias sejam denotadas por letras romanas maiúsculas, e.g. X, Y, Z ; que valores realizados

⁴ Embora, no caso de populações finitas, parâmetros teóricos correspondam a quantidades existentes na realidade.

⁵ George Box. Robustness in the Strategy of Scientific Model Building. In Robert L. Launer and Graham N. Wilkinson, editor, *Robustness in Statistics*, pages 201 – 236. Academic Press, 1979

(números reais) sejam denotados por letras romanas minúsculas, e.g. x, y, z etc. Na literatura de Econometria, todavia, consagrou-se a heresia em que variáveis aleatórias são denotadas por letras minúsculas, e nesse texto seguiremos essa tradição: os símbolos x, y, z, w, u, v (possivelmente acrescidos de um subíndice i ou ij) aqui denotam variáveis aleatórias; se estiverem em negrito (e.g., $\mathbf{x}$) representam *vetores* aleatórios; imitando o muito didático Prof. David Williams,⁶ usaremos os símbolos $x^{\text{obs}}, y^{\text{obs}}, z^{\text{obs}}, w^{\text{obs}}, u^{\text{obs}}, v^{\text{obs}}$ (possivelmente acrescidos de um subíndice i ou ij) para denotar valores realizados⁷ das variáveis aleatórias x, y, z, w, u, v , respectivamente; os símbolos ξ, η, ζ também denotam números reais genéricos, variáveis de integração etc (com sorte não precisaremos muito deles); os símbolos i, j, k, ℓ, n, m, d denotam números inteiros, índices etc.⁸ Nisso, orbitaremos — mas não excessivamente de perto — a notação utilizada por Stachurski⁹ e Hayashi.¹⁰ Em particular, dado um **tamanho de amostra** $n \in \mathbb{N}$, e dadas variáveis aleatórias $y_1, \dots, y_n$ e vetores aleatórios $\mathbf{x}_1, \dots, \mathbf{x}_n$, onde

$$\mathbf{x}'_i = \begin{pmatrix} x_{i1} & x_{i2} & \cdots & x_{id} \end{pmatrix}, \quad i \in \{1, \dots, n\},$$

escreveremos

$$\mathbf{y} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} \quad \text{e} \quad \mathbf{X} = \begin{pmatrix} \mathbf{x}'_1 \\ \mathbf{x}'_2 \\ \vdots \\ \mathbf{x}'_n \end{pmatrix} = \begin{pmatrix} x_{1,1} & x_{1,2} & \cdots & x_{1,d} \\ x_{2,1} & x_{2,2} & \cdots & x_{2,d} \\ \vdots & \vdots & \ddots & \vdots \\ x_{n,1} & x_{n,2} & \cdots & x_{n,d} \end{pmatrix}$$

onde $\mathbf{x}'$ denota o transposto do vetor $\mathbf{x}$ (e, semelhantemente, $\mathbf{X}'$ denota a transposta da matriz $\mathbf{X}$). Note que $\mathbf{y}$ é um vetor $n \times 1$, enquanto $\mathbf{x}_i$ é um vetor $d \times 1$. Note também que o correto aqui seria escrever $x_{i,j}$ ao invés de x_{ij} , mas nos permitiremos esse ligeiro abuso de notação onde não houver possibilidade de trocarmos os pés pelas mãos (por exemplo, de confundirmos o número 11 com o par 1, 1).

UM CONCEITO FUNDAMENTAL na literatura de modelos de regressão, a *esperança condicional* é definida da seguinte maneira: seja y uma variável aleatória **integrável** (quer dizer, y é tal que $\mathbb{E}(|y|) < \infty$), e seja $\mathbf{x}$ um vetor aleatório com valores em $\mathbb{R}^d$. A es-

⁶ David Williams. *Weighing the Odds: A Course in Probability and Statistics*. Cambridge University Press, 2001

⁷ Na prática, o conceito de “valor realizado” é vago; formalmente, $x^{\text{obs}}, y^{\text{obs}}, z^{\text{obs}}, w^{\text{obs}}, u^{\text{obs}}, v^{\text{obs}}$ denotam números reais.

⁸ Comparativamente, enquanto na notação canônica escreveríamos, por exemplo, $F_Y(y) = \mathbb{P}(Y \leq y)$ para denotar o valor da função de distribuição acumulada F_Y da variável aleatória Y no ponto $y \in \mathbb{R}$ (respectivamente, $\mathbb{E}(X) = \int_{-\infty}^{+\infty} x f_X(x) dx$ para denotar o valor esperado da variável aleatória absolutamente contínua X em termos de sua função densidade de probabilidade f_X), com a notação que seguiremos nessas notas escreveremos $F_y(\xi) = \mathbb{P}(y \leq \xi)$ para denotar o valor da função de distribuição acumulada F_y da variável aleatória y no ponto $\xi \in \mathbb{R}$ (respectivamente, $\mathbb{E}(x) = \int_{-\infty}^{+\infty} \xi f_x(\xi) d\xi$ para denotar o valor esperado da variável aleatória x em termos de sua função densidade de probabilidade f_x .)

⁹ John Stachurski. *A primer in econometric theory*. MIT Press, 2016

¹⁰ Fumio Hayashi. *Econometrics*. Princeton University Press, 2000

perança condicional de y dado $\mathbf{x}$ é a variável aleatória $\mathbb{E}(y | \mathbf{x})$ que satisfaz as seguintes condições:

EC1. Tem-se que $\mathbb{E}(y | \mathbf{x}) = \varphi \circ \mathbf{x}$ para alguma função mensurável φ definida em $\mathbb{R}^d$ e tomando valores na reta, com $\mathbb{E}(|\varphi \circ \mathbf{x}|) < \infty$.

EC2. Vale a igualdade $\mathbb{E}(\mathbb{E}(y | \mathbf{x}) \mathbb{I}_{[\mathbf{x} \in B]}) = \mathbb{E}(y \mathbb{I}_{[\mathbf{x} \in B]})$ para todo subconjunto Boreliano $B \subseteq \mathbb{R}^d$.

A função φ dada acima¹¹ é chamada a **função de regressão de y em $\mathbf{x}$** . Nessas condições, se $\mathbf{x}^{\text{obs}} \in \mathbb{R}^d$ é um possível valor assumido pelo vetor aleatório $\mathbf{x}$, definimos

$$\mathbb{E}(y | \mathbf{x} = \mathbf{x}^{\text{obs}}) := \varphi(\mathbf{x}^{\text{obs}}).$$

Atenção para não se confundir!

Remark 1. Por definição, dadas funções $f: A \rightarrow B$ e $h: B \rightarrow C$, denotamos por $h \circ f$ a função cujo domínio é A , cujo contradomínio é C e cujo valor no ponto $a \in A$ é o elemento $h(f(a)) \in C$; isto é, $h \circ f(a) := h(f(a))$. É costumeiro, em Probabilidade, Estatística, Econometria etc, escrever $g(\mathbf{x})$ ao invés de $g \circ \mathbf{x}$ quando $\mathbf{x}$ é um vetor aleatório de dimensão d e g é uma função de $\mathbb{R}^d$ em $\mathbb{R}$. Isso é um abuso de notação. Aqui, a fim de evitar ambiguidades, escreveremos $g \circ \mathbf{x}$ sempre que possível.

PARA FINALIZAR, será conveniente fazer a seguinte suposição ao longo de todo o texto:

TODAS AS VARIÁVEIS ALEATÓRIAS QUE CONSIDERAREMOS TÊM VARIÂNCIA FINITA.¹²

Digo que essa suposição é conveniente por duas razões: primeiro, porque ela implica integrabilidade das variáveis aleatórias que consideraremos (e assim estaremos poupados, ao percorrer o texto, de *exigir* integrabilidade de y antes de nos permitirmos escrever $\mathbb{E}(y)$, por exemplo); segundo, porque ela nos permite¹³ escrever $\mathbb{V}(\mathbf{x})$ para denotar a **matriz de covariância** de um vetor aleatório $\mathbf{x} = (x_1 \ x_2 \ \cdots \ x_d)'$:

$$\mathbb{V}(\mathbf{x}) := \begin{pmatrix} \mathbb{V}(x_1) & \text{cov}(x_1, x_2) & \cdots & \text{cov}(x_1, x_d) \\ \text{cov}(x_2, x_1) & \mathbb{V}(x_2) & \cdots & \text{cov}(x_2, x_d) \\ \vdots & \vdots & \ddots & \vdots \\ \text{cov}(x_d, x_1) & \text{cov}(x_d, x_2) & \cdots & \mathbb{V}(x_d) \end{pmatrix}$$

¹¹ É bem sabido que tal φ é essencialmente única, no seguinte sentido: se $g: \mathbb{R}^d \rightarrow \mathbb{R}$ é outra função (mensurável) satisfazendo a condição

$$\mathbb{E}\{g \circ \mathbf{x} \mathbb{I}_{[\mathbf{x} \in B]}\} = \mathbb{E}(y \mathbb{I}_{[\mathbf{x} \in B]})$$

para todo Boreliano $B \subseteq \mathbb{R}^d$, então vale que

$$\mathbb{P}(g \circ \mathbf{x} = \varphi \circ \mathbf{x}) = 1,$$

onde $\mathbb{I}_{[\mathbf{x} \in B]}$ é a variável aleatória definida pondo-se $\mathbb{I}_{[\mathbf{x} \in B]}(\omega) = 1$ se $\mathbf{x}(\omega) \in B$ e $\mathbb{I}_{[\mathbf{x} \in B]}(\omega) = 0$ se $\mathbf{x}(\omega) \notin B$.

¹² Define-se a **variância** de uma variável aleatória y como sendo a quantidade $\mathbb{V}(y)$ dada por

$$\mathbb{V}(y) := \mathbb{E}(y^2) - (\mathbb{E}y)^2.$$

Essa quantidade é não-negativa mas, possivelmente, infinita.

¹³ De fato, a desigualdade de Cauchy–Schwarz nos dá

$$|\text{cov}(x, y)|^2 \leq \mathbb{V}(x)\mathbb{V}(y).$$

onde $\text{cov}(x, y) := \mathbb{E}(xy) - \mathbb{E}(x)\mathbb{E}(y)$ para variáveis aleatórias x e y . Escreveremos, ainda,

$$\text{cov}(\mathbf{x}, \mathbf{u}) := \mathbb{E}(\mathbf{x}\mathbf{u}') - \mathbb{E}(\mathbf{x})\mathbb{E}(\mathbf{u}'),$$

quando $\mathbf{x}$ e $\mathbf{u}$ forem vetores aleatórios. No caso em que $\mathbf{u}$ tem dimensão 1 — digamos, $\mathbf{u} = (u)$ — escrevemos simplesmente $\text{cov}(\mathbf{x}, u)$ e desconsideramos os sinais de transposição na definição acima.

Exercícios

Quaisquer erros nos enunciados são — obviamente — intencionais.

Exercício 1. Seja $\Omega = (0, 1]$ munido de uma σ -álgebra $\mathcal{A}$ tal que todo sub-intervalo $(a, b] \subseteq \Omega$, com $0 < a \leq b \leq 1$, satisfaz a relação $(a, b] \in \mathcal{A}$. Verifique que a função $x: \Omega \rightarrow \mathbb{R}$ definida, para $\omega \in \Omega$, por $x(\omega) = \omega^2$ é uma variável aleatória (com respeito a $\mathcal{A}$).

Exercício 2. Sejam x e y variáveis aleatórias. Verifique que

$$\mathbb{E}[(x - \mathbb{E}(x)) \cdot (y - \mathbb{E}(y))] = \mathbb{E}(xy) - \mathbb{E}(x)\mathbb{E}(y)$$

Exercício 3. Considere os vetores aleatórios $\mathbf{y} = (y_1 \ \cdots \ y_n)'$, $\mathbf{x} = (x_1 \ x_2 \ \cdots \ x_d)'$ e $\mathbf{x}_i = (x_{i1} \ x_{i2} \ \cdots \ x_{id})'$, $i \in \{1, \dots, n\}$, e seja $\mathbf{X}$ a matriz cuja componente ij é x_{ij} . Sejam ainda Σ uma matriz $m \times d$ e $\boldsymbol{\xi} \in \mathbb{R}^m$, ambos não aleatórios. Verifique as seguintes igualdades:

1. $\mathbb{E}(\Sigma\mathbf{X}) = \Sigma\mathbb{E}(\mathbf{X})$, onde — por definição — a esperança de uma matriz aleatória é tomada componente-a-componente.
2. $\mathbb{E}(\mathbf{X}') = (\mathbb{E}(\mathbf{X}))'$.
3. $\sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i' = \mathbf{X}'\mathbf{X}$.
4. $\sum_{i=1}^n \mathbf{x}_i y_i = \mathbf{X}'\mathbf{y}$.
5. $\mathbb{V}(\mathbf{x}) = \mathbb{E}(\mathbf{x}\mathbf{x}') - \mathbb{E}(\mathbf{x})\mathbb{E}(\mathbf{x}')$.
6. $\mathbb{V}(\Sigma\mathbf{x} + \boldsymbol{\xi}) = \Sigma\mathbb{V}(\mathbf{x})\Sigma'$.

O modelo de regressão linear

“...avistamos de longe e dizemos, O Templo, quando entramos no Pátio dos Gentios tornamos a dizer, O Templo, e agora o carpinteiro José, apoiado à balaustrada, olha e diz, O Templo, e é ele quem tem razão...”

José Saramago, *O Evangelho Segundo Jesus Cristo*

Ao longo de todo o texto, consideramos fixado um espaço de probabilidade $(\Omega, \mathcal{A}, \mathbb{P})$ no qual estão definidas todas as variáveis aleatórias que encontraremos. Ademais, no que segue n é um número natural, $y, y_1, \dots, y_n$ e $u, u_1, \dots, u_n$ são variáveis aleatórias, $\mathbf{x}, \mathbf{x}_1, \dots, \mathbf{x}_n$ são vetores aleatórios de dimensão $d + 1$,¹⁴ onde $d \in \mathbb{N}$ etc. Por conveniência, assumiremos que a primeira componente desses vetores aleatórios é sempre uma constante e, guiados por elevado senso estético, não poderíamos tomar essa constante como sendo qualquer outro valor que não o número 1. Assim temos, por exemplo,

$$\mathbf{x}' = \begin{pmatrix} 1 & x_1 & \dots & x_d \end{pmatrix}, \quad \mathbf{x}'_i = \begin{pmatrix} 1 & x_{i1} & \dots & x_{id} \end{pmatrix}$$

e assim por diante.

Um caso simples, para começar

PODE-SE DIZER QUE lançar os olhos sobre um texto de Econometria é o mesmo que iniciar uma contagem regressiva cujo estopim se dá no momento em que o leitor se depara com a equação

$$y = \beta_0 + \beta_1 x_1 + u. \tag{1}$$

¹⁴ Isso representa um pequeno ajuste com respeito à notação introduzida no Prelúdio.

Comemorai! Já é finda a contagem! Aí está aquele que podemos chamar de “O Modelo de regressão linear” (ao menos em uma versão bem simplezinha).¹⁵ Há diversas maneiras de se chegar a uma tal igualdade, dependendo da ordem com que os ingredientes envolvidos são invocados. Podemos, por exemplo:

1. obtê-la como uma **definição**, considerando $\mathbf{x}$ e u como dados, β_0 e β_1 como constantes pré-fixadas e *definindo* y por $y := \beta_0 + \beta_1 x_1 + u$.
2. atingi-la por uma **proposição**, e.g. tomando $\mathbf{x}$ e y como dados e demonstrando que existem números reais β_0 e β_1 e uma variável aleatória u satisfazendo $\mathbb{E}(u) = 0$ e $\text{cov}(u, x_1) = 0$, tal que $y = \beta_0 + \beta_1 x_1 + u$.
3. considerar que a igualdade (1) é uma **suposição** (uma relação hipotética entre essas variáveis); isso usualmente é feito em um contexto aplicado, onde $\mathbf{x}$ e y modelam quantidades do mundo real, e o pesquisador escreve “suponha que $y = \beta_0 + \beta_1 x_1 + u$ para algum par de parâmetros reais β_0 e β_1 e alguma variável aleatória u satisfazendo isso, isso e aquilo”. É...o papel aceita tudo!

Não há muito mais o que dizer¹⁶, a não ser que podemos dar o pontapé inicial no nosso plano de tudo denotar em notação matricial e reescrever (1) como

$$\mathbf{y} = \mathbf{x}'\boldsymbol{\beta} + u. \quad (2)$$

Agora sim: O Modelo! Daqui, temos dois resultados que ilustram dois dos itens que discutimos logo acima:

Teorema 2. Considere um vetor aleatório $\mathbf{x}' = (1 \ x_1)$ e uma variável aleatória u , e sejam β_0 e β_1 números reais arbitrários mas fixados. Seja ainda y a variável aleatória *definida* pela equação

$$y := \beta_0 + \beta_1 x_1 + u.$$

Se $\mathbb{E}(u) = 0$ e $\text{cov}(x_1, u) = 0$, então

$$\text{cov}(x_1, y) = \beta_1 \mathbb{V}(x_1) \quad \text{e} \quad \mathbb{E}(y) = \beta_0 + \beta_1 \mathbb{E}(x_1).$$

¹⁵ Havíamos de começar em algum lugar, e decidimos começar por uma “amostra de tamanho 1” (por isso dispensamos o índice i) e com apenas um regressor não constante.

¹⁶ Claro, há o jargão: y é dito a **variável resposta** ou o **regressando**; $\mathbf{x}$ é dito o vetor de **covariáveis** / **preditores** / **regressores** etc; u é dito o **termo de erro**.

Teorema 3. Considere um vetor aleatório $\mathbf{x}' = (1 \ x_1)$ e uma variável aleatória y , e assuma que $\mathbb{V}(x_1) > 0$. Sejam ainda β_0 e β_1 *definidos* respectivamente por

$$\beta_1 := \frac{\text{cov}(x_1, y)}{\mathbb{V}(x_1)} \quad \text{e} \quad \beta_0 := \mathbb{E}(y) - \beta_1 \mathbb{E}(x_1).$$

Então, a variável aleatória u *definida* por

$$u := y - (\beta_0 + \beta_1 x_1)$$

satisfaz $\mathbb{E}(u) = 0$ e $\text{cov}(x_1, u) = 0$.

O primeiro dos dois teoremas acima nos diz que, se supusermos que u e x_1 são não-correlacionados então, ao definir y como no enunciado, estamos imediatamente forçando a covariância entre x_1 e y a assumir um valor bem específico — a saber, $\beta_1 \mathbb{V}(x_1)$ — e, semelhantemente, forçamos $\mathbb{E}(y)$ a assumir um valor bem específico — a saber, $\beta_0 + \beta_1 \mathbb{E}(x_1)$. Reciprocamente, o Teorema 3 nos garante que, se y e $\mathbf{x}$ são quaisquer,¹⁷ então sempre é possível representar y como uma transformação linear de $\mathbf{x}$ acrescida de um termo de erro com média zero e não-correlacionado com $\mathbf{x}$ — em certo sentido, isso é dizer que O Modelo de regressão linear está sempre corretamente especificado, e isso nos leva ao terceiro dos itens que discutimos anteriormente: a ironia de que “o papel aceita tudo” foi uma injustiça minha contra aqueles que assumem *coisas* em seus modelos, pois o Teorema 3 os autoriza a escrever a equação (1) sem apelar à necromancia ou às artes ocultas.¹⁸ Agora, um exemplo.

Exemplo 4. Esse poderia facilmente levar a pecha de ser um exemplo artificial, e a acusação não seria sem merecimento! Em qualquer caso, vamos prosseguir e esperar que algo se esclareça com ele. Considere duas variáveis aleatórias x_1 e v mutuamente independentes e com distribuições simétricas em torno de zero, e seja y a variável aleatória *definida* por

$$y := x_1^2 + v.$$

Afirmo que $\text{cov}(x_1, y) = 0$. De fato,

$$\begin{aligned} \text{cov}(x_1, y) &= \text{cov}(x_1, x_1^2 + v) \\ &= \text{cov}(x_1, x_1^2) + \text{cov}(x_1, v) \\ &= \mathbb{E}(x_1^3) = 0 \end{aligned}$$

¹⁷ Não se esqueça da suposição que fizemos — **EM VERMELHO** — no Prelúdio a essas notas.

¹⁸ Mais adiante, quando falarmos sobre endogeneidade e variáveis instrumentais, voltaremos a desconfiar que “o papel aceita tudo”.

Sendo assim, ao escrevermos $y = \beta_0 + \beta_1 x_1 + u$, com β_0 , β_1 e u definidos como no Teorema 3, realmente não estamos fazendo muita coisa, pois aqui ocorre que $\beta_1 = 0$, $\beta_0 = \mathbb{E}(y)$, $u = y - \mathbb{E}(y)$ e aí (pasmem!) $y = \mathbb{E}(y) + y - \mathbb{E}(y)$.

COM ISSO, ENCERRAMOS a apresentação d'O Modelo em sua versão bem simplezinha — isso tudo foi apenas um aquecimento, para que o leitor pudesse se sentir seguro uma última vez antes de nos aventurarmos nos mares turbulentos da Notação Matricial.

Agora temos uma amostra!

Suponhamos agora que $(x_1, y_1), \dots, (x_n, y_n)$ satisfazem as igualdades

$$y_i = \mathbf{x}_i' \boldsymbol{\beta} + u_i, \quad i \in \{1, \dots, n\}, \quad (3)$$

(O Modelo!) onde $u_1, u_2, \dots, u_n$ são termos de erro e onde $\boldsymbol{\beta}' := (\beta_0 \ \beta_1 \ \dots \ \beta_d) \in \mathbb{R}^{d+1}$. Estaremos principalmente¹⁹ interessados no caso em que os vetores aleatórios

$$(\mathbf{x}_1 \ y_1)', \dots, (\mathbf{x}_n \ y_n)' \quad (4)$$

formam uma **amostra aleatória**, isto é, no caso em que sua distribuição conjunta é da forma

$$\begin{aligned} \mathbb{P} \left\{ \bigcap_{i=1}^n [\mathbf{x}_i \in A_i, y_i \in B_i] \right\} &= \prod_{i=1}^n \mathbb{P}(\mathbf{x}_i \in A_i, y_i \in B_i) \\ &= \prod_{i=1}^n \mathbb{P}(\mathbf{x}_1 \in A_i, y_1 \in B_i), \end{aligned}$$

onde A_i e B_i , $i \in \{1, \dots, n\}$ denotam subconjuntos Borelianos de $\mathbb{R}^{d+1}$ e de $\mathbb{R}$, respectivamente. A primeira das duas igualdades acima caracteriza a **independência** entre os vetores aleatórios envolvidos; a segunda nos diz que eles são **identicamente distribuídos**. Por essa razão, quando os vetores em (4) formam uma amostra aleatória, dizemos equivalentemente que eles são **iid**.²⁰

No Prelúdio havíamos convencionado escrever (fazendo aqui uma



Figura 2: J.W. Buel, *Sea and Land*. Muitos alunos de Econometria relatam ter pesadelos em que a Notação Matricial aparece no lugar do Kraken.

¹⁹ Mas não exclusivamente!

²⁰ Fica a cargo do leitor decifrar a sigla.

ligeira correção)

$$\mathbf{y} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} \quad \text{e} \quad \mathbf{X} = \begin{pmatrix} \mathbf{x}'_1 \\ \mathbf{x}'_2 \\ \vdots \\ \mathbf{x}'_n \end{pmatrix} = \begin{pmatrix} x_{1,0} & x_{1,1} & \cdots & x_{1,d} \\ x_{2,0} & x_{2,1} & \cdots & x_{2,d} \\ \vdots & \vdots & \ddots & \vdots \\ x_{n,0} & x_{n,1} & \cdots & x_{n,d} \end{pmatrix}$$

e, acrescentando agora ao nosso estoque mais uma notação, a saber,

$$\mathbf{u}' = (u_1 \quad u_2 \quad \cdots \quad u_n),$$

podemos finalmente — apoiados à balastrada — dizer, O Modelo!, e contemplá-lo em toda sua generalidade:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}. \quad (5)$$

O leitor familiarizado com a Álgebra Linear vai se recordar que há uma dualidade entre sistemas de equações lineares e produtos de matrizes por vetores. Algo parecido ocorre aqui, e é um exercício verificar que as equações (3) e (5) são de fato equivalentes.

Os Teoremas 2 e 3 acima abordam a questão de como obter O Modelo a depender da ordem com que são acrescentados os ingredientes envolvidos. Podemos seguir o mesmo roteiro para obter a equação (5), e é isso que faremos agora.

Teorema 5. Considere uma matriz aleatória $\mathbf{X}$ de dimensões $n \times (d+1)$ e um vetor aleatório $\mathbf{u}$ em $\mathbb{R}^n$, e seja $\boldsymbol{\beta} \in \mathbb{R}^{d+1}$ um vetor arbitrário mas fixado. Seja ainda $\mathbf{y}$ o vetor aleatório *definido* pela igualdade (5). Se $\mathbb{E}(\mathbf{u}) = \mathbf{0} \in \mathbb{R}^n$ e $\mathbb{E}(\mathbf{X}'\mathbf{u}) = \mathbf{0} \in \mathbb{R}^{d+1}$, então

$$\mathbb{E}(\mathbf{y}) = \mathbb{E}(\mathbf{X})\boldsymbol{\beta} \quad \text{e} \quad \mathbb{E}(\mathbf{X}'\mathbf{X})\boldsymbol{\beta} = \mathbb{E}(\mathbf{X}'\mathbf{y}) \quad (6)$$

Teorema 6. Considere uma matriz aleatória $\mathbf{X}$ de dimensões $n \times (d+1)$ e um vetor aleatório $\mathbf{y}$ em $\mathbb{R}^n$, e assumamos que $\mathbb{E}(\mathbf{X}'\mathbf{X})$ é invertível. Seja ainda $\boldsymbol{\beta}$ *definido* por

$$\boldsymbol{\beta} = (\mathbb{E}(\mathbf{X}'\mathbf{X}))^{-1}\mathbb{E}(\mathbf{X}'\mathbf{y}). \quad (7)$$

Então o vetor aleatório $\mathbf{u}$ *definido* por

$$\mathbf{u} := \mathbf{y} - \mathbf{X}\boldsymbol{\beta}$$

satisfaz $\mathbb{E}(\mathbf{X}'\mathbf{u}) = \mathbf{0} \in \mathbb{R}^{d+1}$. Adicionalmente, se

$$\mathbf{x}'_i = (1 \quad x_{i1} \quad \cdots \quad x_{id})$$

para todo i (isto é, a primeira coluna de $\mathbf{X}$ é uma coluna de 1's), e se os vetores em (4) são identicamente distribuídos, então $\mathbb{E}(\mathbf{u}) = \mathbf{0} \in \mathbb{R}^n$.

O Teorema 6 diz que, até certo ponto, O Modelo linear com regressores exógenos²¹ está sempre “corretamente especificado”, no sentido de que exprime uma relação válida entre $\mathbf{y}$ e $\mathbf{x}$; de fato, nos diz explicitamente quem devem ser β e $\mathbf{u}$ para que tal relação valha. Uma outra maneira de se enxergar esse fato é a seguinte:

²¹ Mais sobre jargões logo adiante!

Corolário 7. Considere uma matriz aleatória $\mathbf{X}$ de dimensões $n \times (d+1)$ e um vetor aleatório $\mathbf{y}$ em $\mathbb{R}^n$, e assuma que $\mathbb{E}(\mathbf{X}'\mathbf{X})$ é invertível. Então existe um “vetor de parâmetros” $\beta \in \mathbb{R}^{d+1}$ e termos de erro $u_1, \dots, u_n$ tais que $\mathbb{E}(\mathbf{X}'\mathbf{u}) = \mathbf{0}$ e $\mathbf{y} = \mathbf{X}\beta + \mathbf{u}$.

ATÉ AQUI, fizemos considerações “teóricas”, no sentido de que estamos calculando esperanças, covariâncias etc. Isso nada tem a ver, ainda, com o método de mínimos quadrados ordinários — o qual, em particular, pode ser apresentado sem que seja feita qualquer menção a probabilidades, variáveis aleatórias e *tutti quanti*: de fato, MQO pode ser visto meramente como um método numérico para encontrar um hiperplano afim que se “ajusta de maneira ótima” a um conjunto de pontos no espaço Euclidiano. A despeito disso, a riqueza do método se revela em toda a sua magnitude precisamente quando o enxergamos sob um ponto de vista probabilístico, e é isso que faremos a seguir.

Exercícios

Quaisquer erros nos enunciados são — evidentemente — intencionais.

Exercício 4. Demonstre o Teorema 2.

Exercício 5. Demonstre o Teorema 3.

Exercício 6. Verifique que as equações (3) e (5) são equivalentes.

Exercício 7. Demonstre o Teorema 5.

Exercício 8. Mostre²² que, se a suposição na segunda parte do Teorema 6 for satisfeita (isto é, o vetor de regressores inclui uma constante e os dados são identicamente distribuídos), então

²² Dica: use o Exercício 3.

$$\beta = (\mathbb{E}(\mathbf{x}_1 \mathbf{x}_1'))^{-1} \mathbb{E}(\mathbf{x}_1 y_1).$$

Ademais, utilize a igualdade acima para resolver explicitamente a expressão para β no caso particular em que $d = 1$ — isto é, no caso em que $\mathbf{x}_1' = (1 \quad x_{1,1})$.

Exercício 9. Demonstre o Teorema 6.

Exercício 10. Seja $\mathbf{x}$ um vetor aleatório de dimensão d .

1. Mostre que $\mathbb{V}(\mathbf{x})$ é positiva semidefinida (isto é, a forma quadrática $\boldsymbol{\xi}' \mathbb{V}(\mathbf{x}) \boldsymbol{\xi}$ é não-negativa para todo $\boldsymbol{\xi} \in \mathbb{R}^d$).
2. Mostre que $\mathbb{V}(\mathbf{x})$ é singular (isto é, não-invertível) se, e somente se, existem um vetor (não aleatório) $\boldsymbol{\zeta} \in \mathbb{R}^d$ e uma constante $\eta \in \mathbb{R}$ tais que

$$\mathbb{P}(\boldsymbol{\zeta}' \mathbf{x} = \eta) = 1.$$

Exercício 11. Seja $\mathbf{X}$ uma matriz aleatória de dimensões $n \times (d + 1)$. Encontre condições que garantam a invertibilidade de $\mathbb{E}(\mathbf{X}' \mathbf{X})$.

Esperança condicional e modelos de regressão: um interlúdio

“Ó, vós que entraís, abandonai toda a esperança!”

Dante Alighieri, *Divina Commedia*

Suponha que $\mathbf{y}$, $\mathbf{X}$ e $\mathbf{u}$ satisfazem²³ a equação

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}, \quad (8)$$

para algum vetor $\boldsymbol{\beta} \in \mathbb{R}^{d+1}$. Se nada mais for dito, não há muito por onde seguir, então vamos acrescentar à supracitada suposição as hipóteses de que $\mathbb{E}(\mathbf{u}) = \mathbf{0}$ e $\mathbb{E}(\mathbf{X}'\mathbf{u}) = \mathbf{0}$. Com isso, podemos apelar ao Teorema 5 e concluir que, necessariamente, vale que

$$\boldsymbol{\beta} = (\mathbb{E}(\mathbf{X}'\mathbf{X}))^{-1}\mathbb{E}(\mathbf{X}'\mathbf{y}). \quad (9)$$

Um erro de interpretação bastante comum é tomar, nas condições acima, como verdadeira a sentença que afirma: “ $\mathbf{X}\boldsymbol{\beta}$ é a esperança condicional de $\mathbf{y}$ dado $\mathbf{X}$ ”. Ocorre que essa sentença é, em geral, falsa: essa é uma das conclusões do Teorema 10 logo abaixo. Os exercícios 14 e 22 ilustram casos em que $\mathbb{E}(\mathbf{X}'\mathbf{u}) = \mathbf{0}$ mas $\mathbb{E}(\mathbf{y} | \mathbf{X}) \neq \mathbf{X}\boldsymbol{\beta}$.

NA LITERATURA DE ECONOMETRIA, convencionou-se chamar de **exogeneidade** à propriedade $\mathbb{E}(\mathbf{X}'\mathbf{u}) = \mathbf{0}$, e de **exogeneidade estrita** à propriedade $\mathbb{E}(\mathbf{u} | \mathbf{X}) \stackrel{q.c.}{=} \mathbf{0}$. No primeiro caso dizemos também que $\mathbf{X}$ é um regressor **exógeno** e, no segundo, que é **estritamente exógeno**.²⁴ Afortunadamente, o adjetivo “estrito”



Dante Alighieri, *Apoteosi a Firenze*, Giovanni Guida

²³ Claro, a essa altura já sabemos (ou deveríamos saber) de cor e salteado o que os símbolos $\mathbf{y}$, $\mathbf{X}$ e $\mathbf{u}$ representam.

²⁴ As coisas se tornam mais simples quando estamos em um cenário iid, pois aí temos $\mathbb{E}(\mathbf{X}'\mathbf{u}) = n\mathbb{E}(\mathbf{x}_1 u_1)$. Nesse caso, usamos o jargão recém introduzido para referirmo-nos aos vetores $\mathbf{x}_i$'s.

aqui não é mero jogo de palavras, já que temos o seguinte:

Teorema 8. Se $\mathbf{X}$ é estritamente exógeno, então $\mathbf{X}$ é exógeno.

Em vista desse resultado, quando nos deparamos com O Modelo linear (8) e desejamos interpretá-lo como uma relação que descreve a *esperança condicional* de $\mathbf{y}$ dado $\mathbf{X}$ como uma transformação linear de $\mathbf{X}$, necessitamos forçosamente *supor* que $\mathbf{X}$ é estritamente exógeno.²⁵

Projeções: uma primeira incursão

Um **espaço de Hilbert** é uma estrutura matemática (um conjunto) dotada de propriedades algébricas, topológicas e geométricas semelhantes àquelas que são próprias dos espaços Euclidianos. De fato, para todo $d \in \mathbb{N}$ tem-se que $\mathbb{R}^d$ (munido do produto escalar usual) é um espaço de Hilbert — o interessante é que há espaços de Hilbert de dimensão infinita!

Antes de definir o que é um espaço de Hilbert, precisamos da noção de um **produto interno** em um espaço vetorial²⁶ H , que nada mais é do que uma função $\rho: H \times H \rightarrow \mathbb{R}$ satisfazendo as seguintes condições, para quaisquer $h, f, g \in H$ e $\xi \in \mathbb{R}$:

$$PI1. \quad \rho(h, f) = \rho(f, h)$$

$$PI2. \quad \rho(\xi h, f) = \xi \rho(h, f)$$

$$PI3. \quad \rho(h + f, g) = \rho(h, g) + \rho(f, g)$$

$$PI4. \quad \rho(h, h) \geq 0$$

$$PI5. \quad \text{Se } \rho(h, h) = 0 \text{ então } h = 0 \in H.$$

A partir de um produto interno ρ em um espaço vetorial H , é possível introduzir uma noção de *comprimento* de um elemento $h \in H$: escrevemos

$$\|h\|_\rho = \sqrt{\rho(h, h)}$$

e chamamos essa quantidade a **norma** de h (induzida por ρ). Com essa notação, estamos agora aptos a definir espaços de Hilbert, e a definição é a seguinte: dados um espaço vetorial H e um produto interno ρ nesse espaço, chamamos ao par (H, ρ) um **espaço de Hilbert** se valer a seguinte condição:²⁷

$H1.$ Se $h_0, h_1, h_2, \dots$ é uma sequência de elementos de H tal que

$$\sum_{k=1}^{\infty} \|h_k - h_{k-1}\|_\rho < \infty,$$

²⁵ Isso é particularmente importante em aplicações, onde é costumeiro dar uma interpretação causal ao modelo linear, tomando o coeficiente β_j como o *efeito* parcial do regressor x_{1j} na resposta y_1 .

²⁶ Tomaremos, por simplicidade, espaços vetoriais sobre o corpo dos números reais.

²⁷ A propriedade $H1$ exprime o atributo de *completude* do espaço H : ela significa que, nesse espaço, sequências cujos termos sucessivos tornam-se arbitrariamente próximos uns dos outros necessariamente aproximam-se de algum elemento de H (em certo sentido, isso significa que “não faltam pontos em H ” — pense no conjunto $\mathbb{Q}$ dos números racionais e em como podemos aproximar o número irracional $\sqrt{2}$ por elementos de $\mathbb{Q}$).

então existe um elemento $h \in H$ tal que $\lim_{n \rightarrow \infty} \|h - h_n\|_\rho = 0$.

Exemplo 9. Considere o conjunto H assim definido: um objeto $\mathbf{x}$ pertence a H se, e somente se, $\mathbf{x}$ é um vetor aleatório de dimensão d tal que $\mathbb{E}(\mathbf{x}'\mathbf{x}) < \infty$. Está claro que H é um espaço vetorial. Vamos então definir uma função $\rho: H \times H \rightarrow \mathbb{R}$, pela expressão

$$\rho(\mathbf{x}, \mathbf{z}) = \mathbb{E}(\mathbf{x}'\mathbf{z}), \quad \mathbf{x}, \mathbf{z} \in H. \quad (10)$$

Não é difícil de verificar que ρ assim definida satisfaz as condições *PI1* a *PI4* acima. Quanto a *PI5*, note que se $\mathbb{P}[\mathbf{x} = \mathbf{0}] = 1$, então $\|\mathbf{x}\|_\rho = 0$ e, reciprocamente, se $\|\mathbf{x}\|_\rho = 0$ então necessariamente vale que $\mathbb{P}[\mathbf{x} = \mathbf{0}] = 1$. Por essa razão, identificamos²⁸ elementos de H que sejam iguais com probabilidade 1. Um fato importante (e, infelizmente, aqui não vamos prová-lo) é que o espaço H é um espaço de Hilbert e, exceto no caso em que Ω é um conjunto finito, tipicamente tem dimensão infinita. Aproveitando a deixa, vamos introduzir uma notação e escrever $H =: L^2_{\mathbb{P}}(\Omega; \mathbb{R}^d)$.

PRODUTOS INTERNOS TAMBÉM SÃO ÚTEIS, entre outros motivos, por permitirem que em H sejam introduzidas noções de ângulos entre elementos — de particular interesse para nós é a noção de ortogonalidade: dizemos que $h, f \in H$ são **ortogonais** se $\rho(h, f) = 0$. Espaços de Hilbert possuem a importante propriedade de serem *proximais*, o que quer dizer o seguinte: se $h \in H$ está dado, e se $M \subseteq H$ é um subespaço fechado²⁹ de H , então existe um único elemento $h_M \in M$ tal que

$$\|h - h_M\|_\rho \leq \|h - f\|_\rho, \quad (11)$$

para qualquer $f \in M$. Em particular, a desigualdade acima vale estritamente se tomarmos $f \neq h_M$. O elemento h_M é dito a **projeção ortogonal** de h em M , e vale a decomposição

$$\rho(h_M, h - h_M) = 0, \quad (12)$$

ou seja, h_M e $h - h_M$ são ortogonais entre si. Escreveremos $M^\perp$ para denotar o **complemento ortogonal de M** , definido como o conjunto de todos os elementos de H que são ortogonais a M — isto é, $f \in M^\perp$ se, e somente se, $\rho(f, g) = 0$ para todo $g \in M$. Note que a desigualdade (11) nos diz que h_M é o elemento de M que está o mais próximo possível de h .

²⁸ A matemática aqui tem lá suas nuances: o que estamos fazendo — ávidos por obter a propriedade *PI5* — é o que se chama de *tomar o quociente* do espaço H por uma relação de equivalência, no caso a relação em que dois elementos $\mathbf{x}$ e $\mathbf{z}$ de H são ditos equivalentes se

$$\mathbb{P}[\mathbf{x} = \mathbf{z}] = 1.$$

²⁹ M é dito fechado quando vale a seguinte propriedade: se $h_1, h_2, \dots$ é uma sequência convergente de elementos de M , então o limite é também um elemento de M .

Um caso de particular interesse para nós ocorre quando o subespaço M discutido acima tem dimensão finita (digamos, $\dim(M) = d + 1$); se esse é o caso, então existem elementos $f_0, f_1, \dots, f_d \in M$ tais que o conjunto por eles formado é linearmente independente e tal que todo elemento $f \in M$ se escreve unicamente na forma

$$f = a_0 f_0 + a_1 f_1 + \dots + a_d f_d, \quad (13)$$

para alguma $(d + 1)$ -upla de números reais $a_0, \dots, a_d$. Reciprocamente, se $f_0, f_1, \dots, f_d$ são elementos linearmente independentes em H , então o conjunto formado por todos os elementos $f \in H$ que podem ser expressos na forma (13) tem dimensão finita e igual a $d + 1$.³⁰

Em suma, um subespaço $M \subseteq H$ tem dimensão finita (e igual a $d + 1$) se, e somente se, $M = \text{span}(f_0, \dots, f_d)$, para algum $d \in \mathbb{N}$ e alguma $(d + 1)$ -upla $f_0, \dots, f_d$ de elementos linearmente independentes de H . Se esse é o caso, então $h_M = \sum_{j=0}^d a_j f_j$ para alguma única $(d + 1)$ -upla $a_0, \dots, a_d \in \mathbb{R}$, e a equação (11) nos diz que

$$\left\| h - \sum_{j=0}^d a_j f_j \right\|_\rho \leq \left\| h - \sum_{j=0}^d b_j f_j \right\|_\rho,$$

para qualquer $\mathbf{b} = (b_0 \ b_1 \ \dots \ b_d)' \in \mathbb{R}^{d+1}$. Concluimos, então, que o vetor $\mathbf{a} = (a_0 \ a_1 \ \dots \ a_d)' \in \mathbb{R}^{d+1}$ é a única solução para o problema de minimização

$$\min_{\mathbf{b} \in \mathbb{R}^{d+1}} \left\| h - \sum_{j=0}^d b_j f_j \right\|_\rho$$

Agora que sabemos algo sobre projeções ortogonais, podemos enunciar o seguinte importantíssimo resultado.

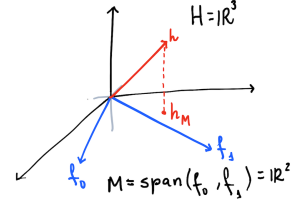


Figura 3: Projeção ortogonal de h no subespaço gerado por f_0 e f_1 .

³⁰ Denotamos tal conjunto por $\text{span}(f_0, \dots, f_d)$ e dizemos que esse é o **espaço gerado** por $f_0, \dots, f_d$.

Teorema 10. Sejam

$$\mathbf{y}' = (y_1 \quad \dots \quad y_n) \quad e \quad \mathbf{x}'_i = (x_{i,0} \quad x_{i,1} \quad \dots \quad x_{i,d}),$$

onde $i \in \{1, \dots, n\}$, vetores aleatórios de dimensões n e $d+1$, respectivamente, onde $n \geq d+1$. Seja ainda

$$\mathbf{X} = \begin{pmatrix} x_{1,0} & x_{1,1} & \dots & x_{1,d} \\ \vdots & \vdots & \ddots & \vdots \\ x_{n,0} & x_{n,1} & \dots & x_{n,d} \end{pmatrix}. \quad (14)$$

Vale o seguinte:

1. No espaço $L^2_{\mathbb{P}}(\Omega; \mathbb{R}^1)$ com o produto interno

$$\rho(z, w) = \mathbb{E}(zw), \quad z, w \in L^2_{\mathbb{P}}(\Omega; \mathbb{R}^1)$$

- (a) a esperança condicional $\mathbb{E}(y_1 | \mathbf{x}_1)$ é a projecção ortogonal da variável aleatória y_1 no subespaço M constituído por todas as variáveis aleatórias z que podem ser representadas na forma $z = g \circ \mathbf{x}_1$ para alguma função $g: \mathbb{R}^{d+1} \rightarrow \mathbb{R}^1$.
- (b) se $\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1)$ é invertível, então o elemento $\mathbf{x}'_1 \boldsymbol{\beta}$, com $\boldsymbol{\beta} := (\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1))^{-1} \mathbb{E}(\mathbf{x}_1 y_1)$, é a projecção ortogonal de y_1 no subespaço $\text{span}(x_{1,0}, \dots, x_{1,d})$.

2. No espaço $L^2_{\mathbb{P}}(\Omega; \mathbb{R}^n)$ com o produto interno

$$\rho(\mathbf{z}, \mathbf{w}) = \mathbb{E}(\mathbf{z}' \mathbf{w}), \quad \mathbf{z}, \mathbf{w} \in L^2_{\mathbb{P}}(\Omega; \mathbb{R}^n),$$

- (a) a esperança condicional $\mathbb{E}(\mathbf{y} | \mathbf{X})$ é a projecção ortogonal do vetor aleatório $\mathbf{y}$ no subespaço M constituído por todos os vetores aleatórios $\mathbf{z}$ que podem ser representados na forma $\mathbf{z} = g \circ \mathbf{X}$ para alguma função $g: \mathbb{R}^{n \times (d+1)} \rightarrow \mathbb{R}^n$.
- (b) se $\mathbb{E}(\mathbf{X}' \mathbf{X})$ é invertível, então o elemento $\mathbf{X} \boldsymbol{\beta}$, com $\boldsymbol{\beta} := (\mathbb{E}(\mathbf{X}' \mathbf{X}))^{-1} \mathbb{E}(\mathbf{X}' \mathbf{y})$, é a projecção ortogonal de $\mathbf{y}$ no subespaço gerado pelas $d+1$ colunas de $\mathbf{X}$.

3. No espaço $\mathbb{R}^n$ com o produto escalar usual, se $\mathbf{X}' \mathbf{X}$ é invertível então o elemento $\mathbf{X} \hat{\boldsymbol{\beta}}$, onde $\hat{\boldsymbol{\beta}} = (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbf{y}$, é a projecção ortogonal do vetor $\mathbf{y}$ no subespaço gerado pelas $d+1$ colunas de $\mathbf{X}$.

Demonstração. Vamos demonstrar o item (2a). O item (1a) segue diretamente tomando-se $n = 1$. As demonstrações dos itens (1b), (2b) e (3) são exercícios.

Seja $\varphi: \mathbb{R}^{n \times (d+1)} \rightarrow \mathbb{R}^n$ a função de regressão de $\mathbf{y}$ em $\mathbf{X}$, quer dizer $\varphi \circ \mathbf{X} = \mathbb{E}(\mathbf{y} | \mathbf{X})$, e seja $g: \mathbb{R}^{n \times (d+1)} \rightarrow \mathbb{R}^n$ uma função mensurável qualquer satisfazendo $\mathbb{E}(|g \circ \mathbf{X}|) < \infty$, onde $|\cdot|$

denota a norma Euclidiana em $\mathbb{R}^n$, $|\boldsymbol{\xi}| := \sqrt{\boldsymbol{\xi}'\boldsymbol{\xi}}$. Escrevendo $\mathbf{z} = \varphi \circ \mathbf{X}$ e $\mathbf{w} = g \circ \mathbf{X}$, temos

$$\begin{aligned}\|\mathbf{y} - g \circ \mathbf{X}\|_\rho^2 &= \mathbb{E}\{|\mathbf{y} - g \circ \mathbf{X}|^2\} \\ &= \sum_{i=1}^n \mathbb{E}\{(y_i - w_i)^2\} \\ &= \sum_{i=1}^n \mathbb{E}\{((y_i - z_i) + (z_i - w_i))^2\} \\ &= \sum_{i=1}^n \mathbb{E}\{(y_i - z_i)^2 + 2(y_i - z_i)(z_i - w_i) + (z_i - w_i)^2\}.\end{aligned}$$

Distribuindo, por linearidade, o operador de esperança e notando que

$$\begin{aligned}\mathbb{E}\{(y_i - z_i)(z_i - w_i)\} &= \mathbb{E}\{\mathbb{E}[(y_i - z_i)(z_i - w_i) | \mathbf{X}]\} \\ &= \mathbb{E}\{(z_i - w_i)\mathbb{E}[(y_i - z_i) | \mathbf{X}]\} \\ &= 0,\end{aligned}$$

ficamos com

$$\begin{aligned}\|\mathbf{y} - g \circ \mathbf{X}\|_\rho^2 &= \sum_{i=1}^n \mathbb{E}\{(y_i - z_i)^2\} + \sum_{i=1}^n \mathbb{E}\{(z_i - w_i)^2\} \\ &= \|\mathbf{y} - \varphi \circ \mathbf{X}\|_\rho^2 + \|\varphi \circ \mathbf{X} - g \circ \mathbf{X}\|_\rho^2.\end{aligned}$$

O primeiro termo no lado direito da igualdade acima independe de g ; o segundo é não negativo e se anula precisamente e tão somente quando tomamos g tal que $\mathbb{P}[g \circ \mathbf{X} = \varphi \circ \mathbf{X}] = 1$. Isso completa a demonstração. $\blacksquare$

Remark 11. Outra maneira de enunciar, digamos, o item (1b) do teorema acima é a seguinte: de todas as variáveis aleatórias z que podem ser escritas na forma

$$z = b_0 x_{1,0} + \cdots + b_d x_{1,d} = \mathbf{x}'_1 \mathbf{b},$$

para algum $\mathbf{b} \in \mathbb{R}^{d+1}$, temos que $\mathbf{x}'_1 \boldsymbol{\beta}$, com $\boldsymbol{\beta} := (\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1))^{-1} \mathbb{E}(\mathbf{x}_1 y_1)$, é aquela que está mais próxima de y_1 no sentido de que

$$\mathbb{E}\{(y_1 - \mathbf{x}'_1 \boldsymbol{\beta})^2\} \leq \mathbb{E}\{(y_1 - \mathbf{x}'_1 \mathbf{b})^2\}$$

para todo $\mathbf{b} \in \mathbb{R}^{d+1}$, em que a desigualdade vale estritamente se tomarmos $\mathbf{b} \neq \boldsymbol{\beta}$. Para fins mnemônicos, convém reescrever cada um dos demais itens do teorema de forma semelhante ao que fizemos aqui.

Exercícios

Quaisquer erros nos enunciados são — é claro! — intencionais.

Exercício 12. Nesse capítulo utilizamos a notação $\mathbb{E}(\mathbf{y} | \mathbf{X})$ sem fornecer uma definição. Ei-la: sejam $\mathbf{y}$ um vetor aleatório de dimensão n , sejam $\mathbf{x}_i$ vetores aleatórios de dimensão $d + 1$, onde $i \in \{1, \dots, n\}$, e seja $\mathbf{X}$ definida como em (14). A **esperança condicional** de $\mathbf{y}$ dado $\mathbf{X}$ (equivalentemente, dados $\mathbf{x}_1, \dots, \mathbf{x}_n$) é o vetor aleatório $\mathbb{E}(\mathbf{y} | \mathbf{X}) \equiv \mathbb{E}(\mathbf{y} | \mathbf{x}_1, \dots, \mathbf{x}_n)$ de dimensão n que satisfaz as seguintes condições:

EC1. $\mathbb{E}(\mathbf{y} | \mathbf{X}) = \varphi \circ \mathbf{X}$, para alguma função mensurável $\varphi: \mathbb{R}^{n \times (d+1)} \rightarrow \mathbb{R}^n$, com $\mathbb{E}(|\varphi \circ \mathbf{X}|) < \infty$.

EC2. Vale a igualdade $\mathbb{E}(\mathbb{E}(\mathbf{y} | \mathbf{X}) \mathbb{I}_{[\mathbf{X} \in B]}) = \mathbb{E}(\mathbf{y} \mathbb{I}_{[\mathbf{X} \in B]})$ para todo subconjunto Boreliano $B \subseteq \mathbb{R}^{n \times (d+1)}$.

1. Observando que qualquer função $g: \mathbb{R}^m \rightarrow \mathbb{R}^k$ é da forma $g = (g_1, \dots, g_k)$, onde $g_i: \mathbb{R}^m \rightarrow \mathbb{R}$, verifique que

$$[\mathbb{E}(\mathbf{y} | \mathbf{X})]_i = \mathbb{E}(y_i | \mathbf{X}), \quad i \in \{1, \dots, n\}$$

2. Mostre que, se z é uma variável aleatória da forma $z = g \circ \mathbf{X}$ para alguma função g de $\mathbb{R}^{n \times (d+1)}$ em $\mathbb{R}$, então $\mathbb{E}(z\mathbf{y} | \mathbf{X}) = z\mathbb{E}(\mathbf{y} | \mathbf{X})$.

Exercício 13. Mostre que, se os vetores aleatórios

$$(\mathbf{x}_1 \ y_1)', \dots, (\mathbf{x}_n \ y_n)'$$

formam uma amostra aleatória, então $[\mathbb{E}(\mathbf{y} | \mathbf{X})]_i = \mathbb{E}(y_i | \mathbf{x}_i)$.

Exercício 14. Suponha que $x_{1,1}, \dots, x_{n,1}$ e $v_1, \dots, v_n$ são uma amostra aleatória de x_1 e v , respectivamente, onde x_1 e v são como no exercício 4: mutuamente independentes e, cada um, simétrico em torno de zero. Para $i \in \{1, \dots, n\}$, defina $y_i = x_{i,1}^2 + v_i$, e ponha

$$\mathbf{X} = \begin{pmatrix} 1 & x_{1,1} \\ \vdots & \vdots \\ 1 & x_{n,1} \end{pmatrix}.$$

Sendo β definido como na equação (9), e pondo $\mathbf{u} = \mathbf{y} - \mathbf{X}\beta$, já sabemos que $\mathbb{E}(\mathbf{X}'\mathbf{u}) = 0$ e, como as variáveis aleatórias aqui são identicamente distribuídas, sabemos de fato que $\mathbb{E}(\mathbf{u}) = 0$.

1. Mostre que $\mathbb{E}(\mathbf{u} | \mathbf{X}) \neq 0$ e conclua que $\mathbb{E}(\mathbf{y} | \mathbf{X}) \neq \mathbf{X}\boldsymbol{\beta}$.
2. Defina agora, para $i \in \{1, \dots, n\}$, $x_{i2} = x_{i1}^2$ e ponha

$$\mathbf{W} := \begin{pmatrix} 1 & x_{1,1} & x_{1,2} \\ \vdots & \vdots & \vdots \\ 1 & x_{n,1} & x_{n,2} \end{pmatrix}, \quad \boldsymbol{\gamma} := (\mathbb{E}(\mathbf{W}'\mathbf{W}))^{-1}\mathbb{E}(\mathbf{W}'\mathbf{y}).$$

Mostre que $\boldsymbol{\gamma}' = (0 \quad 0 \quad 1)$ e $\mathbb{E}(\mathbf{y} - \mathbf{W}\boldsymbol{\gamma} | \mathbf{W}) = \mathbf{0}$.

Exercício 15. Demonstre o Teorema 8.

Exercício 16. Seja $M(n \times d)$ o conjunto formado por todas as matrizes $n \times d$, e considere a transformação $\tau: M(n \times d) \rightarrow \mathbb{R}^{n \times d}$ que “empilha” as colunas de uma matriz $A \in M(n \times d)$ para obter um vetor em $\mathbb{R}^{n \times d}$. Verifique que τ é linear, bijetora, contínua e que τ^{-1} é também contínua.³¹ Esse exercício justifica a notação utilizada para o domínio da função de regressão de $\mathbf{y}$ em $\mathbf{X}$.

³¹ Sobre continuidade: é preciso verificar que, se $A_k \rightarrow A$ em $M(n \times d)$, então $\tau A_k \rightarrow \tau A$ em $\mathbb{R}^{n \times d}$ etc. Em ambos os casos, convergência significa convergência componente-a-componente.

Exercício 17. Considere uma matriz $A \in M(2 \times 2)$,

$$A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix},$$

e seja $\mathbf{a} = \tau A \in \mathbb{R}^4$, onde τ é definida como no exercício acima. Sejam ainda $\boldsymbol{\beta} \in \mathbb{R}^2$ e $B \in M(2 \times 4)$. Comparando os vetores $A\boldsymbol{\beta}$ e $B\mathbf{a}$, conclua que nem toda transformação linear $T: M(2 \times 2) \rightarrow \mathbb{R}^2$ é da forma $A \mapsto A\boldsymbol{\beta}$ para algum $\boldsymbol{\beta} \in \mathbb{R}^2$.

Exercício 18. Verifique que o plano Euclidiano $\mathbb{R}^2$, munido do produto escalar usual, é um espaço de Hilbert.

Exercício 19. Mostre que o espaço $L^2_{\mathbb{P}}(\Omega; \mathbb{R}^1)$ coincide com o conjunto formado por todas as variáveis aleatórias com variâncias finitas.

Exercício 20. Dado um subespaço fechado M em um espaço de Hilbert H , mostre que todo elemento $h \in H$ se exprime de maneira única como $h = h_M + h_{M^\perp}$.

Exercício 21. Demonstre os itens (1b), (2b) e (3) do Teorema 10. Dica: comece pelo item (3), notando que a função real L definida,

para $\mathbf{b} \in \mathbb{R}^{d+1}$, por

$$L(\mathbf{b}) := |\mathbf{y} - \mathbf{X}\mathbf{b}|^2 = \sum_{i=1}^n (y_i - [\mathbf{X}\mathbf{b}]_i)^2,$$

é uma função convexa. Isso permite que encontremos o valor $\hat{\boldsymbol{\beta}} \in \mathbb{R}^{d+1}$ que minimiza globalmente a função L resolvendo o sistema

$$\nabla L(\hat{\boldsymbol{\beta}}) = \mathbf{0},$$

onde ∇L denota o gradiente de L , $[\nabla L]_j := \partial L / \partial b_j$, $j \in \{0, \dots, d\}$.

Para o item (2b), basta seguir um roteiro semelhante ao acima, agora definindo $L(\mathbf{b}) = \mathbb{E}\{|\mathbf{y} - \mathbf{X}\mathbf{b}|^2\}$ para $\mathbf{b} \in \mathbb{R}^{d+1}$, onde $|\cdot|$ denota a norma Euclidiana, e resolver a equação $\nabla L(\boldsymbol{\beta}) = \mathbf{0}$. Sinta-se autorizado a fazer a passagem $\nabla \mathbb{E} = \mathbb{E} \nabla$.

Exercício 22. Seja $(\Omega, \mathcal{A}, \mathbb{P})$ o espaço de probabilidade em que $\Omega = (0, 1]$, $\mathcal{A} = \text{'}\sigma\text{-álgebra de Borel em } (0, 1]\text{'}$ e $\mathbb{P} = \text{'medida de Lebesgue em } (0, 1]\text{'}$ (isto é, $\mathbb{P}(a, b] = b - a$ para qualquer intervalo $(a, b] \subseteq (0, 1]$, $0 \leq a < b \leq 1$). Seja x a variável aleatória definida via $x(\omega) := \omega$, $\omega \in \Omega$, e seja y uma variável aleatória integrável qualquer.

1. Mostre que $\mathbb{E}(y | x) = y$.
2. Sendo $M = \text{span}(x) \subseteq L^2_{\mathbb{P}}(\Omega; \mathbb{R})$, mostre que $y_M = 3x \mathbb{E}(xy)$.

O método de mínimos quadrados

“We all believe fairy-tales, and live in them. Some, with a sumptuous literary turn, believe in the existence of the lady clothed with the sun. Some, with a more rustic, elvish instinct, believe merely in the impossible sun itself.”

G. K. Chesterton, Heretics

A presente seção é um *interim*, durante o qual poderemos esquecer — provisoriamente — que os objetos com os quais estamos lidando são variáveis (e vetores e matrizes) aleatórias. Sendo assim, por ora $\mathbf{y}$ denota um vetor de $\mathbb{R}^n$, $\mathbf{x}_1, \dots, \mathbf{x}_n$ vetores de $\mathbb{R}^{d+1}$ e $\mathbf{X}$ uma matriz $n \times (d+1)$ cuja componente i, j é $x_{ij} = j$ ésima componente de $\mathbf{x}_i$. Tipicamente (embora não necessariamente) convencionou-se que $x_{i0} = 1$ para todo i .

O MÉTODO DOS MÍNIMOS QUADRADOS APARECE PELA PRIMEIRA VEZ EM 1805, em um apêndice intitulado *Sur la Méthode des moindres quarrés*, na *Nouvelles méthodes pour la détermination des orbites des comètes* de Adrien-Marie Legendre. Em 1809, o mesmo método aparece na *Theoria Motus Corporum Coelestium in Sectionibus Conicis Solem Ambientium* de Carl Friedrich Gauss. Tudo indica que Gauss e Legendre “descobriram” o método de forma independente, e há uma contenda a respeito de quem, de fato, foi o primeiro descobridor (foi Gauss).³² O referido método pode ser subdividido em duas categorias, a saber, mínimos quadrados *ordinários* e mínimos quadrados *não-lineares*, e por ora estaremos interessados na primeira delas. Aqui, o termo “ordinários” é empregado em contraposição ao termo “não-lineares” e, portanto, um nome mais adequado (ou ao menos mais explícito) para a primeira categoria seria *mínimos quadrados lineares*. Como veremos, *linea-*

APPENDICE.

Sur la Méthode des moindres quarrés.

Dans la plupart des questions où il s'agit de tirer des mesures données par l'observation, les résultats les plus exacts qu'elles peuvent offrir, on est presque toujours conduit à un système d'équations de la forme

$$E = a + bx + cy + fz + Rc.$$

dans lesquelles $a, b, c, f, Rc.$ sont des coefficients connus, qui varient d'une équation à l'autre, et $x, y, z, Rc.$ sont des inconnues qu'il faut déterminer par la condition que la valeur de E se réduise, pour chaque équation, à une quantité ou nulle ou très-petite.

Si l'on a autant d'équations que d'inconnues $x, y, z, Rc.$, il n'y a aucune difficulté pour la détermination de ces inconnues, et on peut rendre les erreurs E absolument nulles. Mais le plus souvent, le nombre des équations est supérieur à celui des inconnues, et il est impossible d'annuler toutes les erreurs.

Dans cette circonstance, qui est celle de la plupart des problèmes physiques et astronomiques, où l'on cherche à déterminer quelques éléments importants, il entre nécessairement de l'arbitraire dans la distribution des erreurs, et on ne doit pas s'attendre que toutes les hypothèses conduisent exactement aux mêmes résultats; mais il faut sur-tout faire en sorte que les erreurs extrêmes, sans avoir égard à leurs signes, soient renfermées dans les limites les plus étroites qu'il est possible.

De tous les principes qu'on peut proposer pour cet objet, je pense qu'il n'en est pas de plus général, de plus exact, ni d'une application plus facile que celui dont nous avons fait usage dans les recherches précédentes, et qui consiste à rendre

A primeira página do apêndice *Sur la Méthode des moindres quarrés*. Legendre vai direto ao ponto: “Dans la plupart des questions où il s'agit de tirer des mesures données par l'observation, les résultats les plus exacts qu'elles peuvent offrir, on est presque toujours conduit à un système d'équations...”

³² R.L. Plackett. The discovery of the method of least squares. *Biometrika*, 59(2), 1972

ridade e não-linearidade referem-se à maneira como os *parâmetros* se apresentam nos sistemas de equações estudados.³³

³³ Veja os exercícios 32 e 33.

O método

Sem mais delongas, consideremos como dados n pontos³⁴ no espaço Euclidiano $\mathbb{R}^{d+2}$:

$$\begin{pmatrix} y_1 \\ x_{1,0} \\ \vdots \\ x_{1,d} \end{pmatrix}, \begin{pmatrix} y_2 \\ x_{2,0} \\ \vdots \\ x_{2,d} \end{pmatrix}, \dots, \begin{pmatrix} y_n \\ x_{n,0} \\ \vdots \\ x_{n,d} \end{pmatrix}.$$

Para cada vetor $\mathbf{b} = (b_0 \ b_1 \ \dots \ b_d)' \in \mathbb{R}^{d+1}$, podemos escrever um sistema de n equações lineares,

$$\begin{aligned} y_1 &= b_0 x_{1,0} + b_1 x_{1,1} + \dots + b_d x_{1,d} + v_1(\mathbf{b}) \\ y_2 &= b_0 x_{2,0} + b_1 x_{2,1} + \dots + b_d x_{2,d} + v_2(\mathbf{b}) \\ &\vdots \\ y_n &= b_0 x_{n,0} + b_1 x_{n,1} + \dots + b_d x_{n,d} + v_n(\mathbf{b}) \end{aligned} \quad (15)$$

onde os termos $v_i(\mathbf{b})$, $i \in \{1, \dots, n\}$, estão definidos, por tautologia, como

$$v_i(\mathbf{b}) := y_i - \mathbf{x}_i' \mathbf{b}, \quad i \in \{1, \dots, n\}.$$

Visto que estamos considerando $n > d + 1$, ocorre que — exceto em casos mui particulares — não é possível encontrar um vetor $\mathbf{b}^* \in \mathbb{R}^{d+1}$ tal que valha $v_i(\mathbf{b}^*) = 0$ para todo i . Como prosseguir, então? Ora, eis a questão: tentemos encontrar um vetor $\hat{\beta} \in \mathbb{R}^{d+1}$ que torne todos os **resíduos** $\hat{u}_i := v_i(\hat{\beta})$, $i \in \{1, \dots, n\}$, “tão pequenos quanto for possível”. Claro, há infinitas maneiras de se dar sentido à expressão “tão pequenos quanto for possível” — Boscovich e Laplace que o digam! — e a grande sacada de Legendre e Gauss foi considerar que o vetor $\hat{\beta}$ desejado deve ser aquele que soluciona o problema de se minimizar, com respeito a $\mathbf{b} \in \mathbb{R}^{d+1}$, a expressão

$$\sum_{i=1}^n v_i(\mathbf{b})^2 \equiv \sum_{i=1}^n (y_i - \mathbf{x}_i' \mathbf{b})^2 \equiv (\mathbf{y} - \mathbf{X}\mathbf{b})'(\mathbf{y} - \mathbf{X}\mathbf{b}) \equiv |\mathbf{y} - \mathbf{X}\mathbf{b}|^2,$$

onde $|\cdot|$ aqui denota a norma Euclidiana em $\mathbb{R}^n$. Isto é,

$$\hat{\beta} := \arg \min_{\mathbf{b} \in \mathbb{R}^{d+1}} |\mathbf{y} - \mathbf{X}\mathbf{b}|^2. \quad (16)$$

³⁴ Em que pese a crescente importância de se considerar situações onde $n < d + 1$ — por exemplo, na literatura de regularização e seleção de variáveis — por ora vamos nos restringir ao tradicionalíssimo cenário em que $n > d + 1$. Tipicamente, cada um desses n pontos representa uma *observação* de $d + 2$ variáveis.

Felizmente a essa altura já estamos familiarizados com espaços de Hilbert, e sabemos que $\mathbb{R}^n$ é um deles! *Ergo*, pela propriedade de proximalidade, existe um único elemento $\hat{\mathbf{y}}$ pertencente ao espaço gerado pelas colunas de $\mathbf{X}$ tal que

$$|\mathbf{y} - \hat{\mathbf{y}}| \leq |\mathbf{y} - \boldsymbol{\xi}|, \quad \boldsymbol{\xi} \in \mathbb{R}^n.$$

Pelo Teorema 10 (e também pelo exercício 21), sabemos que — desde que a matriz $\mathbf{X}'\mathbf{X}$ seja invertível — o vetor $\hat{\mathbf{y}}$ é dado por

$$\hat{\mathbf{y}} = \mathbf{X}\hat{\boldsymbol{\beta}}, \quad \text{com } \hat{\boldsymbol{\beta}} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y}.$$

Note que isso nos diz quem é a solução na equação (16)! As componentes do vetor $\hat{\mathbf{y}} = (\hat{y}_1 \ \cdots \ \hat{y}_n)'$ são ditas os **valores ajustados** (por mínimos quadrados ordinários). Explicitamente,

$$\hat{y}_i = \mathbf{x}_i'\hat{\boldsymbol{\beta}} = \sum_{j=1}^d \hat{\beta}_j x_{ij}, \quad i \in \{1, \dots, n\}.$$

Lembrando que os *resíduos* foram definidos como $\hat{u}_i = y_i - \mathbf{x}_i'\hat{\boldsymbol{\beta}} = y_i - \hat{y}_i$, e pondo $\hat{\mathbf{u}} = (\hat{u}_1 \ \cdots \ \hat{u}_n)'$, podemos reescrever o sistema de equações (15) — agora com $\hat{\boldsymbol{\beta}}$ no lugar de $\mathbf{b}$ — como

$$\mathbf{y} = \mathbf{X}\hat{\boldsymbol{\beta}} + \hat{\mathbf{u}} = \hat{\mathbf{y}} + \hat{\mathbf{u}}. \quad (17)$$

A igualdade acima nada tem que ver, a princípio, com o modelo linear: não se trata de um modelo, mas de uma relação algébrica sempre válida,³⁵ relacionando os objetos que ali aparecem. De fato, temos o seguinte:

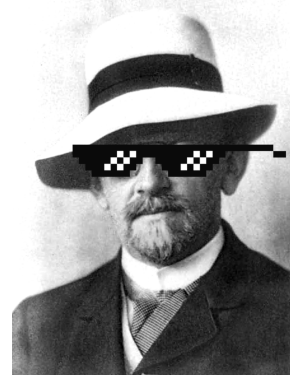
Teorema 12. Seja $M = \text{col}(\mathbf{X}) \subseteq \mathbb{R}^n$ o espaço gerado pelas colunas³⁶ de $\mathbf{X}$, e suponha que $\mathbf{X}'\mathbf{X}$ é invertível. Então $\hat{\mathbf{y}} = \mathbf{y}_M$ e $\hat{\mathbf{u}} = \mathbf{y}_{M^\perp}$.

Demonstração. Não há o que demonstrar: já vimos³⁷ que $\mathbf{X}\hat{\boldsymbol{\beta}}$ é a projeção ortogonal de $\mathbf{y}$ sobre $\text{col}(\mathbf{X})$, e já sabemos³⁸ que a decomposição $\mathbf{y} = \mathbf{y}_M + \mathbf{y}_{M^\perp}$ é única. ■

Corolário 13. Nas condições do Teorema 12, vale que

$$\mathbf{X}'\hat{\mathbf{u}} = \mathbf{0} \in \mathbb{R}^{d+1}.$$

Adicionalmente, se $\mathbf{X}$ possui uma coluna de 1's, então $\sum_{i=1}^n \hat{u}_i = 0$.



David Hilbert, *circa* 1912.

³⁵ Desde que $\mathbf{X}'\mathbf{X}$ seja invertível, de forma a garantir que $\hat{\boldsymbol{\beta}}$ esteja bem definido.

³⁶ Observe que aproveitamos a deixa para introduzir aqui a notação $\text{col}(\mathbf{X})$.

³⁷ Teorema 10.

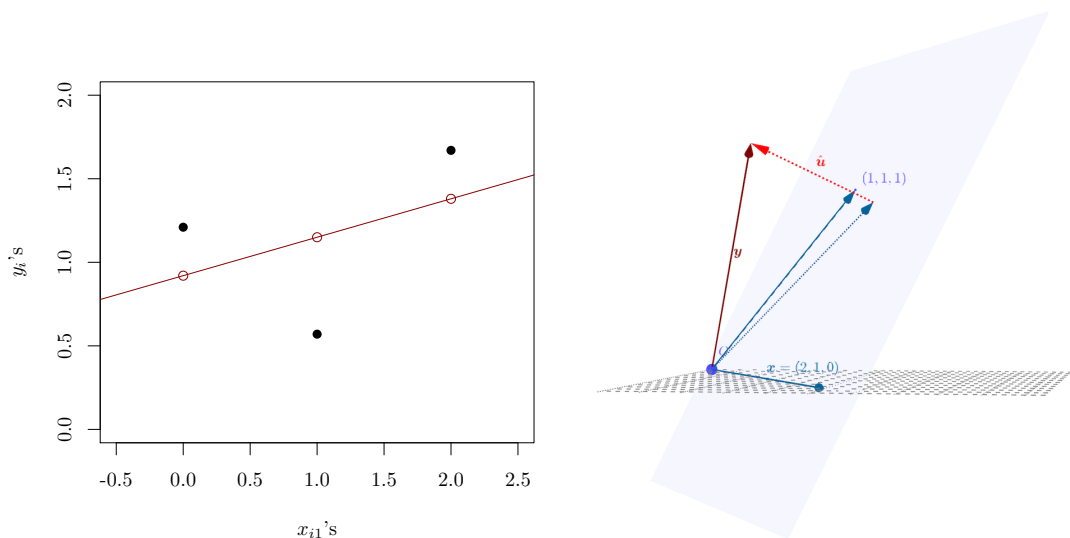
³⁸ Exercício 20.

Exemplo 14. Para ilustrar, consideremos (com $n = 3$ e $d = 1$) o seguinte conjunto de dados:

$$\mathbf{y}' = \begin{pmatrix} 1.67 & 0.57 & 1.21 \end{pmatrix}$$

$$\mathbf{X}' = \begin{pmatrix} 1 & 1 & 1 \\ 2 & 1 & 0 \end{pmatrix}$$

Fazendo as contas encontramos $\hat{\boldsymbol{\beta}} = (0.92 \quad 0.23)'$ e $\hat{\mathbf{y}} = (1.38 \quad 1.15 \quad 0.92)'$.



As duas figuras acima ilustram duas maneiras distintas de se visualizar os objetos envolvidos no problema. Na figura da esquerda, cada um dos $n = 3$ pontos em preto representa um dos pares ordenados (x_{i1}, y_i) . A diagonal em vermelho é a reta dada pela equação $\ell(\xi) = \hat{\beta}_0 + \hat{\beta}_1 \xi$, $\xi \in \mathbb{R}$. Os três círculos situados sobre a reta representam os pares ordenados $(x_{i1}, \hat{y}_i)$. A figura da direita representa os vetores envolvidos como pontos em $\mathbb{R}^n$, isto é, em $\mathbb{R}^3$. Aqui, cada uma das duas colunas da matriz $\mathbf{X}$ está representada por um vetor em azul em linha sólida. O hiperplano (subespaço) gerado pelas colunas de $\mathbf{X}$ aparece em azul-claro; o vetor $\mathbf{y}$, em vermelho-escuro, “aponta para fora” desse subespaço. O vetor de valores ajustados $\hat{\mathbf{y}}$ está representado pelo vetor em azul em linha pontilhada.

O Teorema de Frisch–Waugh–Lovell

Considere o seguinte *setup*: com $d \geq 1$ — de modo a garantir que a matriz $\mathbf{X}$ tenha pelo menos 2 colunas — tomemos inteiros $d_1 > 0$ e

$d_2 > 0$ tais que $d_1 + d_2 = d + 1$. Denotemos por $\mathbf{X}_1$ a matriz formada pelas d_1 primeiras colunas de $\mathbf{X}$ e, semelhantemente, por $\mathbf{X}_2$ a matriz das d_2 colunas remanescentes. Simbolicamente, podemos escrever

$$\mathbf{X} = \begin{bmatrix} \mathbf{X}_1 & \mathbf{X}_2 \end{bmatrix}.$$

Denotemos ainda por $\Pi_1^\perp$ a matriz $(n \times n)$ de projeção ortogonal³⁹ sobre o complemento ortogonal do espaço gerado pelas colunas de $\mathbf{X}_1$. Isto é, assumindo que $\mathbf{X}_1' \mathbf{X}_1$ é invertível, definimos $\Pi_1^\perp := \mathbf{I} - \mathbf{X}_1(\mathbf{X}_1' \mathbf{X}_1)^{-1} \mathbf{X}_1'$, onde $\mathbf{I}$ denota a matriz identidade $n \times n$. *Last but not least*, lembrando que $\hat{\boldsymbol{\beta}} = (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbf{y}$, escrevamos

$$\hat{\boldsymbol{\beta}}' = \begin{bmatrix} \hat{\boldsymbol{\beta}}_1' & \hat{\boldsymbol{\beta}}_2' \end{bmatrix}$$

onde $\hat{\boldsymbol{\beta}}_1$ é o vetor $d_1 \times 1$ cujas componentes correspondem às primeiras d_1 componentes de $\hat{\boldsymbol{\beta}}$ e assim por diante.

Informalmente, o Teorema de Frisch–Waugh–Lovell nos diz que, para computar $\hat{\boldsymbol{\beta}}_2$, podemos seguir a receita abaixo:

1. “Regrida” $\mathbf{y}$ em $\mathbf{X}_1$ e obtenha resíduos $\hat{\mathbf{u}}_1 := \Pi_1^\perp \mathbf{y}$.
2. “Regrida” cada coluna de $\mathbf{X}_2$ em $\mathbf{X}_1$ e obtenha uma “matriz de resíduos” $\hat{\mathbf{U}}_2 := \Pi_1^\perp \mathbf{X}_2$.
3. “Regrida” $\hat{\mathbf{u}}_1$ em $\hat{\mathbf{U}}_2$ para obter $\hat{\boldsymbol{\beta}}_2$.

O enunciado formal é o seguinte:

Teorema 15 (Frisch–Waugh–Lovell). Se $\mathbf{X}' \mathbf{X}$ é invertível, então

$$\hat{\boldsymbol{\beta}}_2 = (\mathbf{X}_2' \Pi_1^\perp \mathbf{X}_2)^{-1} \mathbf{X}_2' \Pi_1^\perp \mathbf{y}.$$

Remark 16. O costume nos concede que, ao escrever a solução

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbf{y},$$

do problema de minimização (16), seja omitida a dependência — muito evidente — que o termo no lado esquerdo dessa identidade tem nos termos que aparecem à direita. Ora, claramente $\hat{\boldsymbol{\beta}}$ é uma função de $\mathbf{y}$ e $\mathbf{X}$, e portanto podemos — quando o clamor por clareza assim demandar — escrever, e.g., $\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$. Com essa notação⁴⁰ — e notando que $\Pi_1^\perp$ é simétrica e idempotente — o Teorema de Frisch–Waugh–Lovell nos revela que

$$\hat{\boldsymbol{\beta}}(\Pi_1^\perp \mathbf{y}, \Pi_1^\perp \mathbf{X}_2) = \begin{bmatrix} \mathbf{0}_{(d_2 \times d_1)} & \mathbf{I}_{(d_2 \times d_2)} \end{bmatrix} \hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}).$$

³⁹ Veja o exercício 29.

⁴⁰ Por sinal, escrever $\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$ é uma maneira de dar sentido ao jargão “regredir $\mathbf{y}$ em $\mathbf{X}$.”

Demonstração. Seja $\Pi^\perp = \mathbf{I} - \mathbf{X}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'$ a projeção ortogonal sobre $\text{col}(\mathbf{X})^\perp$. Notando que $\hat{\mathbf{u}} = \Pi^\perp \mathbf{y}$, podemos escrever

$$\mathbf{y} = \mathbf{X}_1 \hat{\boldsymbol{\beta}}_1 + \mathbf{X}_2 \hat{\boldsymbol{\beta}}_2 + \Pi^\perp \mathbf{y}.$$

Isso nos dá, multiplicando a expressão acima, à esquerda, por $\mathbf{X}_2' \Pi_1^\perp$,

$$\mathbf{X}_2' \Pi_1^\perp \mathbf{y} = \mathbf{X}_2' \Pi_1^\perp \mathbf{X}_1 \hat{\boldsymbol{\beta}}_1 + \mathbf{X}_2' \Pi_1^\perp \mathbf{X}_2 \hat{\boldsymbol{\beta}}_2 + \mathbf{X}_2' \Pi_1^\perp \Pi^\perp \mathbf{y}.$$

Agora, o primeiro termo no lado direito da igualdade acima se anula, pois $\Pi_1^\perp \mathbf{X}_1 = \mathbf{0}$ (por quê?). De maneira análoga, o terceiro termo se anula: seu transposto é $\mathbf{y}' \Pi^\perp \Pi_1^\perp \mathbf{X}_2$, e valem as igualdades $\Pi^\perp \Pi_1^\perp = \Pi^\perp$ e $\Pi^\perp \mathbf{X}_2 = \mathbf{0}$ (por quê?).⁴¹ Disso segue que $\mathbf{X}_2' \Pi_1^\perp \mathbf{y} = \mathbf{X}_2' \Pi_1^\perp \mathbf{X}_2 \hat{\boldsymbol{\beta}}_2$ e a partir daí obtemos a identidade desejada. ■

⁴¹ Esses “por quês?” não são perguntas retóricas: faça o exercício 29!

Exercícios

Quaisquer erros nos enunciados são — *por supuesto!* — intencionais.

Exercício 23. Mostre cuidadosamente que qualquer solução para o problema

$$\min_{\mathbf{b} \in \mathbb{R}^{d+1}} |\mathbf{y} - \mathbf{X}\mathbf{b}|$$

é também uma solução para o problema

$$\min_{\mathbf{b} \in \mathbb{R}^{d+1}} |\mathbf{y} - \mathbf{X}\mathbf{b}|^2$$

e vice-versa.

Exercício 24. Encontre condições suficientes e necessárias para que o sistema de equações (15) admita uma solução $\mathbf{b}^*$ tal que $v_i(\mathbf{b}^*) = 0$ para todo i .

Exercício 25. Encontre condições suficientes e necessárias para que a matriz $\mathbf{X}'\mathbf{X}$ seja invertível.

Exercício 26. Calcule manualmente o valor de $\hat{\boldsymbol{\beta}}$ no exemplo 14. Quem são os vetores $\mathbf{x}_1$, $\mathbf{x}_2$ e $\mathbf{x}_3$ nesse exemplo?

Exercício 27. Demonstre o Corolário 13.

Exercício 28 (Representações alternativas). Verifique as seguintes igualdades:

1. $\hat{\beta} = (\sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i')^{-1} \sum_{i=1}^n \mathbf{x}_i y_i$.
2. $\sum_{i=1}^n \mathbf{x}_i u_i = \mathbf{0}$

Exercício 29. Considere as matrizes Π e $\Pi^\perp$ (ambas $n \times n$) definidas por

$$\Pi = \mathbf{X}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}', \quad \Pi^\perp = \mathbf{I} - \Pi,$$

onde $\mathbf{I}$ denota a matriz identidade (quais as dimensões?). Mostre que:

1. Π e $\Pi^\perp$ são simétricas e idempotentes.⁴²
2. Para qualquer vetor $\xi \in \text{col}(\mathbf{X})$, vale que $\Pi\xi = \xi$ e $\Pi^\perp\xi = \mathbf{0}$. Em particular, $\Pi\mathbf{X} = \mathbf{X}$ e $\Pi^\perp\mathbf{X} = \mathbf{0}$ (matriz de zeros, $n \times (d+1)$).
3. $\Pi\mathbf{y} = \hat{\mathbf{y}}$ e $\Pi^\perp\mathbf{y} = \hat{\mathbf{u}}$.
4. (Teorema de Pitágoras) $|\mathbf{y}|^2 = |\hat{\mathbf{y}}|^2 + |\hat{\mathbf{u}}|^2$.
5. Supondo agora que $\mathbf{X}$ é da forma $\mathbf{X} = [\mathbf{X}_1 \quad \mathbf{X}_2]$, sejam $\Pi_1 = \mathbf{X}_1(\mathbf{X}_1'\mathbf{X}_1)^{-1}\mathbf{X}_1'$ e $\Pi_1^\perp = \mathbf{I} - \Pi_1$. Mostre que $\Pi_1 \cdot \Pi = \Pi \cdot \Pi_1 = \Pi_1$ e $\Pi_1^\perp \Pi^\perp = \Pi^\perp \Pi_1^\perp = \Pi^\perp$.
6. Dê os detalhes da demonstração do Teorema de Frisch–Waugh–Lovell.

⁴² Simetria significa, e.g., $\Pi = \Pi'$ e idempotência, por seu turno, $\Pi \cdot \Pi = \Pi$.

Exercício 30. Seja $\bar{y}_n := n^{-1}(y_1 + \cdots + y_n)$. Verifique as seguintes propriedades:

1. $n^{-1}(\hat{y}_1 + \cdots + \hat{y}_n) = \bar{y}_n$.
2. se $d = 0$ e $x_{i0} = 1$, $i \in \{1, \dots, n\}$, então $\hat{\beta} = \bar{y}_n$.
3. se $d \geq 1$ e $\mathbf{X}'\mathbf{X}$ é invertível, então

$$\sum_{i=1}^n (y_i - \bar{y}_n)^2 = \sum_{i=1}^n \{(\hat{y}_i - \bar{y}_n)^2 + \hat{u}_i^2\}.$$

Exercício 31. O coeficiente de determinação é o número $R^2 \in [0, 1]$ definido por

$$R^2 := 1 - \frac{\sum_{i=1}^n \hat{u}_i^2}{\sum_{i=1}^n (y_i - \bar{y}_n)^2}$$

Mostre que

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

Exercício 32. Considere uma transformação (possivelmente não-linear) $\Phi: \mathbb{R}^{d+1} \rightarrow \mathbb{R}^k$. Para $i \in \{1, \dots, n\}$, sejam $\mathbf{z}_i := \Phi(\mathbf{x}_i)$, e $\mathbf{Z}$ a matriz cuja entrada i, j é z_{ij} . Interprete o vetor $\mathbf{Z}(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{y}$.

Exercício 33 (Computacional). Com a mesma notação do exercício anterior, seja $\Phi: \mathbb{R}^2 \rightarrow \mathbb{R}^3$ definida por

$$\Phi(\xi_0, \xi_1) = (\xi_0, \xi_1, \xi_1^3), \quad \boldsymbol{\xi} \in \mathbb{R}^2.$$

1. Utilize o software RStudio para gerar uma amostra de tamanho $n = 200$ do modelo

$$y_i = 1 + x_{i1} + x_{i1}^3 + u_i,$$

onde x_{i1} tem distribuição uniforme no intervalo $[-2, 2]$ e u_i tem distribuição Normal($0, \sigma^2$), com $\sigma = 0.03$, independente de x_{i1} .

2. Faça o gráfico de dispersão (no plano) dos pares (x_{i1}, y_i) simulados e compare-o com o gráfico de dispersão (em $\mathbb{R}^3$) das ternas (x_{i1}, x_{i1}^3, y_i) .
3. No gráfico 2D, acrescente a reta cuja equação é

$$\ell(\xi) = \hat{\beta}_0 + \hat{\beta}_1 \xi, \quad \xi \in \mathbb{R}.$$

4. No gráfico 2D, acrescente a cúbica cuja equação é

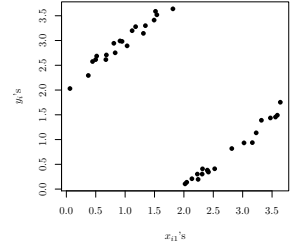
$$q(\xi) = \hat{\gamma}_0 + \hat{\gamma}_1 \xi_1 + \hat{\gamma}_2 \xi_1^3, \quad \xi \in \mathbb{R},$$

onde $\hat{\boldsymbol{\gamma}} = (\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{y}$, com $\mathbf{Z}$ definido como no exercício 32.

Exercício 34. Nesse exercício, vamos escrever $\hat{\boldsymbol{\beta}} = \hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$ para explicitar a dependência da solução (16) no vetor $\mathbf{y}$ e na matriz $\mathbf{X}$. Uma matriz quadrada A é dita **unitária** se $A'A = AA' = \mathbf{I}$. Mostre que, se A é uma matriz $(d+1) \times (d+1)$ unitária, então $\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}A) = A'\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$.

Exercício 35 (Paradoxo de Simpson). Considere o gráfico de dispersão dos pares (y_i, x_{i1}) , $i \in \{1, \dots, n\}$ apresentado ao lado. Seja D_i definida, para $i \in \{1, \dots, n\}$, pondo-se $D_i = 1$ se $x_{i1} > 2$ e $D_i = 0$ se $x_{i1} \leq 2$. Sejam ainda $\mathbf{X}$ e $\mathbf{Z}$ as matrizes definidas por

$$\mathbf{X} = \begin{pmatrix} 1 & x_{1,1} \\ 1 & x_{2,1} \\ \vdots & \vdots \\ 1 & x_{n,1} \end{pmatrix}, \quad \mathbf{Z} = \begin{pmatrix} 1 & x_{1,1} & D_1 \\ 1 & x_{2,1} & D_2 \\ \vdots & \vdots & \vdots \\ 1 & x_{n,1} & D_n \end{pmatrix}.$$



Usando a mesma notação do exercício 34, você consegue intuir, sem fazer as contas, qual será o sinal de $\hat{\beta}_1(\mathbf{y}, \mathbf{X})$? E quanto a $\hat{\beta}_1(\mathbf{y}, \mathbf{Z})$?

CÓDIGO PARA GERAR o gráfico do exercício no RStudio:

```
set.seed(3)
n=40
x = runif(n,0,4)
D = (x>2)
y = 2 - 4*D + x + rnorm(n, sd=.1)
plot(y~x, pch=16)
```

Esse exercício é um exemplo do *Paradoxo de Simpson*. Kadane⁴³ ilustra-o assim:

“Imagine two routes to the summit of a mountain, a difficult route D and an easier route $\bar{D}$. Imagine also two groups of climbers: amateurs A , and experienced climbers, $\bar{A}$. Suppose that a person has probabilities of reaching the summit R , as a function of the route and the experience of the climber as follows:

$$\begin{aligned} P\{R|\bar{D}, \bar{A}\} &= 0.8 & P\{R|\bar{D}, A\} &= 0.7 \\ P\{R|D, \bar{A}\} &= 0.4 & P\{R|D, A\} &= 0.3 \end{aligned}$$

Thus experienced climbers are more likely to reach the summit whichever route they take, and both groups are less likely to reach the summit using the more difficult route.

Further suppose also that the experienced climbers are believed to be more likely to take the more difficult route: $P\{D|\bar{A}\} = 0.75$ [and] $P\{D|A\} = 0.35 \dots$

... The events $R\bar{D}$ and RD are disjoint, and their union is R . Therefore, $P\{R|\bar{A}\} = \dots = 0.5$. Similarly $P\{R|A\} = \dots = 0.56$. Thus amateur climbers have a greater chance of reaching the summit (0.56) than do experienced climbers (0.5), although for each route they have a smaller chance.”

⁴³ Joseph B. Kadane. *Principles of Uncertainty*. Chapman and Hall/CRC, 2011

Propriedades amostrais do estimador de mínimos quadrados

“Se sofres por algo externo a ti, o que te causa sofrimento não é a coisa em si mesma, mas sim a tua *estimativa* a respeito da coisa. Tens o poder de eliminar o sofrimento imediatamente.”

Marcus Aurelius Antoninus, Meditações, Livro VIII.

APÓS FICAREM SUMIDOS durante o capítulo pregresso, os símbolos $\mathbb{P}$ e $\mathbb{E}$ estão de volta! Quer dizer, uma vez mais estamos lidando com objetos (nossos velhos conhecidos $\mathbf{y}$, os $\mathbf{x}_i$'s, $\mathbf{X}$...) que são *aleatórios* — no sentido formal de que são funções cujo domínio é o espaço amostral $(\Omega, \mathcal{A})$ e cujo contradomínio é a reta, ou um dos espaços Euclidianos (do qual, de fato, a reta já é um exemplar e portanto não precisaria ter sido mencionada na abertura dessa lista), ou algum espaço de matrizes (que, já vimos, também é indistinguível,⁴⁴ do ponto de vista de sua estrutura linear e topológica, de um espaço Euclidiano). Significativamente, chamamos esses objetos de *aleatórios* antes mesmo de amunciar-mos o espaço amostral com uma medida de probabilidade $\mathbb{P}$, embora ao fim e ao cabo seja ela quem capture — ao transpormos o simbolismo para a linguagem natural — a ideia de que aqueles objetos são aleatórios. Essa digressão é relevante pois, no decorrer deste capítulo e dos seguintes, faremos suposições distribucionais — por exemplo, de exogeneidade estrita, de Gaussianidade dos termos de erros, de que estamos diante de um mecanismo de amostragem aleatória etc. — e é importante se dar conta que esses atributos distribucionais não estão inteiramente codificados nas variáveis, vetores e matrizes aleatórias com que lidamos, mas sim emergem como uma propriedade em que entram



Estátua equestre de Marco Aurélio, Roma.

⁴⁴ Vide o exercício 16

como ingredientes tanto essas variáveis, vetores, matrizes aleatórias quanto também a medida de probabilidade $\mathbb{P}$.⁴⁵

Estimadores

Estimadores são variáveis (ou vetores, ou matrizes) aleatórias. Simples assim.⁴⁶ Um exemplo é o **estimador de mínimos quadrados ordinários**,⁴⁷ $\hat{\beta}$; o fato de que este é um vetor aleatório está implícito na expressão

$$\hat{\beta} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y} \equiv \left(\sum_{i=1}^n \mathbf{x}_i\mathbf{x}_i'\right)^{-1} \sum_{i=1}^n \mathbf{x}_i y_i.$$

De fato, embora raramente recorramos a tão detalhada notação, podemos escrever⁴⁸

$$\hat{\beta}(\omega) = (\mathbf{X}(\omega)'\mathbf{X}(\omega))^{-1}\mathbf{X}(\omega)'\mathbf{y}(\omega), \quad \omega \in \Omega,$$

em que — por óbvio — $[\mathbf{X}(\omega)]_{ij} = x_{ij}(\omega)$ etc. Vez que outra é conveniente fazer uma distinção, chamando de *estimador* ao vetor aleatório $\hat{\beta}$ e de uma *estimativa* ao valor realizado $\hat{\beta}(\omega) \in \mathbb{R}^{d+1}$. Essa é mesma distinção que fazemos ao chamar de *amostra aleatória* (no caso univariado) às variáveis aleatórias $y_1, \dots, y_n$ e de uma *realização dessa amostra aleatória* aos números reais $y_1(\omega), \dots, y_n(\omega)$.⁴⁹

O QUE NOS INTERESSA do ponto de vista de inferência estatística, contudo e sobretudo, é estabelecer alguma relação entre estimadores e *parâmetros*. Um **parâmetro** é um objeto da forma $\theta_{\mathbb{P}} \in \Theta$, em que Θ é — tipicamente — um subconjunto de $\mathbb{R}^k$, onde $k \in \mathbb{N}$,⁵⁰ ou, mais geralmente, podemos enxergar o parâmetro como sendo a aplicação $\mathbb{P} \mapsto \theta_{\mathbb{P}}$, no ímpeto de manter sempre na memória o fato de que *parâmetros dependem da medida de probabilidade* $\mathbb{P}$. Um exemplo de parâmetro é o vetor⁵¹

$$\beta = [\mathbb{E}(\mathbf{X}'\mathbf{X})]^{-1}\mathbb{E}(\mathbf{X}'\mathbf{y}) \equiv \left(\sum_{i=1}^n \mathbb{E}(\mathbf{x}_i\mathbf{x}_i')\right)^{-1} \sum_{i=1}^n \mathbb{E}(\mathbf{x}_i y_i).$$

Em particular, se

$$(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n) \quad (18)$$

é uma amostra iid, então vale a igualdade⁵²

$$\beta = (\mathbb{E}(\mathbf{x}_1\mathbf{x}_1'))^{-1}\mathbb{E}(\mathbf{x}_1 y_1). \quad (19)$$

⁴⁵ Veja o exercício 36.

⁴⁶ Bom, não tão simples porque a variável aleatória deve estar definida “antes” de introduzirmos uma medida de probabilidade no espaço amostral; variáveis aleatórias do tipo $\vartheta = \mathbb{E}(x)$ “não valem”!

⁴⁷ Poderíamos tê-lo assim batizado já nos capítulos anteriores, mas optamos por não fazê-lo pois o epíteto ‘estimador’ em nada nos auxiliaria a compreensão.

⁴⁸ Anteriormente havíamos introduzido a notação $\hat{\beta}(\mathbf{y}, \mathbf{X})$, cuja dependência em ω aqui apareceria, com um ligeiro abuso notacional, assim: $\hat{\beta}(\omega) = \hat{\beta}(\mathbf{y}(\omega), \mathbf{X}(\omega))$.

⁴⁹ Na nossa notação, escreveríamos $y_i(\omega) = y_i^{\text{obs}}$ etc.

⁵⁰ E nada impede $\Theta = \mathbb{R}^k$!

⁵¹ A dependência de β em $\mathbb{P}$ está implícita, já que $\mathbb{E}$ depende de $\mathbb{P}$.

⁵² Recorde o exercício 8!

Uma relação de particular interesse entre um estimador $\vartheta: \Omega \rightarrow \mathbb{R}^k$ e um parâmetro $\theta \equiv \theta_{\mathbb{P}} \in \mathbb{R}^k$ é o **viés**, definido por

$$\text{bias}(\vartheta, \theta) := \mathbb{E}(\vartheta) - \theta. \quad (20)$$

Um estimador ϑ é dito **não-enviesado** para um parâmetro θ se $\mathbb{E}(\vartheta) = \theta$, isto é, se $\text{bias}(\vartheta, \theta) = \mathbf{0}$.⁵³

De posse desse vocabulário, podemos agora nos perguntar se o estimador de MQO é não-enviesado para o β populacional. A resposta é que, em geral, $\text{bias}(\hat{\beta}, \beta) \neq \mathbf{0}$; o seguinte resultado oferece, como corolário, condições suficientes e necessárias para que esse viés seja nulo.

Teorema 17. Sejam $\mathbf{y}$ e $\mathbf{x}_i, i \in \{1, \dots, n\}$, vetores aleatórios em $\mathbb{R}^n$ e $\mathbb{R}^{d+1}$, respectivamente, e seja $\mathbf{X}$ a matriz aleatória $n \times (d+1)$ cuja componente i, j é x_{ij} . Seja ainda $\mathbb{P}$ uma medida de probabilidade em $(\Omega, \mathcal{A})$ satisfazendo

(i) $\mathbb{P}(\xi' \mathbf{X}' \mathbf{X} \xi > 0 \text{ para todo } \xi \in \mathbb{R}^{d+1} \text{ não nulo}) = 1$.

(ii) $\mathbb{E}_{\mathbb{P}}|\hat{\beta}| < \infty$, onde $\hat{\beta} := (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbf{y}$.

Então,

$$\mathbb{E}(\hat{\beta} | \mathbf{X}) = \beta + (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbb{E}(\mathbf{u} | \mathbf{X}), \quad (21)$$

em que $\beta := [\mathbb{E}(\mathbf{X}' \mathbf{X})]^{-1} \mathbb{E}(\mathbf{X}' \mathbf{y}) \in \mathbb{R}^{d+1}$ e onde $\mathbf{u}$ é o vetor aleatório em $\mathbb{R}^n$ definido por $\mathbf{u} := \mathbf{y} - \mathbf{X}\beta$.

Corolário 18. Nas condições do Teorema 17, tem-se $\text{bias}(\hat{\beta}, \beta) = \mathbf{0}$ se, e somente se, $\mathbb{E}((\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbf{u}) = \mathbf{0}$. Em particular, se $\mathbf{X}$ é estritamente exógeno, então $\hat{\beta}$ é não-enviesado para β .

Demonstração. Primeiramente, verifiquemos que o estimador $\hat{\beta}$ e o parâmetro β estão bem definidos. Quanto a $\hat{\beta}$, é imediato, pois a condição (i) impõe nada mais que a invertibilidade da matriz $\mathbf{X}(\omega)' \mathbf{X}(\omega)$, para ω pertencente a um evento $E^* \subseteq \Omega$ com $\mathbb{P}(E^*) = 1$. Quanto a β , vemos que, para qualquer $\xi \in \mathbb{R}^{d+1}$, vale a desigualdade $\mathbb{E}(\xi' \mathbf{X}' \mathbf{X} \xi) > 0$, novamente pela condição (i). Disso segue, usando a identidade $\xi' \mathbb{E}(\mathbf{X}' \mathbf{X}) \xi = \mathbb{E}(\xi' \mathbf{X}' \mathbf{X} \xi)$, que a matriz $\mathbb{E}(\mathbf{X}' \mathbf{X})$ é invertível.

Para a tese temos, das definições, $\mathbf{y} = \mathbf{X}\beta + \mathbf{u}$. Logo,

$$\hat{\beta} := (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' (\mathbf{X}\beta + \mathbf{u}) = \beta + (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbf{u}.$$

⁵³ Note que a expressão definindo o viés depende da medida de probabilidade $\mathbb{P}$. Em particular — e isso não é lá muito óbvio — um estimador pode ser não enviesado para um parâmetro θ sob uma medida de probabilidade e ser enviesado sob outra. Veja o exercício 37.

Tomando a esperança incondicional em ambos os lados dessa identidade, obtemos a caracterização declarada no corolário. Tomando a esperança condicional em $\mathbf{X}$, seguem o teorema e a condição suficiente para anular o viés no corolário. ■

Quanto à dispersão de $\widehat{\beta}$, temos o seguinte.

Teorema 19. Nas condições do Teorema 17, seja

$$\mathbf{X}_* := (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'.$$

Então

1. $\mathbb{V}(\widehat{\beta}) = \mathbb{E}(\mathbf{X}_* \mathbf{u} \mathbf{u}' \mathbf{X}_*) - \mathbb{E}(\mathbf{X}_* \mathbf{u}) \mathbb{E}(\mathbf{u}' \mathbf{X}_*)$.
2. $\mathbb{V}(\widehat{\beta} | \mathbf{X}) = \mathbf{X}_* \mathbb{V}(\mathbf{u} | \mathbf{X}) \mathbf{X}_*$

Adicionalmente,

3. se bias $(\widehat{\beta}, \beta) = 0$, então

$$\mathbb{V}(\widehat{\beta}) = \mathbb{E}\left\{\mathbb{V}(\widehat{\beta} | \mathbf{X})\right\} + \mathbb{E}\left\{\mathbf{X}_* \mathbb{E}(\mathbf{u} | \mathbf{X}) \mathbb{E}(\mathbf{u}' | \mathbf{X}) \mathbf{X}_*\right\}.$$

4. se $\mathbf{X}$ é estritamente exógeno, então

$$\mathbb{V}(\widehat{\beta} | \mathbf{X}) = \mathbf{X}_* \mathbb{E}(\mathbf{u} \mathbf{u}' | \mathbf{X}) \mathbf{X}_*.$$

5. se os termos de erro são **homoscedásticos**, no sentido de que $\mathbb{V}(\widehat{\mathbf{u}} | \mathbf{X}) = \sigma^2 \mathbf{I}$ para alguma constante $\sigma > 0$, então $\mathbb{V}(\widehat{\beta} | \mathbf{X}) = \sigma^2 (\mathbf{X}'\mathbf{X})^{-1}$.

O Teorema de Gauss–Markov

O Teorema de Gauss–Markov é a cereja do bolo no estudo das propriedades estatísticas do estimador de MQO em amostras finitas. Ele afirma que, sob certas condições, o estimador $\widehat{\beta}$ possui a propriedade de ser o *melhor estimador linear não-enviesado* para β — uma propriedade que, em inglês, recebe a carinhosa alcunha de BLUE (*best linear unbiased estimator*). Em que pese sua importância, esse resultado — tal qual uma cereja não figurativa em um bolo não figurativo — deve ser usado com ponderação. Não devemos nos empolgar com a aparente invencibilidade do nosso amigo “beta-chapéu”. Por três razões, pelo menos: primeiro, uma das hipóteses

do teorema é de que os termos de erro na equação $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$ são homoscedásticos, isto é satisfazem a condição $\mathbb{V}(\hat{\mathbf{u}} \mid \mathbf{X}) = \sigma^2 \mathbf{I}$. Essa hipótese é, em um sem fim de casos, inadequada, pois impede que a variância da resposta dependa do nível dos regressores, algo que raramente ocorre em dados observacionais. Segundo, sua tese restringe-se à classe de estimadores *lineares* e *não-enviesados* e, sem a ressalva de que podem existir estimadores *não-lineares* ou *enviesados* para $\boldsymbol{\beta}$ que sejam “melhores” do que $\hat{\boldsymbol{\beta}}$, o teorema soa mais poderoso do que realmente é. Terceiro, a noção de melhor aqui refere-se à ordem parcial no espaço das matrizes de covariância $(d+1) \times (d+1)$, em que uma matriz A é dita “maior” do que uma matriz B sse $A - B$ é positiva semidefinida; nada impede, todavia, que avaliemos a qualidade de um estimador segundo outros critérios.

Teorema 20 (Gauss–Markov). Nas condições do Teorema 17, suponha adicionalmente que

1. $\mathbb{E}(\mathbf{u} \mid \mathbf{X}) = \mathbf{0}$ quase certamente. (exogeneidade)
2. $\mathbb{E}(\mathbf{u}\mathbf{u}' \mid \mathbf{X}) = \sigma^2 \mathbf{I}$, para algum $\sigma > 0$. (homoscedasticidade)

Seja ainda $\mathbf{A}$ uma matriz $n \times (d+1)$ da forma $\mathbf{A} = \varphi \circ \mathbf{X}$, onde φ é uma aplicação de $\mathbb{R}^{n \times (d+1)}$ em $\mathbb{R}^{n \times (d+1)}$. Se $\mathbb{E}(\mathbf{A}'\mathbf{y} \mid \mathbf{X}) \stackrel{q.c.}{=} \boldsymbol{\beta}$, então $\mathbb{V}(\mathbf{A}'\mathbf{y} \mid \mathbf{X}) - \mathbb{V}(\hat{\boldsymbol{\beta}} \mid \mathbf{X})$ é positiva semidefinida.

Exercícios

Quaisquer erros nos enunciados são — provavelmente! — intencionais.

Exercício 36. Considere o espaço amostral

$$(\Omega, \mathcal{A}) := ([0, 1], \text{“Borelianos em } [0, 1]\text{”})$$

e a variável aleatória definida, para $\omega \in \Omega$, por $x(\omega) := \omega$. Sejam $\mathbb{P}_1$ e $\mathbb{P}_2$ as medidas de probabilidade em $(\Omega, \mathcal{A})$ definidas, respectivamente, por

$$\mathbb{P}_1[a, b] = b - a, \quad 0 \leq a < b \leq 1$$

e

$$\mathbb{P}_2[a, b] = \frac{1}{2} \mathbb{I}_{\{0 \in [a, b]\}} + \frac{1}{2} \mathbb{I}_{\{1 \in [a, b]\}}, \quad 0 \leq a < b \leq 1.$$

Verifique que $x \sim \text{Uniforme}[0, 1]$ no espaço de probabilidade $(\Omega, \mathcal{A}, \mathbb{P}_1)$ e que $x \sim \text{Bernoulli}(1/2)$ no espaço de probabilidade $(\Omega, \mathcal{A}, \mathbb{P}_2)$.

Exercício 37. Nas mesmas condições do exercício acima, seja $\mathbb{P}_3$ a medida de probabilidade definida por

$$\mathbb{P}_3[a, b] = b^2 - a^2, \quad 0 \leq a < b \leq 1.$$

Considere o parâmetro $\theta_{\mathbb{P}} = \mathbb{E}_{\mathbb{P}}(x)$, onde o subscrito na esperança indica a medida de probabilidade subjacente. Considere ainda os estimadores ϑ_1 e ϑ_2 definidos, para $\omega \in \Omega$, respectivamente por

$$\vartheta_1(\omega) := x(\omega) \quad \text{e} \quad \vartheta_2(\omega) := \frac{1}{2}.$$

1. Mostre que ϑ_1 é não-enviesado para $\theta_{\mathbb{P}}$ sob qualquer uma das medidas de probabilidade $\mathbb{P}_1$, $\mathbb{P}_2$ e $\mathbb{P}_3$.
2. Mostre que ϑ_2 é não-enviesado para $\theta_{\mathbb{P}}$ sob as medidas de probabilidade $\mathbb{P}_1$ e $\mathbb{P}_2$, mas é enviesado para $\theta_{\mathbb{P}_3}$.

Exercício 38. Mostre que $\mathbb{V}(z) = \mathbb{E}(\mathbb{V}(z | \mathbf{X})) + \mathbb{V}(\mathbb{E}(z | \mathbf{X}))$.

Exercício 39. Nas condições do Teorema 17, mostre que

$$\text{cov}(\hat{\beta}, \hat{u} | \mathbf{X}) = \mathbf{0}.$$

Exercício 40. Demonstre o Teorema 19.

Exercício 41. Mostre que, nas condições do Teorema 17, se $\mathbf{u}$ e $\mathbf{X}$ são independentes, com $\mathbf{u} \sim N(\mathbf{0}, \sigma^2 \mathbf{I})$, então

$$\hat{\beta} | \mathbf{X} \sim N(\beta, \sigma^2 (\mathbf{X}' \mathbf{X})^{-1}).$$

Conclua que a distribuição incondicional de $\hat{\beta}$ é uma **mistura** de distribuições Normais.

Exercício 42. Nas condições do Teorema 17, considere o estimador

$$\hat{\sigma}^2 := \frac{|\hat{\mathbf{u}}|^2}{n - (d + 1)}.$$

Mostre que, se $\mathbf{X}$ é estritamente exógeno e $\mathbb{V}(\mathbf{u} | \mathbf{X}) = \sigma^2 \mathbf{I}$, para algum $\sigma > 0$, então $\mathbb{E}(\hat{\sigma}^2) = \sigma^2$. A raiz $\hat{\sigma}$ é dita o **erro padrão da regressão**.

Exercício 43. Nas condições do Teorema 17, seja Π a matriz de projeção ortogonal sobre $\text{col}(\mathbf{X})$, isto é, $\Pi := \mathbf{X}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'$, e seja $\Pi^\perp := \mathbf{I} - \Pi$ a projeção ortogonal sobre $\text{col}(\mathbf{X})^\perp$. Mostre que $\hat{\mathbf{u}} = \Pi^\perp \mathbf{u}$. Conclua que $|\hat{\mathbf{u}}|^2 = \mathbf{u}'\Pi^\perp \mathbf{u}$.

Exercício 44 (Medidas empíricas). Seja $n \in \mathbb{N}$ fixado. Suponha que

$$(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n) \quad (22)$$

é uma amostra aleatória de $(\mathbf{x}, y)$. Para cada $\omega \in \Omega$, seja $\hat{\mathbb{P}}_{n,\omega}$ a medida de probabilidade (discreta) dada por

$$\hat{\mathbb{P}}_{n,\omega}[(\mathbf{x}, y) = (\mathbf{x}_i(\omega), y_i(\omega))] = \frac{1}{n}, \quad \forall i \in \{1, \dots, n\}.$$

Isto é, $\hat{\mathbb{P}}_n$ é a medida de probabilidade (aleatória) em Ω que atribui massa $1/n$ para cada um dos n pontos da amostra (22).⁵⁴ Escreva $\hat{\mathbb{E}}_{n,\omega}$ para denotar o valor esperado com respeito à medida de probabilidade $\hat{\mathbb{P}}_{n,\omega}$. Mostre que, para qualquer $\varphi: \mathbb{R}^{d+1} \times \mathbb{R} \rightarrow \mathbb{R}^k$ contínua, vale

$$\hat{\mathbb{E}}_{n,\omega}(\varphi \circ (\mathbf{x}, y)) = \frac{1}{n} \sum_{i=1}^n \varphi(\mathbf{x}_i(\omega), y_i(\omega)), \quad \omega \in \Omega.$$

Conclua que

$$\hat{\beta}(\omega) = \left(\hat{\mathbb{E}}_{n,\omega}(\mathbf{x}_1 \mathbf{x}_1') \right)^{-1} \hat{\mathbb{E}}_{n,\omega}(\mathbf{x}_1 y_1)$$

e compare com a equação (19).

Exercício 45. No contexto de um modelo linear com regressores estocásticos, que é o caso em quase toda a literatura de econometria, o Teorema de Gauss–Markov costuma ser enunciado de maneira não completamente rigorosa, algo que só se percebe ao ler cuidadosamente a sua demonstração. Em particular, o termo “não-enviesado” é empregado sem deixar claro que, na prova, será usada a propriedade mais forte de que o estimador deve ser *condicionalmente não-enviesado*. De fato, se considerarmos estimadores incondicionalmente não-enviesados, o Teorema de Gauss–Markov pode deixar de valer.⁵⁵ Aqui está um enunciado rigoroso do teorema (o exercício é demonstrá-lo):

⁵⁴ $\hat{\mathbb{P}}_n$ é dita uma *medida empírica*. Um detalhe é que essa medida de probabilidade, em geral, não pode ser definida na σ -álgebra original $\mathcal{A}$, mas somente naquela induzida pela amostra aleatória, isto é, em $\sigma(\mathbf{y}, \mathbf{X})$.

⁵⁵ Juliet Popper Shaffer. The Gauss–Markov Theorem and random regressors. *The American Statistician*, 45(4):269–273, 1991

Teorema 21. Considere dados um espaço amostral $(\Omega, \mathcal{A})$, uma matriz $n \times d$ aleatória, $\mathbf{X}$, e um vetor $n \times 1$ aleatório, $\mathbf{v}$. Seja $\mathfrak{M}$ a coleção de todas as medidas de probabilidade $\mathbb{P}: \mathcal{A} \rightarrow \mathbb{R}$ que satisfazem o seguinte:

1. $\mathbb{P}(\text{rank}(\mathbf{X}) = d) = 1$
2. $\mathbb{E}_{\mathbb{P}}(\mathbf{v}'\mathbf{v}) < \infty$
3. $\mathbb{E}_{\mathbb{P}}(\mathbf{v}|\mathbf{X}) = \mathbf{0}$
4. $\mathbb{E}_{\mathbb{P}}(\mathbf{v}\mathbf{v}'|\mathbf{X}) = \mathbf{Id}$

Seja ainda $\psi : M(n \times d) \rightarrow M(n \times d)$ uma aplicação mensurável (onde $M(n \times d)$ denota o espaço vetorial das matrizes $n \times d$), e ponha $\mathbf{X}_{\psi} = \psi \circ \mathbf{X}$.

Suponha que, para toda $\mathbb{P} \in \mathfrak{M}$, todo $\boldsymbol{\beta} \in \mathbb{R}^d$ e todo $\sigma > 0$, vale

$$\mathbb{E}_{\mathbb{P}}(\mathbf{X}'_{\psi}(\mathbf{X}\boldsymbol{\beta} + \sigma\mathbf{v}) | \mathbf{X}) = \boldsymbol{\beta}, \quad \mathbb{P}\text{-q.c.}$$

Então, para qualquer $\mathbb{P} \in \mathfrak{M}$, qualquer $\boldsymbol{\beta} \in \mathbb{R}^d$ e qualquer $\sigma > 0$, tem-se que a matriz $d \times d$

$$\mathbb{V}_{\mathbb{P}}(\mathbf{X}'_{\psi}(\mathbf{X}\boldsymbol{\beta} + \sigma\mathbf{v}) | \mathbf{X}) - \mathbb{V}_{\mathbb{P}}((\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'(\mathbf{X}\boldsymbol{\beta} + \sigma\mathbf{v}) | \mathbf{X})$$

é positiva semidefinida.

περὶ τῶν ἀσυμπτῶτων (sobre assíntotas)

“En la parte inferior del escalón, hacia la derecha, vi una pequeña esfera tornasolada, de casi intolerable fulgor. Al principio la creí giratoria; luego comprendí que ese movimiento era una ilusión producida por los vertiginosos espectáculos que encerraba. El diámetro del Aleph sería de dos o tres centímetros, pero el espacio cósmico estaba ahí, sin disminución de tamaño.”

J. L. Borges, *El Aleph*.



O conto *El Aleph*, de Borges, inicia com uma citação da famosa passagem em que Hamlet considera a possibilidade de viver recluso numa casca de noz e, ainda assim, proclamar-se Rei do espaço infinito. Ilustrado nesse trecho está o truísmo de que não compreendemos, intuitivamente, o conceito de infinito — ou ao menos o fato de que esse conceito pode ser enganoso! Pois bem, iniciaremos neste capítulo uma jornada, que será cumprida no próximo, cuja meta é estudar as propriedades distribucionais do estimador de mínimos quadrados ordinários *quando o tamanho da amostra vai para o infinito*.⁵⁶ Ocorre que amostras são sempre finitas: estamos sempre confinados ao interior, às fronteiras, duma casca de um fruto da noqueira. Por essa razão, faz sentido ter em mente que aquilo que estaremos estudando são, na verdade, aproximações as quais são válidas em amostras “suficientemente grandes” (mas ainda finitas). As duas mais importantes dessas propriedades — pertinentes ao estimador de MQO, sob certas suposições — são sua *consistência* e *normalidade assintóticas*. O primeiro passo será introduzir e estudar os conceitos de convergência a partir dos quais serão tornadas precisas as noções de consistência e normalidade assintótica mencionadas acima. Trataremos do estimador de mínimos quadrados no capítulo seguinte.

⁵⁶ A essas costumamos chamar de *propriedades assintóticas*.

Convergência em probabilidade

Suponha que nos é dada uma sequência $(z_n)_{n \geq 1} \equiv (z_1, z_2, \dots)$ de vetores aleatórios em $\mathbb{R}^k$, onde $k \in \mathbb{N}$ está fixado, e seja z também um vetor aleatório em $\mathbb{R}^k$. Dizemos⁵⁷ que a sequência $(z_n)_{n \geq 1}$ **converge em probabilidade para z** se, e somente se, valer a seguinte propriedade: dado $\delta > 0$ arbitrário, é possível encontrar um número natural n_δ tal que valha a desigualdade

$$\mathbb{P}(|z_n - z| > \delta) < \delta \quad \text{para todo } n \geq n_\delta.$$

Isso é o mesmo que dizer que $\lim_{n \rightarrow \infty} \mathbb{P}(|z_n - z| > \delta) = 0$, para qualquer $\delta > 0$.⁵⁸ Nessas condições, dizemos também que z é o **limite em probabilidade** da sequência $(z_n)_{n \geq 1}$, e escrevemos

- (i) $z_n \xrightarrow{p} z$;
- (ii) $\text{plim}_{n \rightarrow \infty} z_n = z$;
- (iii) “ $z_n \rightarrow z$ em probabilidade”,

etc. Se z é uma variável aleatória degenerada (digamos, com $\mathbb{P}(z = \xi) = 1$ para algum $\xi \in \mathbb{R}^k$) então dizer que $z_n \xrightarrow{p} \xi$ é o mesmo que afirmar que, para valores de n “suficientemente grandes”, a probabilidade (ou, para sermos bem precisos, a $\mathbb{P}$ -probabilidade) de encontrarmos a variável aleatória z_n na região

$$\overline{\text{ball}}(\xi; \delta) := \{\eta \in \mathbb{R}^k : |\eta - \xi| \leq \delta\}$$

pode ser tomada tão próxima de 1 quanto queiramos (ênfatizando: desde que o índice n seja tomado suficientemente grande).

Exemplo 22. No espaço de probabilidade do exercício 22, considere a sequência de variáveis aleatórias $(z_n)_{n \geq 1}$ definidas, para $\omega \in (0, 1]$ e $n \geq 1$, via

$$z_n(\omega) := \begin{cases} 1, & \text{se } n \leq (1 + \omega) \cdot 2^{f(n)} < n + 1, \\ 0, & \text{caso contrário,} \end{cases}$$

onde $f(n) := \lfloor \log_2(n) \rfloor$ = “maior número inteiro majorado por $\log_2(n)$ ”. Vejamos que $z_n \xrightarrow{p} 0$. Seja $\delta > 0$. Se $\delta \geq 1$ podemos tomar $n_\delta = 1$, pois nesse caso de fato temos, para todo $n \geq 1$,

$$\mathbb{P}(|z_n - 0| > \delta) = \mathbb{P}(\emptyset) = 0 < \delta,$$

⁵⁷ Se quisermos ser extremamente precisos na terminologia, devemos perceber que a medida de probabilidade $\mathbb{P}$ está aí implícita: poderíamos com maior exatidão dizer que $(z_n)_{n \geq 1}$ converge em $\mathbb{P}$ -probabilidade para z .

⁵⁸ Veja o exercício 46.

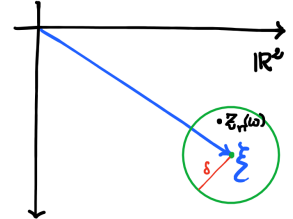


Figura 4: A figura exemplifica um ω pertencente ao evento $\{|z_n - \xi| \leq \delta\}$, o qual poderíamos reescrever como $\{z_n \in \overline{\text{ball}}(\xi; \delta)\}$. Se $z_n \xrightarrow{p} \xi$, o conjunto de tais ω 's tem probabilidade próxima de 1, para todo índice n a partir de n_δ .

já que $z_n(\cdot) \leq 1$ para todo n . Por outro lado, se $0 < \delta < 1$, então temos a igualdade de eventos

$$[|z_n - 0| > \delta] = [z_n = 1] = \{\omega : n \leq (1 + \omega) \cdot 2^{f(n)} < n + 1\}.$$

Quer dizer, para $\delta \in (0, 1)$ e $n \geq 1$ o evento $[|z_n - 0| > \delta]$ coincide com o intervalo

$$\left(\frac{n}{2^{f(n)}} - 1, \frac{n+1}{2^{f(n)}} - 1 \right].$$

Daí, pela definição de $\mathbb{P}$, temos $\mathbb{P}(|z_n - 0| > \delta) = 2^{-f(n)}$. Escrevendo $f(n) = \log_2(n) - h(n)$, onde $h(n) \in [0, 1)$, temos $2^{-f(n)} = 2^{h(n)}/n \leq 2/n$ e, tomando $n_\delta := \lceil 4/\delta \rceil$, obtemos finalmente $\mathbb{P}(|z_n - 0| > \delta) < \delta$ para todo $n \geq n_\delta$.

Convergência em distribuição

Antes de definirmos convergência em distribuição, precisamos de dois lembretes. O primeiro deles, de que a **função de distribuição acumulada**⁵⁹ de um vetor aleatório $\mathbf{z}$ em $\mathbb{R}^k$ é definida como sendo a função $F_{\mathbf{z}}: \mathbb{R}^k \rightarrow \mathbb{R}$ cuja regra é

$$F_{\mathbf{z}}(\boldsymbol{\xi}) := \mathbb{P}(z_1 \leq \xi_1, \dots, z_k \leq \xi_k), \quad \boldsymbol{\xi} \equiv (\xi_1 \ \cdots \ \xi_k)' \in \mathbb{R}^k.$$

O segundo, de que uma função $h: \mathbb{R}^k \rightarrow \mathbb{R}^d$ é dita **contínua no ponto** $\boldsymbol{\xi} \in \mathbb{R}^k$ se, e somente se, vale a implicação

$$\boldsymbol{\xi}_n \rightarrow \boldsymbol{\xi} \in \mathbb{R}^k \implies h(\boldsymbol{\xi}_n) \rightarrow h(\boldsymbol{\xi}) \in \mathbb{R}^d.$$

Se h é contínua no ponto $\boldsymbol{\xi}$, chamemo-lo então – mui apropriadamente – de um **ponto de continuidade de h** . Caso h seja contínua em todos os pontos do seu domínio, dizemos que h é **contínua**.⁶⁰

SEJAM $(\mathbf{z}_n)_{n \geq 1}$ e $\mathbf{z}$ como na seção anterior. Podemos, enfim, definir que a sequência $(\mathbf{z}_n)_{n \geq 1}$ **converge em distribuição** para $\mathbf{z}$ se, e somente se, para cada ponto de continuidade $\boldsymbol{\xi} \in \mathbb{R}^k$ da função de distribuição acumulada $F_{\mathbf{z}}$, a sequência numérica $F_{\mathbf{z}_n}(\boldsymbol{\xi})$ convergir para $F_{\mathbf{z}}(\boldsymbol{\xi})$, isto é, se, e somente se, valer o seguinte:

$$\text{“se } \boldsymbol{\xi} \text{ é ponto de continuidade de } F_{\mathbf{z}}, \text{ então } \lim_{n \rightarrow \infty} F_{\mathbf{z}_n}(\boldsymbol{\xi}) = F_{\mathbf{z}}(\boldsymbol{\xi})\text{”}.$$

Nesse contexto, dizemos que $\mathbf{z}$ é o **limite em distribuição** da sequência $(\mathbf{z}_n)_{n \geq 1}$, e escrevemos, e.g.

⁵⁹ A importância dessa noção é que uma medida de probabilidade em $\mathbb{R}^k$, $k \geq 1$, é unicamente determinadas por sua função de distribuição acumulada (e reciprocamente).

⁶⁰ Uma função h é contínua nesse sentido se, e somente se, $h^{-1}(U)$ é um subconjunto aberto de $\mathbb{R}^k$, para todo subconjunto U aberto em $\mathbb{R}^d$.

- (a) $\mathbf{z}_n \xrightarrow{d} \mathbf{z}$;
 (b) $\mathbf{z}_n \xrightarrow{\mathcal{L}} \mathbf{z}$;
 (c) $\mathcal{L}(\mathbf{z}_n) \rightarrow \mathcal{L}(\mathbf{z})$; (o $\mathcal{L}$ é para convergência em $\mathcal{L}$ ei)
 (d) “ $\mathbf{z}_n \rightarrow \mathbf{z}$ em distribuição”,

etc.

Na literatura de teoria da probabilidade, convergência em distribuição também é chamada de *convergência em lei*, de *convergência fraca*, e de *convergência fraca** (lê-se “convergência-fraca-estrela” mesmo), sendo essa última terminologia — proveniente da análise funcional — a mais adequada. É importante notar que convergência em distribuição realmente não diz respeito aos vetores aleatórios em si, mas sim às suas respectivas funções de distribuição acumuladas, e nesse sentido o vetor aleatório limítrofe, $\mathbf{z}$, aparece apenas de maneira auxiliar (de fato, até mesmo a sequência $(\mathbf{z}_n)_{n \geq 1}$ é dispensável, mas vamos com calma): se H é uma função de distribuição acumulada qualquer (com domínio $\mathbb{R}^k$) tal que $F_{\mathbf{z}_n}(\boldsymbol{\xi}) \rightarrow H(\boldsymbol{\xi})$ para qualquer $\boldsymbol{\xi}$ que seja ponto de continuidade de H , então dizemos que $(\mathbf{z}_n)_{n \geq 1}$ **converge em distribuição para H** , e escrevemos

$$(e) \mathbf{z}_n \xrightarrow{d} H.$$

Apropriadamente, a distribuição H é dita a **distribuição limite** da sequência $(\mathbf{z}_n)_{n \geq 1}$. Adicionalmente, se H possui um “código padrão” (por exemplo,⁶¹ $N(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ ou $\text{Uniforme}(0, 1]$ etc), escrevemos também, e.g.

$$(f) \mathbf{z}_n \xrightarrow{d} N(\boldsymbol{\mu}, \boldsymbol{\Sigma}).$$

Nesse último caso, dizemos que a sequência $(\mathbf{z}_n)_{n \geq 1}$ é **assintoticamente normal** (com *vetor de médias $\boldsymbol{\mu}$* e *matriz de covariâncias $\boldsymbol{\Sigma}$*) e, no caso particular em que $\boldsymbol{\mu} = \mathbf{0}$ e $\boldsymbol{\Sigma} = \mathbf{Id}$, que a sequência é **assintoticamente normal padrão**. O conceito de convergência em distribuição é particularmente importante pois permite que calculemos probabilidades aproximadas quando a distribuição exata da variável aleatória $\mathbf{z}_n$ é desconhecida e difícil de se obter analiticamente. Quer dizer, se $\mathbf{z}_n \xrightarrow{d} H$, então podemos usar a aproximação $\mathbb{P}(\mathbf{z}_n \leq \boldsymbol{\xi}) \approx H(\boldsymbol{\xi})$, $\boldsymbol{\xi} \in \mathbb{R}^k$:

Proposição 23. Se $\mathbf{z}_n \xrightarrow{d} H$ e H é contínua, então para qualquer $\delta > 0$ é possível encontrar um número natural n_δ tal que a

⁶¹ Lembre: uma variável aleatória z é dita ter **distribuição Normal (ou Gaussiana)** se z é degenerada ou se z possui função densidade de probabilidade dada, para $\xi \in \mathbb{R}$, por

$$f_z(\xi) = \frac{1}{s\sqrt{2\pi}} e^{-(\xi-m)^2/(2s^2)}$$

para algum par de parâmetros $m \in \mathbb{R}$ e $s > 0$. Um vetor aleatório $\mathbf{z}$ em $\mathbb{R}^k$ é dito ter **distribuição Normal multivariada** se, para todo $\boldsymbol{\xi} \in \mathbb{R}^k$, a variável aleatória $\boldsymbol{\xi}'\mathbf{z}$ tem distribuição Normal.

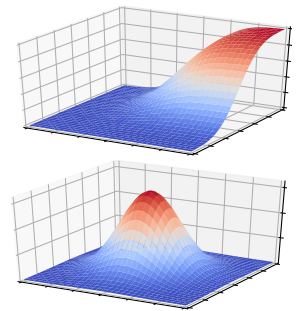


Figura 5: A função de distribuição acumulada de uma $N(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ (topo) e a respectiva função densidade de probabilidade (abaixo).

desigualdade

$$\sup_{\xi \in \mathbb{R}^d} |\mathbb{P}(z_n \leq \xi) - H(\xi)| < \delta$$

vale para todo $n \geq n_\delta$.

NESSE PONTO É IMPORTANTE RESSALTAR que o espaço amostral $(\Omega, \mathcal{A})$ que estamos considerando deve comportar tais sequências infinitas de variáveis aleatórias. Não é trivial que um tal espaço exista; o *Teorema de Existência de Kolmogorov* nos garante que não temos com o que nos preocupar — de fato, o teorema é construtivo: ver Billingsley⁶² para mais detalhes. Para nós, é suficiente enunciar a seguinte versão desse resultado:

Teorema 24. Seja $(H_n)_{n \geq 1}$ uma sequência de funções de distribuição acumuladas com

$$H_n: (\mathbb{R}^k)^n \rightarrow \mathbb{R}, \quad n \geq 1.$$

Suponha que o critério de compatibilidade

$$H_{n+1}(\xi_1, \dots, \xi_n, +\infty) = H_n(\xi_1, \dots, \xi_n)$$

é satisfeito para toda n -upla $\xi_1, \dots, \xi_n$ de vetores em $\mathbb{R}^k$. Então existe um espaço de probabilidade $(\Omega, \mathcal{A}, \mathbb{P})$ e uma sequência $(z_n)_{n \geq 1}$ de vetores aleatórios em $\mathbb{R}^k$ tais que, para todo $n \geq 1$ e toda n -upla $\xi_1, \dots, \xi_n$ de vetores em $\mathbb{R}^k$, vale a igualdade

$$\mathbb{P}(z_1 \leq \xi_1, \dots, z_n \leq \xi_n) = H_n(\xi_1, \dots, \xi_n).$$

UMA IMPORTANTE APLICAÇÃO do teorema acima é a seguinte: tomando uma função de distribuição acumulada fixa $H: \mathbb{R}^k \rightarrow \mathbb{R}$ e a partir dela definindo a sequência $(H_n)_{n \geq 1}$ via

$$H_n(\xi_1, \dots, \xi_n) := H(\xi_1) \times \dots \times H(\xi_n),$$

para $\xi_1, \dots, \xi_n \in \mathbb{R}^k$ e $n \geq 1$, o teorema nos garante que existe uma sequência iid⁶³ $(z_n)_{n \geq 1}$ de vetores aleatórios em $\mathbb{R}^k$ com $z_n \sim H$ para todo n . Quando fazemos teoria assintótica para dados *cross section*, embora na maior parte do tempo não tenhamos (nem queiramos ter) em mente qual é o espaço amostral sendo utilizado, necessitamos ao menos garantir que os objetos matemáticos com os quais estamos trabalhando sejam bem definidos — a presente aplicação do teorema de Kolmogorov supre essa necessidade.

⁶² Patrick Billingsley. *Probability and measure*. John Wiley & Sons, 3rd edition, 1995

⁶³ Uma sequência infinita de vetores aleatórios $(z_n)_{n \geq 1}$ é iid se, e somente se, para cada $n \geq 1$ os vetores aleatórios $z_1, \dots, z_n$ são iid, no sentido que vimos anteriormente.

Elementos de teoria assintótica

No próximo capítulo, estudaremos as propriedades assintóticas do estimador de mínimos quadrados ordinários. Veremos que, sob certas condições, vale que $\hat{\beta}_n \xrightarrow{p} \beta$ e que $\sqrt{n}(\hat{\beta}_n - \beta) \xrightarrow{d} N(0, \Sigma^{-1} \mathbf{S} \Sigma^{-1})$, onde a matriz de covariâncias $\Sigma^{-1} \mathbf{S} \Sigma^{-1}$ será especificada oportunamente. Note que escrevemos $\hat{\beta}_n$ ao invés de simplesmente $\hat{\beta}$ para enfatizar que “o” estimador de MQO é na verdade um termo em uma sequência de estimadores: de fato, na fórmula $\hat{\beta} = (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbf{y}$ está implícita a dependência no tamanho da amostra, através das dimensões das matrizes e vetores que aparecem nessa expressão. Veja: com as definições dadas nas duas seções anteriores, já é possível compreender as propriedades de consistência e normalidade assintótica do estimador de MQO! Para provar essas propriedades, entretanto, vamos recorrer a uma dose adicional de teoria. Enunciaremos os principais resultados que serão usados posteriormente, sem prová-los. Algumas dessas demonstrações estão espalhadas nos exercícios logo adiante; os detalhes estão no *Asymptotic Statistics* do van der Vaart⁶⁴ e no já mencionado *Econometrics* do Hayashi.⁶⁵ Os três primeiros desses resultados garantem que os modos de convergência estudados são “bem comportados” mediante transformações contínuas (em particular as operações algébricas usuais dos espaços Euclidianos), e estabelecem algumas relações entre esses modos de convergência.

Teorema 25 (Continuous mapping theorem). Sejam $\mathbf{z}_n$, $n \geq 1$, e $\mathbf{z}$ vetores aleatórios em $\mathbb{R}^k$. Suponha que h é uma função de $\mathbb{R}^k$ em $\mathbb{R}^d$ cujo conjunto de pontos de continuidade $\Theta \subseteq \mathbb{R}^k$ satisfaz $\mathbb{P}(\mathbf{z} \in \Theta) = 1$. Então vale o seguinte:

1. Se $\mathbf{z}_n \xrightarrow{p} \mathbf{z}$, então $h \circ \mathbf{z}_n \xrightarrow{p} h \circ \mathbf{z}$.
2. Se $\mathbf{z}_n \xrightarrow{d} \mathbf{z}$, então $h \circ \mathbf{z}_n \xrightarrow{d} h \circ \mathbf{z}$.

Nos teorema e corolário abaixo, estão implícitas as dimensões dos espaços correspondentes a cada um dos vetores aleatórios que ali aparecem. Essas dimensões podem variar de um item para outro de forma que as expressões façam sentido.

Teorema 26. Sejam $\mathbf{z}_n$, $\mathbf{w}_n$, $n \geq 1$, e $\mathbf{z}$ vetores aleatórios e seja ξ um vetor constante. Então vale o seguinte:

⁶⁴ A. W. van der Vaart. *Asymptotic Statistics*. Cambridge Univ. Press, 1998

⁶⁵ Fumio Hayashi. *Econometrics*. Princeton University Press, 2000

1. Se $z_n \xrightarrow{p} z$, então $z_n \xrightarrow{d} z$.
2. Se $z_n \xrightarrow{d} \xi$, então $z_n \xrightarrow{p} \xi$.
3. Se $z_n \xrightarrow{d} z$ e $|z_n - w_n| \xrightarrow{p} 0$, então $w_n \xrightarrow{d} z$.
4. Se $z_n \xrightarrow{d} z$ e $w_n \xrightarrow{p} \xi$, então $(z'_n \ w'_n)' \xrightarrow{d} (z' \ \xi')'$.

Corolário 27 (Lema de Slutsky). Nas condições do Teorema 26, se $z_n \xrightarrow{d} z$ e $w_n \xrightarrow{d} \xi$, então

1. $z_n + w_n \xrightarrow{d} z + \xi$.
2. $w'_n z_n \xrightarrow{d} \xi' z$.
3. $z_n/w_{n,1} \xrightarrow{d} z/\xi_1$ desde que $\xi_1 \neq 0$.

O método Delta

A ideia do método Delta é utilizar uma expansão de Taylor para aproximar a distribuição de um vetor aleatório da forma $h \circ z_n$. Nessa seção, para mantermo-nos fiéis à notação mais usual e a despeito da ressalva que fizemos na seção introdutória destas notas, escreveremos provisoriamente $h(z_n)$ ao invés de $h \circ z_n$. Seja, então, $h = (h_1 \ h_2 \ \dots \ h_d)'$ uma transformação com domínio $\Theta \subseteq \mathbb{R}^k$ e contradomínio $\mathbb{R}^d$. Suponha que as derivadas parciais

$$\partial_j^\ell h(\theta) := \lim_{\delta \rightarrow 0} \delta^{-1} (h_\ell(\theta + \delta e_j) - h_\ell(\theta))$$

existem e são contínuas numa vizinhança⁶⁶ $U \subseteq \Theta$ do ponto $\theta \in \Theta$ para $j \in \{1, \dots, k\}$ e $\ell \in \{1, \dots, d\}$, onde e_j é o j -ésimo elemento da base canônica de $\mathbb{R}^k$ — por exemplo,

$$e_2 = (0 \ 1 \ 0 \ \dots \ 0)'$$

etc. Defina ∂h_θ como sendo a matriz $(d \times k)$ cuja componente (j, ℓ) é $\partial_j^\ell h(\theta)$. Nessas condições, temos

$$h(\theta + \xi) = h(\theta) + \partial h_\theta \xi + r(\xi), \quad \xi \in \mathbb{R}^k, \quad (23)$$

em que $\lim_{\xi \rightarrow 0} r(\xi) = 0$.⁶⁷

No enunciado abaixo, a sequência $(a_n)_{n \geq 1}$ é usualmente tomada com $a_n = \sqrt{n}$, mas não há razão para não sermos mais generalistas:

Teorema 28 (Método Delta). Seja h como introduzida acima. Sejam $(z_n)_{n \geq 1}$ uma sequência de vetores aleatórios em $\mathbb{R}^k$, com

⁶⁶ Usaremos o termo *vizinhança* de um ponto θ no seguinte sentido: para $\delta > 0$ suficientemente pequeno, os conjuntos $\text{ball}(\theta; \delta)$ estão inteiramente contidos em U .

⁶⁷ Veja o exercício 64.

$\mathbb{P}(\mathbf{z}_n \in \Theta) = 1$ para todo n , e $(a_n)_{n \geq 1}$ uma sequência de números reais com $a_n \rightarrow +\infty$. Se $a_n(\mathbf{z}_n - \boldsymbol{\theta}) \xrightarrow{d} \mathbf{z}$, então

$$a_n(h(\mathbf{z}_n) - h(\boldsymbol{\theta})) \xrightarrow{d} \partial h_{\boldsymbol{\theta}} \mathbf{z}.$$

Ademais,

$$\left| a_n(h(\mathbf{z}_n) - h(\boldsymbol{\theta})) - \partial h_{\boldsymbol{\theta}}(a_n(\mathbf{z}_n - \boldsymbol{\theta})) \right| \xrightarrow{p} 0.$$

Os Reis do infinito

A *Lei Forte dos Grandes Números* e o *Teorema Central do Limite* são os dois principais resultados da Probabilidade Axiomática; reiam, absolutos, sobre toda a teoria assintótica. Abaixo enunciamos esses resultados, em sua versão multivariada, para sequências iid: para modelos em que assume-se um sistema de amostragem *cross section*, isso é tudo que precisaremos.

Teorema 29 (2ª Lei Forte dos Grandes Números de Kolmogorov). Seja $(\mathbf{z}_n)_{n \geq 1}$ uma sequência iid de vetores aleatórios em $\mathbb{R}^k$, com $\mathbb{E}|\mathbf{z}_n| < \infty$. Então

$$\mathbb{P}\left\{\omega \in \Omega: \lim_{n \rightarrow \infty} |\bar{\mathbf{z}}_n(\omega) - \mathbb{E}(\mathbf{z}_1)| = 0\right\} = 1,$$

em que $\bar{\mathbf{z}}_n := n^{-1} \sum_{i=1}^n \mathbf{z}_i$.

O vetor aleatório $\bar{\mathbf{z}}_n := n^{-1} \sum_{i=1}^n \mathbf{z}_i$ que aparece no teorema acima é chamado a **média amostral** (da amostra $\mathbf{z}_1, \dots, \mathbf{z}_n$). O teorema afirma que, com $\mathbb{P}$ -probabilidade 1, a média amostral converge para o valor esperado (comum) dos vetores aleatórios $\mathbf{z}_n$, $n \geq 1$. Nos exercícios abaixo o leitor será convidado a provar que, sob essas condições, vale também a convergência $\bar{\mathbf{z}}_n \xrightarrow{p} \mathbb{E}(\mathbf{z}_1)$.

Teorema 30 (Teorema Central do Limite de Lindeberg–Levy). Seja $(\mathbf{z}_n)_{n \geq 1}$ uma sequência iid de vetores aleatórios em $\mathbb{R}^k$. Suponha que $\mathbb{E}(\mathbf{z}'_1 \mathbf{z}_1) < \infty$. Então

$$\sqrt{n}(\bar{\mathbf{z}}_n - \boldsymbol{\mu}) \equiv \frac{1}{\sqrt{n}} \sum_{i=1}^n (\mathbf{z}_i - \boldsymbol{\mu}) \xrightarrow{d} \mathbf{N}(\mathbf{0}, \boldsymbol{\Sigma}), \quad (24)$$

onde $\boldsymbol{\mu} := \mathbb{E}(\mathbf{z}_1)$ e $\boldsymbol{\Sigma} := \mathbb{V}(\mathbf{z}_1)$.

O teorema acima pode não parecer tão impressionante à primeira vista, mas sua força torna-se mais evidente após um segundo olhar. Primeiramente, note que podemos reinterpretar a convergência em (24) como uma afirmação de que a média amostral $\bar{z}_n$ possui distribuição aproximadamente $N(\mu, n^{-1/2}\Sigma)$, desde que n seja “suficientemente grande”. Pensemos no caso univariado (i.e., $k = 1$) para ilustrar: nesse caso, escrevendo z_n para a primeira componente do vetor $\mathbf{z}_n$, a condição $\mathbb{E}(\mathbf{z}'_1 \mathbf{z}_1) < \infty$ se resume a $\mathbb{E}(z_1^2) < \infty$, e já sabemos que isso é equivalente à exigência de que $\sigma^2 := \mathbb{V}(z_1) < \infty$. Escrevendo ainda $\mu := \mathbb{E}(z_1)$, o Teorema Central do Limite nos dá a aproximação

$$\mathbb{P}(\bar{z}_n \leq \xi) \approx \int_{-\infty}^{\xi} \frac{1}{\sigma\sqrt{2\pi}} e^{-(t-\mu)^2/(2\sigma^2)} dt, \quad \xi \in \mathbb{R},$$

para n suficientemente grande. Veja bem, as únicas hipóteses feitas sobre a distribuição conjunta das variáveis aleatórias $z_1, \dots, z_n$ é que estas são independentes, identicamente distribuídas, e que possuem variâncias finitas. Afora isso, a distribuição marginal das variáveis aleatórias z_i é arbitrária. O surpreendente é que essa arbitrariedade não importa quando olhamos para a distribuição de probabilidades da média amostral $\bar{z}_n$: essa distribuição é aproximadamente normal. Dito de outra forma, ao tomar uma média aritmética de variáveis aleatórias quadrado-integráveis, independentes e identicamente distribuídas, ocorre que não importa qual seja a distribuição marginal subjacente, essa média necessariamente se distribuirá de acordo com uma lei muito próxima da distribuição normal!

Exercícios

Exercício 46. Sejam $\mathbf{z}_n$, $n \geq 1$ e $\mathbf{z}$ vetores aleatórios em $\mathbb{R}^k$, e seja $(a_n)_{n \geq 1}$ uma sequência de números reais positivos. Escrevemos $\mathbf{z}_n = \mathbf{z} + o_{\mathbb{P}}(a_n)$ se, e somente se, $(\mathbf{z}_n - \mathbf{z})/a_n \rightarrow 0$ em probabilidade.⁶⁸ Verifique que são equivalentes as seguintes afirmações:

1. $\mathbf{z}_n \xrightarrow{P} \mathbf{z}$.
2. $\mathbf{z}_n - \mathbf{z} \xrightarrow{P} \mathbf{0} \in \mathbb{R}^k$.
3. $|\mathbf{z}_n - \mathbf{z}| \xrightarrow{P} 0 \in \mathbb{R}$.
4. Para todo $\delta > 0$, $\mathbb{P}(|\mathbf{z}_n - \mathbf{z}| > \delta) \rightarrow 0 \in \mathbb{R}$.⁶⁹

⁶⁸ Há aqui um ligeiro abuso de notação — seria mais preciso escrever, e.g., $\mathbf{z}_n - \mathbf{z} \in o_{\mathbb{P}}(a_n)$ etc — mas, em que pese essa incorreção, um que é extremamente útil.

⁶⁹ Lembre que uma sequência

$$(\boldsymbol{\xi}_n)_{n \geq 1} = (\boldsymbol{\xi}_1, \boldsymbol{\xi}_2, \dots)$$

em $\mathbb{R}^k$ **converge** para um vetor $\boldsymbol{\xi} \in \mathbb{R}^k$ se, e somente se, para todo $\delta > 0$ existe um número natural n_δ tal que a desigualdade $|\boldsymbol{\xi}_n - \boldsymbol{\xi}| < \delta$ é válida para todo $n \geq n_\delta$.

5. Para todo par de números positivos δ e δ' existe um natural $n(\delta, \delta')$ tal que $\sup_{n \geq n(\delta, \delta')} \mathbb{P}(|z_n - z| > \delta) < \delta'$.
6. Para todo $\delta > 0$, $\mathbb{P}(|z_n - z| \leq \delta) \rightarrow 1 \in \mathbb{R}$.
7. $z_n = z + o_{\mathbb{P}}(1)$.

Exercício 47. Nas mesmas condições do exercício acima, escrevemos $z_n = z + O_{\mathbb{P}}(a_n)$ se, e somente se, para todo $\delta > 0$ existir uma constante $C_\delta > 0$ tal que

$$\mathbb{P}\left(\frac{|z_n - z|}{a_n} > C_\delta\right) < \delta$$

para todo $n \geq 1$. Mostre que

1. $o_{\mathbb{P}}(1) + o_{\mathbb{P}}(1) = o_{\mathbb{P}}(1)$.
2. $o_{\mathbb{P}}(1) + O_{\mathbb{P}}(1) = O_{\mathbb{P}}(1)$.
3. $o_{\mathbb{P}}(1) \cdot O_{\mathbb{P}}(1) = o_{\mathbb{P}}(1)$.
4. $(1 + o_{\mathbb{P}}(1))^{-1} = O_{\mathbb{P}}(1)$.
5. $o_{\mathbb{P}}(a_n) = a_n \cdot o_{\mathbb{P}}(1)$.
6. $O_{\mathbb{P}}(a_n) = a_n \cdot O_{\mathbb{P}}(1)$.
7. $o_{\mathbb{P}}(O_{\mathbb{P}}(1)) = o_{\mathbb{P}}(1)$.

(Conforme mencionado acima, há um abuso de notação aqui, e as “igualdades” representam, na verdade, uma implicação — isto é, essas igualdades devem ser lidas “da esquerda para a direita”. Preencha os detalhes relativos a notação.)

Exercício 48. *Prove que $z_n \xrightarrow{d} z$ se, e somente se, vale o seguinte: para toda função $h: \mathbb{R}^k \rightarrow \mathbb{R}$ contínua e limitada⁷⁰ vale $\lim_{n \rightarrow \infty} \mathbb{E}(h \circ z_n) = \mathbb{E}(h \circ z)$.

Exercício 49. Dizemos que $(z_n)_{n \geq 1}$ **converge quase certamente** para z , e que z é o **limite quase certo** da sequência $(z_n)_{n \geq 1}$ se, e somente se, valer

$$\mathbb{P}\left\{\omega \in \Omega: \lim_{n \rightarrow \infty} z_n(\omega) = z(\omega)\right\} = 1.$$

Notação: $z_n \xrightarrow{q.c.} z$. Já vimos que a sequência de variáveis aleatórias $(z_n)_{n \geq 1}$ do exemplo 22 converge em probabilidade para 0 (dê os detalhes se não estiver convencido). Mostre que $z_n \not\xrightarrow{q.c.} 0$ quase certamente.

⁷⁰ h ser limitada significa que $\sup_{\xi \in \mathbb{R}^k} |h(\xi)| < \infty$ ou, o que é dizer o mesmo, que existe uma constante $c > 0$ tal que $|h(\xi)| \leq c$ para todo $\xi \in \mathbb{R}^k$.

Exercício 50. Seja $q \geq 1$. Dizemos que a sequência $(z_n)_{n \geq 1}$ **converge em q -norma** para z se, e somente se, $\mathbb{E}(|z_n - z|^q) \rightarrow 0$. Notação: $z_n \xrightarrow{q} z$, “ $z_n \rightarrow z$ em q -norma” etc.

1. Mostre que, se $z_n \rightarrow z$ em q -norma, então $z_n \rightarrow z$ em probabilidade.
2. Nas mesmas condições do exemplo 22, sejam w_n as variáveis aleatórias definidas por $w_n(\omega) := 2^n \cdot z_n(\omega)$, $n \geq 1$. Verifique que $w_n \xrightarrow{p} 0$ mas $\mathbb{E}w_n \rightarrow +\infty$.⁷¹ Conclua que convergência em probabilidade não implica convergência em q -norma.

Exercício 51. Mostre que, se $z_n \xrightarrow{q.c.} z$, então $z_n \xrightarrow{p} z$.

Exercício 52. Mostre que, se $z_n \xrightarrow{p} z$ e $w_n \xrightarrow{p} w$, então

$$(z'_n \ w'_n)' \xrightarrow{p} (z' \ w')'.$$

Exercício 53. Demonstre o Teorema 25.

Exercício 54. *Demonstre o Teorema 26.

Exercício 55. Os modos de convergência introduzidos para sequências de vetores aleatórios podem ser facilmente estendidos para sequências de matrizes aleatórias, utilizando a aplicação τ do exercício 16. Seja $(A_n)_{n \geq 1}$ uma sequência de matrizes aleatórias, e seja A uma matriz não aleatória.⁷² Mostre que:

1. Se $z_n \xrightarrow{d} z$ e $A_n \xrightarrow{p} A$, então $A_n z_n \xrightarrow{d} A z$.
2. Nas condições do item acima, se adicionalmente $z \sim N(0, \Sigma)$, então $A_n z_n \xrightarrow{d} N(0, A \Sigma A')$.
3. Se $z_n \xrightarrow{d} z$ e $A_n \xrightarrow{p} A$, com A invertível, então $A_n^{-1} z_n \xrightarrow{d} A^{-1} z$ e $z'_n A_n^{-1} z_n \xrightarrow{d} z' A^{-1} z$.

Exercício 56. Prove a seguinte *lei fraca dos grandes números*: se $(z_n)_{n \geq 1}$ é uma sequência iid de vetores aleatórios em $\mathbb{R}^k$ com $\mathbb{E}(z'_n z_n) < \infty$, então $\bar{z}_n \xrightarrow{p} \mathbb{E}(z_1)$, onde $\bar{z}_n := n^{-1} \sum_{i=1}^n z_i$. Dica: comece com o caso $\mathbb{E}(z_1) = 0$ e use a *desigualdade de Markov*: se z é uma variável aleatória não negativa, então $\delta \mathbb{P}(z \geq \delta) \leq \mathbb{E}(z)$.

Exercício 57. Com a mesma notação do exercício acima (mas sem a hipótese de que a sequência é iid), prove a seguinte *lei fraca dos grandes números*: se $\mathbb{E}(\bar{z}_n) \rightarrow \mu \in \mathbb{R}^k$ e $\mathbb{V}(\bar{z}_n) \rightarrow 0$, então $\bar{z}_n \xrightarrow{p} \mu$.

⁷¹ Se $(a_n)_{n \geq 1}$ é uma sequência de números reais, escrevemos $a_n \rightarrow +\infty$ se, e somente se, $(1 + a_n)^{-1} \rightarrow 0$.

⁷² Preencha os detalhes quanto às dimensões das matrizes e vetores envolvidos. Essas dimensões podem depender de n ?

Exercício 58. Suponha que $(z_n)_{n \geq 1}$ é uma sequência iid de variáveis aleatórias com distribuição normal padrão. Verifique que $z_n \xrightarrow{d} z_1$ mas $z_n \not\rightarrow z_1$ em probabilidade.

Exercício 59. Seja $(z_n)_{n \geq 1}$ uma sequência de variáveis aleatórias com

$$\mathbb{P}(z_n \leq \xi) = \int_{-\infty}^{\xi} \frac{\Gamma((n+1)/2)}{\sqrt{n\pi}\Gamma(n/2)} (1+t^2/n)^{-(n+1)/2} dt, \quad \xi \in \mathbb{R}, n \geq 1,$$

onde Γ é a *função gamma*,

$$\Gamma(\xi) = \int_0^{\infty} t^{\xi-1} e^{-t} dt, \quad \xi \in \mathbb{R}.$$

Isto é, para $n \geq 1$, a variável aleatória z_n tem distribuição Student- t com n graus de liberdade. Mostre que $z_n \xrightarrow{d} N(0, 1)$.

Exercício 60. Em que sentido é verdade que $o_{\mathbb{P}}(1) = O_{\mathbb{P}}(1)$? E quanto à igualdade $O_{\mathbb{P}}(1) = o_{\mathbb{P}}(1)$, podemos afirmar que ela é verdadeira?

Exercício 61. Fixado $\theta \in \mathbb{R}^k$, mostre que, se $\sqrt{n}(z_n - \theta) \xrightarrow{d} z$, então $z_n \xrightarrow{p} \theta$.

Exercício 62. Fixado $\theta \in \mathbb{R}^k$, mostre que, se $\mathbb{E}(z_n) \rightarrow \theta$ e $|\mathbb{V}(z_n)| \rightarrow 0$, então $z_n \xrightarrow{p} \theta$. Faça o caso $k = 1$ primeiro.

Exercício 63. Prove o seguinte: se, para todo $\delta > 0$ a série $\sum_{n=1}^{\infty} \mathbb{P}(|z_n| > \delta)$ converge, então $z_n \xrightarrow{q.c.} \mathbf{0}$.

Exercício 64. Mostre que o resto $r(\cdot)$ na equação (23) satisfaz $\lim_{\xi \rightarrow 0} r(\xi) = \mathbf{0}$.

Exercício 65. *Prove o Teorema 28.

Exercício 66. *Mostre que, se o Teorema Central do Limite de Lindeberg-Levy é válido para $k = 1$, então o teorema é válido para todo $k \geq 1$.

Consistência e Gaussianidade assintóticas

“A skeptic, I would ask for consistency first of all.”

Sylvia Plath, *The Unabridged Journals of Sylvia Plath*.



Sylvia Plath.

CONTA A ANEDOTA que o ganhador do prêmio Nobel de Economia Clive W. J. Granger certa vez teria dito: “*If you can’t get it right as n goes to infinity you shouldn’t be in this business.*” Felizmente nossos esforços até aqui não foram em vão, e veremos que — sob certas condições — o estimador de mínimos quadrados ordinários *gets it right* ao $n \rightarrow \infty$. Enunciaremos os teoremas de consistência e Normalidade assintótica com bastante generalidade, para que o leitor tenha na manga um coringa a ser usado em cenários menos restritivos do que aquele sobre o qual essas notas se concentram (a saber, a situação em que os dados são *cross section*). Nessa seção, consideraremos que está dada uma sequência (infinita) $\{(y_i, \mathbf{x}_i)\}_{i \geq 1}$ onde, como antes, os y_i ’s são variáveis aleatórias reais e os $\mathbf{x}_i$ ’s são vetores aleatórios em $\mathbb{R}^{d+1}$. Também como antes, $\mathbf{X}$ denota a matriz $n \times (d+1)$ cuja componente ij é x_{ij} , e $\mathbf{y}$ denota o vetor $n \times 1$ cuja i -ésima componente é y_i .

Consistência

Um requerimento que parece ser o mínimo a se fazer quando desejamos desenvolver a teoria assintótica para o estimador de MQO é que a sequência $\{(y_i, \mathbf{x}_i)\}_{i \geq 1}$ seja **estacionária**,⁷³ no seguinte sentido: para qualquer $n \geq 1$, qualquer $t \geq 1$ e quaisquer subconjuntos

⁷³ A noção que introduzimos aqui é a de *estacionariedade forte*, em contraposição à noção de *estacionariedade fraca* (ou *em covariância*) que é prevalente em textos de séries temporais. Observe que a definição depende da medida de probabilidade $\mathbb{P}$.

Borelianos $A_1, A_2, \dots, A_n \subseteq \mathbb{R} \times \mathbb{R}^{d+1}$, vale a igualdade

$$\begin{aligned} \mathbb{P}((y_1, \mathbf{x}_1) \in A_1, \dots, (y_n, \mathbf{x}_n) \in A_n) \\ = \mathbb{P}((y_{1+t}, \mathbf{x}_{1+t}) \in A_1, \dots, (y_{n+t}, \mathbf{x}_{n+t}) \in A_n). \end{aligned}$$

Considerando, por conveniência, que as unidades amostrais estão indexadas no “tempo”, a identidade acima pode ser interpretada como uma exigência (ou uma suposição) de que o procedimento de amostragem⁷⁴ independa, do ponto de vista distribucional, do instante de tempo em que é iniciado. Em particular, sob estacionariedade vemos que amostras de tamanho n possuem sempre uma mesma distribuição conjunta subjacente, que os vetores aleatórios $(y_i, \mathbf{x}_i)$, $i \geq 1$, possuem uma distribuição marginal comum etc.

Seria fortuito se, com estacionariedade apenas, pudéssemos obter alguma propriedade assintótica interessante para o estimador $\hat{\beta}_n$, mas coisas interessantes raramente se revelam tão facilmente: será preciso introduzir mais estrutura na medida de probabilidade $\mathbb{P}$ se quisermos chegar a algum lugar. O próximo atributo a exigir é aquilo que chamamos de *ergodicidade* — um termo que soa misterioso mas que é, de fato, o requerimento mínimo para que se possa fazer qualquer inferência quanto à distribuição subjacente a um conjunto de dados observados: o requerimento de que, em amostras “grandes”, observemos aproximadamente $n \times \mathbb{P}((y_1, \mathbf{x}_1) \in A)$ pontos amostrais na região $A \subseteq \mathbb{R} \times \mathbb{R}^{d+1}$. Dizemos que uma sequência estacionária $\{\mathbf{z}_i\}_{i \geq 1} \equiv \{(y_i, \mathbf{x}_i)\}_{i \geq 1}$ é **ergódica** se, e somente se, para qualquer $t \geq 1$ e quaisquer subconjuntos Borelianos $A_0, A_1, \dots, A_t \subseteq \mathbb{R} \times \mathbb{R}^{d+1}$, vale

$$\frac{1}{n} \sum_{i=1}^n \mathbb{I}[\mathbf{z}_i \in A_0, \dots, \mathbf{z}_{i+t} \in A_t] \xrightarrow{q.c.} \mathbb{P}(\mathbf{z}_1 \in A_0, \dots, \mathbf{z}_{1+t} \in A_t).$$

Uma fato importante e de fácil verificação é que sequências iid são ergódicas — ver as páginas 488 e 489 em Karlin & Taylor.⁷⁵ Portanto, os resultados que enunciamos abaixo têm validade também no caso de dados *cross section*.

Com ergodicidade obtemos a consistência do estimador $\hat{\beta}_n$, mas antes de prosseguirmos convém recordar o Teorema 3, segundo o qual $\mathbb{E}(\mathbf{X}'(\mathbf{y} - \mathbf{X}\beta)) = \mathbf{0}$ desde que $\mathbb{E}(\mathbf{X}'\mathbf{X})$ seja invertível (e onde $\beta := (\mathbb{E}(\mathbf{X}'\mathbf{X}))^{-1}\mathbb{E}(\mathbf{X}'\mathbf{y})$). Como, sob estacionariedade, valem⁷⁶ as

⁷⁴ O qual pode ser puramente observacional, é claro.

⁷⁵ Samuel Karlin and Howard M. Taylor. *A First Course in Stochastic Processes*. Academic Press, Boston, 2nd edition, 1975

⁷⁶ Veja o exercício 68.

igualdades $\mathbb{E}(\mathbf{X}'\mathbf{X}) = n\mathbb{E}(\mathbf{x}_1\mathbf{x}_1')$, $\mathbb{E}(\mathbf{X}'\mathbf{y}) = n\mathbb{E}(\mathbf{x}_1y_1)$ e também

$$\boldsymbol{\beta} = (\mathbb{E}(\mathbf{x}_1\mathbf{x}_1'))^{-1}\mathbb{E}(\mathbf{x}_1y_1), \quad (25)$$

vemos, do mencionado teorema, que sob essas condições é satisfeita a identidade

$$\mathbb{E}(\mathbf{x}_i(y_i - \mathbf{x}_i'\boldsymbol{\beta})) = \mathbf{0}, \quad i \geq 1,$$

algo que no jargão é chamado de “regressores predeterminados”. Isso é importante porque há uma linha alternativa através da qual se pode conduzir o argumento até chegarmos à consistência de $\hat{\boldsymbol{\beta}}_n$: poderíamos equivalentemente colocar como uma *suposição* que valem as igualdades $y_i = \mathbf{x}_i'\mathbf{b} + v_i$, $i \geq 1$, para algum *vetor de parâmetros desconhecidos* $\mathbf{b}$ e alguma sequência $(v_i)_{i \geq 1}$ de *termos de erro não observáveis*, e após isso exigir que os regressores sejam predeterminados com respeito a essa sequência, isto é, que $\mathbb{E}(\mathbf{x}_iv_i) = \mathbf{0}$ para todo $i \geq 1$. Ocorre que, se a sequência $\{(y_i, \mathbf{x}_i)\}_{i \geq 1}$ é estacionária e se $\mathbb{E}(\mathbf{x}_1\mathbf{x}_1')$ é invertível, então essa última exigência somente pode ocorrer quando $\mathbf{b} = \boldsymbol{\beta}$, com $\boldsymbol{\beta}$ dado pela equação (25).

Estamos prontos para enunciar o teorema de consistência, mas antes de prosseguir, lembre que

$$\hat{\boldsymbol{\beta}}_n = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{y} = \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i\mathbf{x}_i'\right)^{-1} \frac{1}{n} \sum_{i=1}^n \mathbf{x}_iy_i.$$

Teorema 31 (Consistência do estimador de MQO). Suponha que $\{(y_i, \mathbf{x}_i)_{i \geq 1}\}$ é uma sequência ergódica, e que $\mathbb{E}(\mathbf{x}_1\mathbf{x}_1')$ é invertível. Então

$$\hat{\boldsymbol{\beta}}_n \xrightarrow{q.c.} (\mathbb{E}(\mathbf{x}_1\mathbf{x}_1'))^{-1}\mathbb{E}(\mathbf{x}_1y_1).$$

Sketch da demonstração. Uma adaptação do exercício 70 nos dá

$$\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i\mathbf{x}_i' \xrightarrow{P} \mathbb{E}(\mathbf{x}_1\mathbf{x}_1') \quad (26)$$

e, como o conjunto das matrizes invertíveis é um subconjunto aberto (verifique) do espaço $M((d+1) \times (d+1))$, vemos que $\mathbf{X}'\mathbf{X}$ é invertível para n suficientemente grande. Pelo mesmo exercício temos a convergência $n^{-1} \sum_{i=1}^n \mathbf{x}_iy_i \xrightarrow{P} \mathbb{E}(\mathbf{x}_1y_1)$, e isso é quase tudo que precisamos para obter $\hat{\boldsymbol{\beta}} \xrightarrow{P} \boldsymbol{\beta}$: agora resta apenas aplicar o Teorema 25 algumas vezes. Para convergência quase certa, ver o Teorema 5.6, seção 9, em Karlin & Taylor.⁷⁷ ■

⁷⁷ Samuel Karlin and Howard M. Taylor. *A First Course in Stochastic Processes*. Academic Press, Boston, 2nd edition, 1975

Normalidade assintótica

Quanto à distribuição assintótica do estimador de mínimos quadrados ordinários, temos o seguinte:

Teorema 32 (Normalidade assintótica do estimador de MQO). Nas condições do Teorema 31, seja $u_i := y_i - \mathbf{x}_i' \boldsymbol{\beta}$, $i \geq 1$, onde $\boldsymbol{\beta} := (\mathbb{E}(\mathbf{x}_1 \mathbf{x}_1'))^{-1} \mathbb{E}(\mathbf{x}_1 y_1)$. Assuma adicionalmente que

A1. $\mathbb{E}(|u_1^2 \mathbf{x}_1 \mathbf{x}_1'|) < \infty$.

A2. $\mathbb{E}(\mathbf{x}_{\ell+1} u_{\ell+1} \mid (\mathbf{x}_i, u_i), 1 \leq i \leq \ell) = \mathbf{0}$ para todo $\ell \geq 1$.

Então

$$\sqrt{n}(\hat{\boldsymbol{\beta}}_n - \boldsymbol{\beta}) \xrightarrow{d} N(\mathbf{0}, \boldsymbol{\Sigma}),$$

onde $\boldsymbol{\Sigma} := (\mathbb{E}(\mathbf{x}_1 \mathbf{x}_1'))^{-1} \mathbb{E}(u_1^2 \mathbf{x}_1 \mathbf{x}_1') (\mathbb{E}(\mathbf{x}_1 \mathbf{x}_1'))^{-1}$.

Demonstração. Ver a Proposição 2.1 no *Econometrics* do Hayashi⁷⁸ para a versão geral enunciada acima. Para o cenário iid, primeiramente observe que nesse caso (sob as suposições do Teorema 31) a condição A2 é sempre satisfeita, de forma que somente precisamos impor A1. Note também que $n^{-1/2} \sum_{i=1}^n \mathbf{x}_i u_i \xrightarrow{d} N(\mathbf{0}, \mathbb{E}(u_1^2 \mathbf{x}_1 \mathbf{x}_1'))$, pelo Teorema 30. Ademais, no presente caso podemos obter a convergência em (26) usando a Lei Forte dos Grandes Números para sequências iid. Agora basta notar que

$$\sqrt{n}(\hat{\boldsymbol{\beta}}_n - \boldsymbol{\beta}) = \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i' \right)^{-1} \frac{1}{\sqrt{n}} \sum_{i=1}^n \mathbf{x}_i u_i.$$

e aplicar os teoremas vistos no capítulo anterior para obter a convergência desejada. ■

Exercícios

Quaisquer erros nos enunciados são — por questões de consistência — intencionais.

Exercício 67. Mostre que, se a sequência $\{(y_i, \mathbf{x}_i)\}_{i \geq 1}$ é iid, então ela é estacionária.

Exercício 68. Mostre que, se a sequência $\{(y_i, \mathbf{x}_i)\}_{i \geq 1}$ é estacionária, então

⁷⁸ Fumio Hayashi. *Econometrics*. Princeton University Press, 2000

$$1. \mathbb{E}(\mathbf{X}'\mathbf{X}) = n\mathbb{E}(\mathbf{x}_1\mathbf{x}_1').$$

$$2. \mathbb{E}(\mathbf{X}'\mathbf{y}) = n\mathbb{E}(\mathbf{x}_1y_1).$$

Conclua que, sob estacionariedade, vale $\beta = (\mathbb{E}(\mathbf{x}_1\mathbf{x}_1'))^{-1}\mathbb{E}(\mathbf{x}_1y_1)$ e $\mathbb{E}(\mathbf{X}'\mathbf{u}) = n\mathbb{E}(\mathbf{x}_1u_1) = \mathbf{0}$, onde $\mathbf{u}' = (u_1 \ \cdots \ u_n)$ com $u_i := y_i - \mathbf{x}_i'\beta$, $i \geq 1$.

Exercício 69. Suponha que a sequência $\{(y_i, \mathbf{x}_i)\}_{i \geq 1}$ é ergódica. Seja A um subconjunto Boreliano de $\mathbb{R} \times \mathbb{R}^{d+1}$, e ponha $\alpha := \mathbb{P}((y_1, \mathbf{x}_1) \in A)$. Mostre que, dado $\delta > 0$ arbitrário, podemos encontrar um evento Ω_δ com $\mathbb{P}(\Omega_\delta) = 1$ e um número natural n_δ tal que, para $n \geq n_\delta$ e $\omega \in \Omega_\delta$, a proporção dos pontos

$$(y_1(\omega), \mathbf{x}_1(\omega)), \dots, (y_n(\omega), \mathbf{x}_n(\omega))$$

que encontramos na região A é algo entre $\alpha - \delta$ e $\alpha + \delta$. Isto é, sob ergodicidade, o número de pontos em um *scatterplot* que se situam numa região A qualquer é aproximadamente $n \cdot \mathbb{P}((y_1, \mathbf{x}_1) \in A)$, desde que n seja suficientemente grande.

Exercício 70. Seja $h: \mathbb{R} \times \mathbb{R}^{d+1} \rightarrow \mathbb{R}$ uma função que pode ser expressa como uma combinação linear de funções indicadoras. Isto é, existem um $k \geq 1$, subconjuntos Borelianos $A_1, \dots, A_k \subseteq \mathbb{R} \times \mathbb{R}^{d+1}$, e constantes $a_1, \dots, a_k \in \mathbb{R}$ tais que

$$h(\xi, \zeta) = \sum_{\ell=1}^k a_\ell \cdot \mathbb{I}[(\xi, \zeta) \in A_\ell], \quad (\xi, \zeta) \in \mathbb{R} \times \mathbb{R}^{d+1}.$$

Mostre que, se a sequência $\{\mathbf{z}_i\}_{i \geq 1} \equiv \{(y_i, \mathbf{x}_i)\}_{i \geq 1}$ é ergódica, então

$$\frac{1}{n} \sum_{i=1}^n h(\mathbf{z}_i) \xrightarrow{q.c.} \mathbb{E}h(\mathbf{z}_1).$$

Com algum esforço adicional, verifique que $\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i \xrightarrow{p} \mathbb{E}\mathbf{z}_1$, desde que $\mathbb{E}|\mathbf{z}_1| < \infty$.

Exercício 71. Dê os detalhes da demonstração do Teorema 31 (apenas para o caso de convergência em probabilidade).

Exercício 72. Dê os detalhes da demonstração do Teorema 32 (apenas para o caso iid).

Exercício 73. Mostre que a condição

$$A2'. \mathbb{E}(\mathbf{x}_{\ell+1}u_{\ell+1} \mid \mathbf{x}_{\ell+1}, (\mathbf{x}_i, u_i), 1 \leq i \leq \ell) = \mathbf{0} \text{ para todo } \ell \geq 1,$$

implica a condição A2.

Exercício 74. Nas condições do Teorema 32, se $\mathbb{P}(x_{1,0} = 1) = 1$ então $\mathbb{E}(u_{\ell+1} \mid u_\ell, \dots, u_1) = 0$ para $\ell \geq 1$. Conclua que, nessas condições, os termos de erro $(u_i)_{i \geq 1}$ não apresentam autocorrelação serial.

Exercício 75. Verifique que, se $\mathbb{E}(|\mathbf{x}_1|^4)$ e $\mathbb{E}(|u_1|^4)$ são ambas finitas, então vale A1 no Teorema 32.

Exercício 76. Compare a variância assintótica de $\hat{\boldsymbol{\beta}}_n$ dada no Teorema 32 com a variância em amostras finitas que havíamos obtido no Teorema 19.

Sob a tese...

“El diccionario se basa en la hipótesis, obviamente no comprobada, de que los idiomas están compuestos de sinónimos equivalentes.”

J. L. Borges.

Introdução

Uma *hipótese estatística* nada mais é do que uma sentença (formal) envolvendo parâmetros de um modelo. Por exemplo, definindo $\beta := (\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1))^{-1} \mathbb{E}(\mathbf{x}_1 y_1)$, fixada uma aplicação $h: \mathbb{R}^{d+1} \rightarrow \mathbb{R}^k$, e sendo dado $\theta \in \mathbb{R}^k$ arbitrário, a sentença

$$“h(\beta) = \theta” \quad (H_0)$$

é uma hipótese estatística, usualmente chamada de *hipótese nula*. Note que, embora não explicitamente, a sentença H_0 acima refere-se à medida de probabilidade $\mathbb{P}$.⁷⁹ Por essa razão, vamos adotar um ligeiro abuso de notação aqui e identificar cada sentença com o conjunto que ela define, isto é, escreveremos, e.g.

$$H_0 \equiv \{\mathbb{P} \in H_* : H_0 \text{ é satisfeita por } \mathbb{P}\}$$

etc, onde H_* denota uma hipótese/suposição que está implícita na discussão⁸⁰ (a qual assumimos como verdadeira e que não será testada). Por exemplo, H_* pode ser a suposição de que a amostra é iid, de que as variáveis aleatórias têm variâncias finitas etc.

Um *teste de hipóteses* é uma regra de decisão cujo espaço de escolha, $\mathcal{E}$, é binário; uma dessas escolhas é interpretada como *rejeição da hipótese* H_0 e a outra como *não-rejeição da hipótese* H_0 .



J. L. Borges, 1969.

⁷⁹ Pois β está definido em termos do operador $\mathbb{E}$, o qual por sua vez é definido em termos de $\mathbb{P}$.

⁸⁰ H_* exerce aqui o papel de *domínio de discurso* e garante que o conjunto H_0 especificado pelo *axioma de compreensão* está bem definido.

A regra de decisão é representada por uma função $D: \Omega \rightarrow \mathcal{E}$, em geral expressa na forma⁸¹

$$D(\omega) = \mathbb{I}(\mathbf{z}_n(\omega) \in R_\alpha), \quad \omega \in \Omega,$$

onde $\mathbf{z}_n$ é uma *estatística de teste*⁸² e onde, para um *nível de significância* $0 < \alpha < 1$, a *região de rejeição* $R_\alpha \subseteq \mathbb{R}^k$ é escolhida de tal forma que valha a implicação

$$\mathbb{P} \in H_* \wedge H_0 \Rightarrow \mathbb{P}(\mathbf{z}_n \in R_\alpha) \approx \alpha. \quad (27)$$

A pergunta que se apresenta naturalmente aqui é: “de onde vem a região R_α ?” Ora, se pudermos obter, assumindo $H_* \wedge H_0$, a distribuição do vetor aleatório $\mathbf{z}_n$, então é claro que podemos encontrar uma região R_α como acima: basta tomar $\xi > 0$ suficientemente grande e definir R_α implicitamente através da igualdade de eventos

$$(\mathbf{z}_n \in R_\alpha) = (|\mathbf{z}_n| > \xi).$$

De fato, se pudermos (como dito) obter a distribuição de $\mathbf{z}_n$ sob $H_* \wedge H_0$, então é evidente que a probabilidade do evento acima converge para zero ao $\xi \rightarrow \infty$ e, em particular, se a mencionada distribuição for contínua, é possível até mesmo encontrar um valor ξ_α para o qual a referida probabilidade é *igual* a α .

MAS SOMOS HUMANOS e, como tais, através de exemplos é que pensamos. Vejamos um, para ilustrar: em um contexto de regressão linear, consideremos a hipótese⁸³

$$H_0 = \{\mathbb{P} \in H_* : \beta_1 = 0\} \quad (28)$$

O Teorema 32 nos garante que $\sqrt{n}(\hat{\beta}_n - \beta) \xrightarrow{d} N(\mathbf{0}, \Sigma)$, desde que as suposições (H_*) no seu enunciado estejam satisfeitas. Em particular, se esse for o caso e *também valer* H_0 , então a variável aleatória

$$z_n := \frac{\sqrt{n}\hat{\beta}_{n1}}{\sqrt{\sigma_{11}}}, \quad (29)$$

onde $\sigma_{\ell j}$ denota a componente (ℓ, j) da matriz Σ , tem distribuição limite aproximadamente Normal padrão. Aqui, utilizando a aproximação Normal, podemos (por exemplo) fixar $0 < \alpha < 1$ e definir

⁸¹ Sem perda de generalidade, podemos identificar $\mathcal{E} \equiv \{0, 1\}$.

⁸² Em geral, estatísticas de teste são variáveis (ou vetores) aleatórias que dependem de ω somente através da amostra $(\mathbf{y}(\omega), \mathbf{X}(\omega))$. Por isso, a regra de decisão D também é, em geral, da forma $D = \phi(\mathbf{y}, \mathbf{X})$ para alguma função ϕ .

⁸³ Nesse contexto, onde $y_1 = \sum_{j=0}^d \beta_j x_{1j} + u_1$, tal hipótese é comumente interpretada como a afirmação de que “ $x_{1,1}$ não exerce efeito linear sobre y_1 .”

nossa regra de decisão via

$$D(\omega) = \begin{cases} \text{rejeitar,} & \text{se } z_n(\omega) < -\xi_\alpha \text{ ou } z_n(\omega) > \xi_\alpha; \\ \text{não rejeitar,} & \text{caso contrário,} \end{cases}$$

onde $\xi_\alpha > 0$ é a única solução para a equação $F(\xi_\alpha) = 1 - \alpha/2$ e onde F denota aqui a função de distribuição acumulada de uma variável aleatória Normal padrão. Quer dizer, nesse exemplo nossa região de rejeição é dada pela união de dois intervalos: $R_\alpha = (-\infty, -\xi_\alpha) \cup (\xi_\alpha, +\infty)$. Se o leitor ainda anseia por concretude, podemos tomar $\alpha = 0.01$, como é costumeiro, e nesse caso temos $\xi_\alpha \approx 2.575829$, o que nos leva a rejeitar a hipótese nula se (e somente se), $|z_n^{\text{obs}}| > 2.575829$.

O teste descrito acima baseia-se em uma aproximação, e portanto é chamado de um *teste assintótico*.⁸⁴ Por essa razão, o nível α nesse contexto é chamado de *nível de significância nominal* do teste (em contraposição a um nível exato que seria obtível se conhecêssemos a distribuição *exata* de z_n). Para justificar com mais detalhes a aproximação empregada acima, recorreremos à Proposição 23, segundo a qual podemos tomar $\delta > 0$ tão pequeno quanto queiramos e tal que, para todo n suficientemente grande ($n \geq n_\delta$), valham as quotas

$$F(\xi) - \delta \leq \mathbb{P}(z_n \leq \xi) \leq F(\xi) + \delta, \quad \xi \in \mathbb{R}.$$

Com algumas linhas de conta e usando a simetria da distribuição Normal, as duas desigualdades acima nos dão, para $\xi \geq 0$,

$$2(1 - F(\xi) - 2\delta) \leq \mathbb{P}(|z_n| > \xi) \leq 2(1 - F(\xi) + 2\delta). \quad (30)$$

Assim, tomando ξ_α como anteriormente, a equação (30) nos dá

$$\mathbb{P}(|z_n| > \xi_\alpha) = \alpha + r(\delta),$$

com $|r(\delta)| \leq 4\delta$, para qualquer $\alpha \in (0, 1)$ e todo n suficientemente grande.⁸⁵ Segue que, para $\mathbb{P} \in \mathbf{H}_* \wedge \mathbf{H}_0$ e $n \geq n_\delta$, temos

$$\mathbb{P}(\text{rejeitar } H_0) \approx \alpha.$$

O teste que acabamos de descrever se resume à regra de decisão “rejeitar se, e somente se, $|z_n^{\text{obs}}| > \xi_\alpha$ ”. Com efeito, um teste de hipóteses pode ser interpretado como um algoritmo, constituído das seguintes etapas:

⁸⁴ É importante ressaltar também que esse é um exemplo de teste *bilateral*.

⁸⁵ A questão sobre “quão grande deve ser o tamanho da amostra para que valham as referidas aproximações?” é interessantíssima, mas foge da nossa linha de exposição. O leitor interessado pode pesquisar sobre o Teorema de Berry–Esseen e também sobre expansões de Edgeworth. É importante também perceber que o n_δ depende, implicitamente, da medida de probabilidade $\mathbb{P}$.

1. Formule uma hipótese H_0 , declare um nível de significância $\alpha \in (0, 1)$, determine a região de rejeição R_α correspondente e escolha uma estatística de teste $z_n: \Omega \rightarrow \mathbb{R}^k$.
2. Colete dados (isto é, observe o resultado ω do experimento ou estudo observacional⁸⁶) e calcule o valor correspondente da estatística de teste $z_n^{\text{obs}} \equiv z_n(\omega)$.
3. Se $z_n^{\text{obs}} \in R_\alpha$, rejeite a hipótese H_0 ; caso contrário, não rejeite.

Esse procedimento não requer, a princípio, qualquer teoria que o justifique.⁸⁷ Todavia, a heurística segundo a qual propõe-se a operacionalização acima é a seguinte:

1. Num primeiro momento, escolhemos a região de rejeição R_α de tal forma que a implicação (27) valha. Na prática, a região de rejeição tipicamente delimita uma magnitude a partir da qual a estatística de teste é considerada “grande demais”, ou seja, temos $z_n^{\text{obs}} \in R_\alpha$ se, e somente se, $|z_n^{\text{obs}}| > \xi_\alpha$ para algum $\xi_\alpha > 0$ o qual é escolhido de tal forma que a sentença formal

$$\mathbb{P} \in H_* \wedge H_0 \Rightarrow \mathbb{P}(|z_n| > \xi_\alpha) \approx \alpha$$

seja verdadeira.

2. Se observarmos $z_n^{\text{obs}} \in R_\alpha$ e se a hipótese H_0 for verdadeira, então essa observação pode ser considerada *atípica* (ou *extrema*) no sentido de que, sob $H_* \wedge H_0$, é uma observação situada numa região que está na “cauda” da distribuição de probabilidades da estatística de teste, conforme a implicação no item anterior (lembre que α é tipicamente um número “pequeno”). Por essa razão — e supondo que não queiramos (ou que não tenhamos motivos para) abrir mão de H_* — consideramos nessas condições que a hipótese H_0 é “inverossímil” e que é “razoável” rejeitá-la.⁸⁸
3. Se, por outro lado, observamos $z_n^{\text{obs}} \notin R_\alpha$, então consideramos que “não há evidências suficientes para rejeitar a hipótese nula”, no sentido de que essa observação é “compatível” com H_0 :

$$\mathbb{P} \in H_* \wedge H_0 \Rightarrow \mathbb{P}(z_n \notin R_\alpha) \approx 1 - \alpha.$$

Com isso encerramos essa breve (e incompleta) recapitulação da teoria de testes de hipóteses. O leitor pode encontrar uma exposição

⁸⁶ Ou, se a partir do experimento/procedimento observacional não for possível determinar o valor ω , observe por exemplo $(y(\omega), X(\omega))$ etc.

⁸⁷ De fato, do ponto de vista algorítmico, não há qualquer relação *a priori* entre a hipótese e a região de rejeição. Tampouco há qualquer menção a probabilidades.

⁸⁸ A sentença $\neg H_0$ é usualmente chamada de *hipótese alternativa* e denotada por H_1 . No jargão, dizemos que “rejeitamos H_0 em favor de H_1 ”.

suficientemente rigorosa e acessível no livro clássico de Casella & Berger.⁸⁹ Agora, vamos passar ao problema de encontrar um estimador consistente para Σ .

Um estimador para a covariância assintótica

Nosso objetivo nessa seção é propor um estimador $\hat{\Sigma}_n$ satisfazendo $\hat{\Sigma}_n \xrightarrow{p} \Sigma$, isto é, queremos encontrar um estimador consistente para a matriz de covariância assintótica⁹⁰ do estimador de mínimos quadrados ordinários, a qual é dada por

$$\Sigma := (\mathbb{E}(\mathbf{x}_1 \mathbf{x}_1'))^{-1} \mathbb{E}(u_1^2 \mathbf{x}_1 \mathbf{x}_1') (\mathbb{E}(\mathbf{x}_1 \mathbf{x}_1'))^{-1}. \quad (31)$$

A motivação é a seguinte: o Teorema 32 nos dá, sob certas condições, a convergência $\sqrt{n}(\hat{\beta}_n - \beta) \xrightarrow{d} N(\mathbf{0}, \Sigma)$. Em particular, se Σ é positiva definida, então vale que

$$\sqrt{n} \Sigma^{-\frac{1}{2}} (\hat{\beta}_n - \beta) \xrightarrow{d} N(\mathbf{0}, \text{Id}), \quad (32)$$

onde $\Sigma^{\frac{1}{2}}$ é a única matriz positiva definida $\mathbf{A}$ que satisfaz $\mathbf{A}^2 = \Sigma$. Nessas condições, se pudermos encontrar um estimador $\hat{\Sigma}_n$ como acima, então uma simples aplicação de um dos resultados derivados no exercício 55 nos garante que podemos substituir Σ por $\hat{\Sigma}_n$ em (32), e com essa substituição obtemos uma estatística de teste cujo valor é computável sempre que estivermos diante de uma hipótese nula da forma $H_0 = \{\mathbb{P}: \beta = \mathbf{b}\}$.

Agora, uma rápida inspeção da expressão (31) já sugere (ou, a essa altura, já deveria sugerir) como estimar Σ : temos dois “ingredientes”, o primeiro deles sendo $n^{-1} \mathbf{X}' \mathbf{X}$, o qual satisfaz

$$\frac{1}{n} \mathbf{X}' \mathbf{X} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i' \xrightarrow{p} \mathbb{E}(\mathbf{x}_1 \mathbf{x}_1')$$

se $(\mathbf{x}_i)_{i \geq 1}$ é ergódica (com respeito a $\mathbb{P}$).⁹¹ O segundo desses ingredientes seria, idealmente, a média amostral $n^{-1} \sum_{i=1}^n u_i^2 \mathbf{x}_i \mathbf{x}_i'$, pois sob ergodicidade temos

$$\frac{1}{n} \sum_{i=1}^n u_i^2 \mathbf{x}_i \mathbf{x}_i' \xrightarrow{p} \mathbb{E}(u_1^2 \mathbf{x}_1 \mathbf{x}_1'),$$

desde que $\mathbb{E}|u_1^2 \mathbf{x}_1 \mathbf{x}_1'| < \infty$. O problema aqui é que os termos de erro u_i não são observáveis e, portanto, o estimador que aparece

⁸⁹ George Casella and Roger L Berger. *Statistical inference*, volume 2. Duxbury Pacific Grove, CA, 2002

⁹⁰ O leitor mais atento certamente deu-se conta de que a estatística de teste (29) não é computável, a menos que o desvio padrão σ_{11} seja conhecido.

⁹¹ E, note, vale também que $(\frac{1}{n} \mathbf{X}' \mathbf{X})^{-1}$ converge em probabilidade para $(\mathbb{E}(\mathbf{x}_1 \mathbf{x}_1'))^{-1}$.

no lado esquerdo na expressão acima não é computável a partir da amostra $(\mathbf{y}, \mathbf{X})$. Bom, na verdade isso não é de fato um problema, pois temos o seguinte:

Teorema 33. Suponha que $\{(y_i, \mathbf{x}_i)_{i \geq 1}\}$ é uma sequência ergódica satisfazendo as seguintes condições:

1. $\mathbb{E}(|\mathbf{x}_1 \mathbf{x}'_1|) < \infty$.
2. $\mathbb{E}(|u_1^2 \mathbf{x}_1 \mathbf{x}'_1|) < \infty$.

Então⁹² $n^{-1} \sum_{i=1}^n \hat{u}_i^2 \mathbf{x}_i \mathbf{x}'_i \xrightarrow{P} \mathbb{E}(u_1^2 \mathbf{x}_1 \mathbf{x}'_1)$.

Corolário 34. Nas condições do teorema acima, suponha adicionalmente que $\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1)$ é invertível. Então

$$\hat{\Sigma}_n \xrightarrow{P} \Sigma,$$

onde $\hat{\Sigma}_n := n(\mathbf{X}'\mathbf{X})^{-1}(\sum_{i=1}^n \hat{u}_i^2 \mathbf{x}_i \mathbf{x}'_i)(\mathbf{X}'\mathbf{X})^{-1}$.

Exemplo 35. Vamos explorar novamente, agora com um pouco mais de generalidade, o exemplo visto na seção introdutória deste capítulo. Seja $\hat{\sigma}_{\ell j}^2$ a componente (ℓ, j) da matriz $\hat{\Sigma}_n$ introduzida acima.⁹³ Consideremos a hipótese nula $H_0 = \{\mathbb{P}: \beta_k = b\}$, onde estão fixados $0 \leq k \leq d$ e $b \in \mathbb{R}$. Satisfeitas as devidas hipóteses (H_*) temos, do Teorema 32, que

$$\mathbb{P} \in H_* \wedge H_0 \Rightarrow \frac{\sqrt{n}(\hat{\beta}_{nk} - b)}{\hat{\sigma}_{kk}} \xrightarrow{d} N(0, 1).$$

A importância disso reside, como dissemos, no fato de que a estatística de teste acima é da forma $z_n(\omega) = h(\mathbf{y}(\omega), \mathbf{X}(\omega), b)$, onde h é uma função “universal”, no sentido de que independe de quaisquer parâmetros associados a $\mathbb{P}$. Isso quer dizer que *podemos computar o valor $z_n(\omega)$ para qualquer amostra $(\mathbf{y}(\omega), \mathbf{X}(\omega))$ e qualquer valor hipotético $b \in \mathbb{R}$* .⁹⁴

O exemplo acima é útil pois, como já mencionamos anteriormente, pode ser do interesse do pesquisador testar a hipótese de que um particular regressor, digamos x_{1k} , “não tem efeito linear sobre a resposta” ($H_0 = \{\mathbb{P}: \beta_k = 0\}$). Todavia, uma situação mais realista é aquela em que gostaríamos de testar hipóteses sobre mais de um dos coeficientes em β (possivelmente todos); aí é preciso ter em mente que — conforme se aprende após alguma prática no

⁹² Nesse enunciado, adotamos a mesma notação já usada anteriormente: definimos $\hat{u}_i := y_i - \mathbf{x}'_i \hat{\beta}$ e $u_i := y_i - \mathbf{x}'_i \beta$, onde $\hat{\beta}$ é o estimador de mínimos quadrados ordinários e onde $\beta = [\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1)]^{-1} \mathbb{E}(\mathbf{x}_1 y_1)$.

⁹³ Omitiremos o índice n em $\hat{\sigma}_{\ell j}$ para evitar que a notação fique excessivamente carregada — o leitor pode considerar que, sempre que o símbolo “chapéu” der as caras na notação, implicitamente há um n ali escondido.

⁹⁴ Aqui parece ser um bom lugar para introduzir mais um jargão: a variável aleatória

$$\widehat{\text{se}}(\hat{\beta}_k) := \frac{\hat{\sigma}_{kk}}{\sqrt{n}}$$

é dita o **erro padrão de White**, ou o **erro padrão robusto a heteroscedasticidade**. Trata-se de um estimador para o desvio padrão $\sqrt{\mathbb{V}(\hat{\beta}_k)}$.

“jogo” de teste de hipóteses — *não se deve testar múltiplas hipóteses isoladamente*. O teorema abaixo nos diz como proceder:

Teorema 36. Suponha que $\{(y_i, \mathbf{x}_i)_{i \geq 1}\}$ é uma sequência ergódica satisfazendo as seguintes condições ($H_* = \text{ergodicidade} \wedge a \wedge b \wedge c \wedge d$):

- (a) $\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1)$ é invertível.
- (b) $\mathbb{E}(|u_1^2 \mathbf{x}_1 \mathbf{x}'_1|) < \infty$.
- (c) $\mathbb{E}(\mathbf{x}_{\ell+1} u_{\ell+1} \mid (\mathbf{x}_i, u_i), 1 \leq i \leq \ell) = \mathbf{0}$ para todo $\ell \geq 1$.
- (d) $\mathbb{E}(u_1^2 \mathbf{x}_1 \mathbf{x}'_1)$ é invertível,

onde $u_i := y_i - \mathbf{x}'_i \boldsymbol{\beta}$, com $\boldsymbol{\beta} = (\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1))^{-1} \mathbb{E}(\mathbf{x}_1 y_1)$. Sejam ainda $\mathbf{a}$ um vetor de $\mathbb{R}^k$ arbitrário mas fixado e $\mathbf{A}$ uma matriz não aleatória de dimensões $k \times (d+1)$, com $\text{rank}(\mathbf{A}) = k$, e considere a hipótese nula $H_0 = \{\mathbb{P}: \mathbf{A}\boldsymbol{\beta} = \mathbf{a}\}$. Então, para toda $\mathbb{P} \in H_* \wedge H_0$, vale que

$$\hat{w}_n \xrightarrow{d} \chi^2(k),$$

onde $\hat{w}_n := n(\mathbf{A}\hat{\boldsymbol{\beta}} - \mathbf{a})'(\mathbf{A}\hat{\boldsymbol{\Sigma}}_n \mathbf{A}')^{-1}(\mathbf{A}\hat{\boldsymbol{\beta}} - \mathbf{a})$ e onde $\chi^2(k)$ denota a distribuição qui-quadrado com k graus de liberdade.

Exercícios

Quaisquer erros nos enunciados são — por hipótese — intencionais.

Exercício 77. Verifique a validade da equação (30).

Exercício 78. Justifique por que a convergência em (32) é válida.

Exercício 79. Mostre que, sob ergodicidade, $n^{-1} \hat{\mathbf{u}}' \hat{\mathbf{u}} \xrightarrow{P} \mathbb{E}(u_1^2)$.

Exercício 80. Suponha que $\{(y_n, \mathbf{x}_n)\}_{n \geq 1}$ é uma sequência iid satisfazendo as seguintes condições:

- 1. $\mathbb{E}(y_1^4) < \infty$.
- 2. $\mathbb{E}(|\mathbf{x}_1|^r) < \infty$ para algum $r \geq 4$.
- 3. $\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1)$ é positiva definida.

Mostre que

$$\max_{1 \leq i \leq n} |\hat{u}_i - u_i| = o_{\mathbb{P}}(\sqrt[4]{n}/\sqrt{n}).$$

Exercício 81. Demonstre o Teorema 33 no caso em que $d = 0$ (uma covariável, sem termo constante). A norma que aparece nos itens 2. e 3. do teorema pode ser tomada como qualquer norma matricial.

Exercício 82. Demonstre o Corolário 34.

Exercício 83. Demonstre o Teorema 36.

Exercício 84. Nas condições do exemplo 35, mostre que

$$\mathbb{P} \in H_* \wedge (\neg H_0) \Rightarrow \lim_{n \rightarrow \infty} \mathbb{P} \left(\left| \sqrt{n} \hat{\sigma}_{kk}^{-1} (\hat{\beta}_{nk} - b) \right| > \xi_\alpha \right) = 1,$$

onde ξ_α satisfaz $\mathbb{P}(N(0, 1) \leq \xi_\alpha) = 1 - \alpha/2$, $\alpha \in (0, 1)$. A função

$$\text{power}(\mathbb{P}) := \mathbb{P} \left(\left| \sqrt{n} \hat{\sigma}_{kk}^{-1} (\hat{\beta}_{nk} - b) \right| > \xi_\alpha \right), \quad \mathbb{P} \in H_* \wedge (\neg H_0),$$

é dita *poder* do teste.

Exercício 85. Prove o seguinte resultado:

Teorema 37. Suponha que $\{(y_i, \mathbf{x}_i)_{i \geq 1}\}$ é uma sequência ergódica satisfazendo as seguintes condições ($H_* = \text{ergodicidade} \wedge a \wedge b \wedge c \wedge d$):

- (a) $\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1)$ é invertível.
- (b) $\mathbb{E}(|u_1^2 \mathbf{x}_1 \mathbf{x}'_1|) < \infty$.
- (c) $\mathbb{E}(\mathbf{x}_{\ell+1} u_{\ell+1} \mid (\mathbf{x}_i, u_i), 1 \leq i \leq \ell) = \mathbf{0}$ para todo $\ell \geq 1$.
- (d) $\mathbb{E}(u_1^2 \mathbf{x}_1 \mathbf{x}'_1)$ é invertível,

onde $u_i := y_i - \mathbf{x}'_i \boldsymbol{\beta}$, com $\boldsymbol{\beta} = (\mathbb{E}(\mathbf{x}_1 \mathbf{x}'_1))^{-1} \mathbb{E}(\mathbf{x}_1 y_1)$. Seja ainda $h: \mathbb{R}^{d+1} \rightarrow \mathbb{R}^k$ uma função cujas derivadas parciais existem e tal que a matriz de derivadas parciais $\partial h_{\boldsymbol{\beta}}$ satisfaça $\text{rank}(\partial h_{\boldsymbol{\beta}}) = k$. Considere a hipótese nula $H_0 = \{\mathbb{P}: h(\boldsymbol{\beta}) = \mathbf{0}\}$. Então, para toda $\mathbb{P} \in H_* \wedge H_0$, vale que

$$\hat{w}_n \xrightarrow{d} \chi^2(k),$$

onde $\hat{w}_n := nh(\hat{\boldsymbol{\beta}})'(\partial h_{\hat{\boldsymbol{\beta}}} \hat{\boldsymbol{\Sigma}}_n \partial h'_{\hat{\boldsymbol{\beta}}})^{-1} h(\hat{\boldsymbol{\beta}})$ e onde $\chi^2(k)$ denota a distribuição qui-quadrado com k graus de liberdade.

Exercício 86. Proponha um intervalo de confiança para β_k . Dê os detalhes, em particular indicando as propriedades assintóticas do intervalo proposto. Você conseguiria propor uma *região* de confiança para $\boldsymbol{\beta}$? E um intervalo de confiança (condicional) para $\boldsymbol{\xi}'\boldsymbol{\beta}$?

Referências Bibliográficas

- Patrick Billingsley. *Probability and measure*. John Wiley & Sons, 3rd edition, 1995.
- Olav Bjerkholt. On the Founding of the Econometric Society. *Journal of the History of Economic Thought*, 39(2), 2017.
- George Box. Robustness in the Strategy of Scientific Model Building. In Robert L. Launer and Graham N. Wilkinson, editor, *Robustness in Statistics*, pages 201 – 236. Academic Press, 1979.
- George Casella and Roger L Berger. *Statistical inference*, volume 2. Duxbury Pacific Grove, CA, 2002.
- John Geweke, Joel L. Horowitz, and Hashem Pesaran. *Econometrics*, pages 1–44. Palgrave Macmillan UK, London, 2017.
- Trygve Haavelmo. The Probability Approach in Econometrics. *Econometrica*, 12:iii–115, 1944.
- Fumio Hayashi. *Econometrics*. Princeton University Press, 2000.
- David F. Hendry. Econometrics – Alchemy or Science? *Economica*, 47(188):387–406, 1980.
- Kevin D. Hoover. The Methodology of Econometrics. In Terence C. Mills and Kerry Patterson, editors, *Palgrave Handbook of Econometrics*, volume 1. Palgrave Macmillan UK, 2006.
- Joseph B. Kadane. *Principles of Uncertainty*. Chapman and Hall/CRC, 2011.
- Samuel Karlin and Howard M. Taylor. *A First Course in Stochastic Processes*. Academic Press, Boston, 2nd edition, 1975.

R.L. Plackett. The discovery of the method of least squares. *Biometrika*, 59(2), 1972.

Juliet Popper Shaffer. The Gauss-Markov Theorem and random regressors. *The American Statistician*, 45(4):269–273, 1991.

John Stachurski. *A primer in econometric theory*. MIT Press, 2016.

A. W. van der Vaart. *Asymptotic Statistics*. Cambridge Univ. Press, 1998.

David Williams. *Weighing the Odds: A Course in Probability and Statistics*. Cambridge University Press, 2001.