

Índice general

1	De dónde venimos: un plano, un problema sin resolver	2
2	Primera condición: el residuo tiene media cero	3
3	Segunda condición: el residuo ortogonal a x	4
4	$\hat{\beta}_2$: covarianza sobre varianza	6
5	Vía alternativa: generadores ortogonales	7
6	$\hat{\beta}_1$ y el lugar de las medias	8
7	Ejemplo numérico completo	10
8	Recapitulación y guiño a la sesión siguiente	12

Lección 7. Estimación MCO sin matrices: dos ortogonalidades, dos coeficientes

Marcos Bujosa

19 de septiembre de 2026

Resumen

Hoy resolvemos el sistema de dos ortogonalidades planteado en la lección anterior: de él surgen las fórmulas de $\hat{\beta}_2$ y $\hat{\beta}_1$. La pendiente resulta ser la covarianza dividida por la varianza —o, equivalentemente, la correlación reescalada—. Por ahora sin matrices de por medio.

“Imponer perpendicularidad, dos veces, y despejar: así nace la recta de mínimos cuadrados.”

- ([slides](#)) — ([html](#)) — ([pdf](#))

1 De dónde venimos: un plano, un problema sin resolver

En la lección 6 planteamos la regresión simple como una proyección sobre el *plano* de todas las combinaciones lineales de $\mathbf{1}$ y $\mathbf{x}$, es decir, sobre $\mathcal{L}(\mathbf{1}, \mathbf{x})$:

$$\mathbf{y} = \hat{\beta}_1 \mathbf{1} + \hat{\beta}_2 \mathbf{x} + \hat{\mathbf{e}} = \hat{\mathbf{y}} + \hat{\mathbf{e}}.$$

- $\hat{\mathbf{e}}$ debe ser ortogonal a *todo* el plano.
- Eso equivale a ser ortogonal a cada uno de sus generadores:

$$\langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = 0, \quad \langle \hat{\mathbf{e}}, \mathbf{x} \rangle_s = 0.$$

En la lección 6 nos quedamos ahí: dos condiciones, ningún coeficiente despejado.

Hoy resolvemos ese sistema. Sin matrices: dos ecuaciones, dos incógnitas ($\hat{\beta}_1, \hat{\beta}_2$).

De dónde venimos: un plano, un problema sin resolver

Conviene recordar el camino recorrido antes de avanzar. En la lección 4 resolvimos el problema más sencillo posible de proyección: buscar, dentro de la recta $\mathcal{L}(\mathbf{1})$, el vector más cercano a $\mathbf{y}$. La condición de ortogonalidad del residuo¹ frente al único generador $\mathbf{1}$ nos dio, sin apenas esfuerzo, la media $\mu_{\mathbf{y}}$. En la lección 6 dimos el siguiente paso: añadimos un segundo generador, $\mathbf{x}$, y pasamos de una recta a un *plano*, $\mathcal{L}(\mathbf{1}, \mathbf{x})$. La lógica no cambió en absoluto —seguíamos buscando el vector

Licencia: Creative Commons Attribution-ShareAlike 4.0 International (CC BY-SA 4.0).

¹que en aquel caso denominamos vector en desviaciones

del subespacio más cercano a $\mathbf{y}$ —, pero ahora la condición de ortogonalidad del residuo frente a *todo el plano* se traduce en *dos* condiciones, una por cada generador:

$$\langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = 0, \quad \langle \hat{\mathbf{e}}, \mathbf{x} \rangle_s = 0.$$

La lección 6 dejó planteado este sistema —explicando por qué debía cumplirse— pero no lo resolvió. Esa es precisamente la tarea de hoy: determinar $\hat{\beta}_1$ y $\hat{\beta}_2$ a partir de esas dos ecuaciones, sin recurrir aún a ninguna maquinaria matricial (que llegará en la lección 10, y solo como notación compacta para expresar esto mismo).

Merece la pena insistir en que no hay ningún concepto nuevo en esta sesión. Todo lo que necesitamos —el producto escalar, la ortogonalidad, la media, la varianza, la covarianza— ya fue introducido y demostrado en las lecciones 2 a 5. Hoy simplemente *aplicamos* esas herramientas a un sistema de dos ecuaciones. Es, en cierto sentido, la culminación aritmética de todo lo anterior: la recompensa de haber invertido tiempo en entender qué es realmente un producto escalar.

También conviene recordar la motivación última de todo esto, sembrada ya en la lección 3: MCO no es más que la búsqueda del vector $\hat{\mathbf{y}}$ del subespacio $\mathcal{L}(\mathbf{1}, \mathbf{x})$ que minimiza la distancia $\|\mathbf{y} - \hat{\mathbf{y}}\|$. Esa minimización —un problema de mínimos cuadrados— es equivalente a imponer que el vector de error $\hat{\mathbf{e}} = \mathbf{y} - \hat{\mathbf{y}}$ sea ortogonal al subespacio. Hoy vemos, por fin, adónde lleva esa condición cuando se traduce en números.

Para profundizar:

- Wooldridge, J. M. (2020). *Introductory Econometrics*, cap. 2, sección 2.2: “Deriving the Ordinary Least Squares Estimates” (la misma derivación algebraica, en notación de sumatorios).

2 Primera condición: el residuo tiene media cero

$$\text{Recordemos: } \hat{\mathbf{e}} = \mathbf{y} - \hat{\mathbf{y}} = \mathbf{y} - \hat{\beta}_1 \mathbf{1} - \hat{\beta}_2 \mathbf{x}.$$

De la primera condición, $\langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = 0$:

$$\langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = \left\langle (\mathbf{y} - \hat{\beta}_1 \mathbf{1} - \hat{\beta}_2 \mathbf{x}), \mathbf{1} \right\rangle_s = \mu_{\mathbf{y}} - \hat{\beta}_1 - \hat{\beta}_2 \mu_{\mathbf{x}} = 0.$$

Despejando

$$\boxed{\hat{\beta}_1 = \mu_{\mathbf{y}} - \hat{\beta}_2 \mu_{\mathbf{x}}.}$$

El intercepto (la constante) es *función de* la pendiente (aún falta una ecuación más).

Primera condición: el residuo tiene media cero

La primera condición de ortogonalidad, $\langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = 0$ ya la hemos visto en otro contexto: es la misma condición que, en la lección 4, hacía que el vector en desviaciones tuviera suma cero. Aquí el residuo es más complejo —depende de dos coeficientes, no de uno—, pero la condición algebraica es idéntica: la media de los residuos es cero, por tanto, *la suma de los residuos es cero*.

Recordando que: $\hat{\mathbf{e}} = \mathbf{y} - (\hat{\beta}_1 \mathbf{1} + \hat{\beta}_2 \mathbf{x})$, la primera condición queda en:

$$\langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = \left\langle (\mathbf{y} - \hat{\beta}_1 \mathbf{1} - \hat{\beta}_2 \mathbf{x}), \mathbf{1} \right\rangle_s = 0;$$

y aplicando la propiedad distributiva:

$$\langle \mathbf{y}, \mathbf{1} \rangle_s - \hat{\beta}_1 \langle \mathbf{1}, \mathbf{1} \rangle_s - \hat{\beta}_2 \langle \mathbf{x}, \mathbf{1} \rangle_s = 0.$$

Recordando la definición de media de un vector: $\langle \mathbf{v}, \mathbf{1} \rangle_s = \mu_{\mathbf{v}}$, y que $\langle \mathbf{1}, \mathbf{1} \rangle_s = \|\mathbf{1}\|_s^2 = 1$, llegamos a

$$\mu_{\mathbf{y}} - \hat{\beta}_1 - \hat{\beta}_2 \mu_{\mathbf{x}} = 0,$$

de donde despejamos inmediatamente

$$\hat{\beta}_1 = \mu_{\mathbf{y}} - \hat{\beta}_2 \mu_{\mathbf{x}}.$$

Este resultado, todavía incompleto (pues depende del valor —aún desconocido— de $\hat{\beta}_2$), revelará el papel que juega $\hat{\beta}_1$ en el ajuste de la recta de regresión: sea cual sea la pendiente $\hat{\beta}_2$, la constante $\hat{\beta}_1$ siempre tomará el valor necesario para que la recta ajustada pase por el punto $(\mu_{\mathbf{x}}, \mu_{\mathbf{y}})$. Volveremos sobre esta idea, con su lectura geométrica completa, más adelante en la sesión.

Hay otra lectura de este mismo resultado, igual de importante. Recordemos que en la lección 6 vimos que $\mu_{\hat{\mathbf{e}}} = \langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = 0$ (el vector de residuos tiene media cero). Esto permite deducir que la ortogonalidad del residuo frente a $\mathbf{1}$ implica que el ajuste *preserva la media* de $\mathbf{y}$. En efecto: como $\mathbf{y} = \hat{\mathbf{y}} + \hat{\mathbf{e}}$, tomando medias en ambos lados (la media es lineal, pues es un producto escalar) obtenemos

$$\mu_{\mathbf{y}} = \mu_{\hat{\mathbf{y}}} + \mu_{\hat{\mathbf{e}}} = \mu_{\hat{\mathbf{y}}}.$$

Es decir: *la media del vector ajustado coincide con la media del vector original*. Este hecho, que puede parecer una curiosidad menor, será la clave geométrica de la transparencia dedicada al lugar de las medias, más adelante en esta misma sesión.

Para profundizar:

- Bujosa, M. *Curso de Álgebra Lineal*, cap. 11: linealidad del producto escalar, aplicada aquí a la media de una suma de vectores.

3 Segunda condición: el residuo ortogonal a $\mathbf{x}$

De la segunda condición, $\langle \hat{\mathbf{e}}, \mathbf{x} \rangle_s = 0$:

$$\langle \hat{\mathbf{e}}, \mathbf{x} \rangle_s = \langle (\mathbf{y} - \hat{\beta}_1 \mathbf{1} - \hat{\beta}_2 \mathbf{x}), \mathbf{x} \rangle_s = \langle \mathbf{y}, \mathbf{x} \rangle_s - \hat{\beta}_1 \mu_{\mathbf{x}} - \hat{\beta}_2 \mu_{\mathbf{x}^2} = 0.$$

Sustituyendo $\hat{\beta}_1 = \mu_{\mathbf{y}} - \hat{\beta}_2 \mu_{\mathbf{x}}$ (transparencia anterior) y agrupando:

$$\langle \mathbf{y}, \mathbf{x} \rangle_s - \mu_{\mathbf{x}} \mu_{\mathbf{y}} = \hat{\beta}_2 (\mu_{\mathbf{x}^2} - \mu_{\mathbf{x}}^2).$$

Como $\langle \mathbf{x}, \mathbf{y} \rangle_s = \mu_{\mathbf{x} \odot \mathbf{y}}$ (media del producto componente a componente):

$$\underbrace{\mu_{\mathbf{x} \odot \mathbf{y}} - \mu_{\mathbf{x}} \mu_{\mathbf{y}}}_{\sigma_{\mathbf{xy}}} = \hat{\beta}_2 \underbrace{(\mu_{\mathbf{x}^2} - \mu_{\mathbf{x}}^2)}_{\sigma_{\mathbf{x}}^2};$$

donde aparecen las “fórmulas de cálculo” de covarianza y varianza. Así:

$$\boxed{\sigma_{\mathbf{xy}} = \hat{\beta}_2 \sigma_{\mathbf{x}}^2.}$$

Segunda condición: el residuo ortogonal a $\mathbf{x}$

La segunda condición de ortogonalidad, $\langle \hat{\mathbf{e}}, \mathbf{x} \rangle_s = 0$, se desarrolla con la misma herramienta que la primera —la propiedad distributiva (por la bilinealidad) del producto escalar—, aplicada ahora sobre $\mathbf{x}$ en lugar de sobre $\mathbf{1}$:

$$\langle \hat{\mathbf{e}}, \mathbf{x} \rangle_s = \langle (\mathbf{y} - \hat{\beta}_1 \mathbf{1} - \hat{\beta}_2 \mathbf{x}), \mathbf{x} \rangle_s = \langle \mathbf{y}, \mathbf{x} \rangle_s - \hat{\beta}_1 \langle \mathbf{1}, \mathbf{x} \rangle_s - \hat{\beta}_2 \langle \mathbf{x}, \mathbf{x} \rangle_s = 0.$$

Necesitamos identificar los dos productos escalares que involucran a $\mathbf{x}$ consigo mismo o con $\mathbf{1}$. El primero ya lo conocemos: $\langle \mathbf{1}, \mathbf{x} \rangle_s = \mu_{\mathbf{x}}$ (la media, lección 4). El segundo lo vimos en la lección 5: $\langle \mathbf{x}, \mathbf{x} \rangle_s = \frac{1}{n} \sum_{i=1}^n x_i x_i = \mu_{\mathbf{x}^2}$, donde $\mathbf{x}^2$ denotaba el vector cuyas componentes son los cuadrados de las de $\mathbf{x}$ —aunque allí no lo dijimos con este nombre, $\mathbf{x}^2$ es $\mathbf{x} \odot \mathbf{x}$, el *producto componente a componente* de $\mathbf{x}$ consigo mismo—. Del mismo modo tenemos que $\langle \mathbf{x}, \mathbf{y} \rangle_s = \frac{1}{n} \sum_{i=1}^n x_i y_i = \mu_{\mathbf{x} \odot \mathbf{y}}$, donde $\mathbf{x} \odot \mathbf{y}$ denota el vector cuyas componentes son los productos componente a componente de $\mathbf{x}$ e $\mathbf{y}$. Sustituyendo productos escalares tenemos:

$$\mu_{\mathbf{x} \odot \mathbf{y}} - \hat{\beta}_1 \mu_{\mathbf{x}} - \hat{\beta}_2 \mu_{\mathbf{x}^2} = 0.$$

Sustituimos ahora $\hat{\beta}_1 = \mu_{\mathbf{y}} - \hat{\beta}_2 \mu_{\mathbf{x}}$ (transparencia anterior) y agrupamos los términos con $\hat{\beta}_2$:

$$\mu_{\mathbf{x} \odot \mathbf{y}} - (\mu_{\mathbf{y}} - \hat{\beta}_2 \mu_{\mathbf{x}}) \mu_{\mathbf{x}} - \hat{\beta}_2 \mu_{\mathbf{x}^2} = 0 \implies \mu_{\mathbf{x} \odot \mathbf{y}} - \mu_{\mathbf{x}} \mu_{\mathbf{y}} = \hat{\beta}_2 (\mu_{\mathbf{x}^2} - \mu_{\mathbf{x}}^2).$$

El paréntesis de la derecha ya lo reconocemos: es exactamente $\sigma_{\mathbf{x}}^2 = \mu_{\mathbf{x}^2} - \mu_{\mathbf{x}}^2$, la “fórmula de cálculo” de la varianza que obtuvimos en la lección 5 a partir de Pitágoras. Así que

$$\mu_{\mathbf{x} \odot \mathbf{y}} - \mu_{\mathbf{x}} \mu_{\mathbf{y}} = \hat{\beta}_2 \sigma_{\mathbf{x}}^2.$$

Falta reconocer el lado izquierdo; y resulta ser la “fórmula de cálculo” de la *covarianza* —el análogo exacto de la fórmula de la varianza, ahora con dos vectores distintos en lugar de uno consigo mismo—. En efecto, por la misma bilinealidad que hemos usado en toda esta lección (y que ya usamos en la lección 2):

$$\sigma_{\mathbf{x}\mathbf{y}} = \langle \mathbf{x} - \mu_{\mathbf{x}} \mathbf{1}, \mathbf{y} - \mu_{\mathbf{y}} \mathbf{1} \rangle_s = \langle \mathbf{x}, \mathbf{y} \rangle_s - \mu_{\mathbf{x}} \langle \mathbf{1}, \mathbf{y} \rangle_s - \mu_{\mathbf{y}} \langle \mathbf{x}, \mathbf{1} \rangle_s + \mu_{\mathbf{x}} \mu_{\mathbf{y}} \langle \mathbf{1}, \mathbf{1} \rangle_s = \mu_{\mathbf{x} \odot \mathbf{y}} - \mu_{\mathbf{x}} \mu_{\mathbf{y}}.$$

Así que nuestra ecuación se reescribe, sin ningún ingrediente nuevo, como

$$\sigma_{\mathbf{x}\mathbf{y}} = \hat{\beta}_2 \sigma_{\mathbf{x}}^2,$$

que es ya, prácticamente, el resultado final. Solo falta despejar $\hat{\beta}_2$, lo que hacemos en la siguiente transparencia.

Nota de notación. El subíndice de $\sigma_{\mathbf{x}\mathbf{y}}$ (covarianza *entre* $\mathbf{x}$ e $\mathbf{y}$, vectores ya centrados) y el de $\mu_{\mathbf{x} \odot \mathbf{y}}$ (media *del vector* $\mathbf{x} \odot \mathbf{y}$, sin centrar) se parecen tipográficamente, pero indican cosas distintas: el primero nombra una relación entre dos vectores; el segundo nombra la media de un vector nuevo, construido multiplicando componente a componente. El símbolo $\odot$ hace explícita esa construcción para evitar la confusión entre ambos subíndices.

Para profundizar:

- Bujosa, M. *Curso de Álgebra Lineal*, [cap. 11](#): bilinealidad del producto escalar.

4 $\hat{\beta}_2$: covarianza sobre varianza

Despejando de $\sigma_{xy} = \hat{\beta}_2 \sigma_x^2$:

$$\hat{\beta}_2 = \frac{\sigma_{xy}}{\sigma_x^2}.$$

Y recordando $\rho_{xy} = \frac{\sigma_{xy}}{\sigma_x \sigma_y}$ (lección 5):

$$\hat{\beta}_2 = \rho_{xy} \frac{\sigma_y}{\sigma_x}.$$

La pendiente MCO es *la correlación, reescalada por el ratio de las desviaciones típicas*: su signo es el de ρ_{xy} .

$\hat{\beta}_2$: covarianza sobre varianza

De la ecuación $\sigma_{xy} = \hat{\beta}_2 \sigma_x^2$ obtenida en la transparencia anterior se despeja de inmediato

$$\hat{\beta}_2 = \frac{\sigma_{xy}}{\sigma_x^2}.$$

Esta es la fórmula que probablemente recuerde de un curso de estadística: la pendiente de la recta de mínimos cuadrados es la covarianza entre $\mathbf{x}$ e $\mathbf{y}$ dividida por la varianza de $\mathbf{x}$. Lo que hemos hecho aquí es mostrar *de dónde sale*: no es una regla mnemotécnica aislada, sino la consecuencia directa —vía Pitágoras y bilinealidad del producto escalar— de imponer la segunda condición de ortogonalidad.

Podemos reescribir esta fórmula de una manera todavía más reveladora. Recordando la definición de correlación de la lección 5, $\rho_{xy} = \sigma_{xy}/(\sigma_x \sigma_y)$, despejamos $\sigma_{xy} = \rho_{xy} \sigma_x \sigma_y$ y sustituimos:

$$\hat{\beta}_2 = \frac{\rho_{xy} \sigma_x \sigma_y}{\sigma_x^2} = \rho_{xy} \frac{\sigma_y}{\sigma_x}.$$

Esta identidad, ya anticipada como promesa en las lecciones 5 y 6 (“pendiente = correlación reescalada”), tiene una lectura muy intuitiva: la pendiente MCO es el coseno del ángulo θ entre los vectores en desviaciones de $\mathbf{x}$ e $\mathbf{y}$ (la correlación), corregido por el cociente de las escalas de ambas variables (σ_y/σ_x). Si $\mathbf{x}$ e $\mathbf{y}$ tuvieran la misma dispersión ($\sigma_x = \sigma_y$), la pendiente coincidiría exactamente con la correlación.

Una consecuencia inmediata que conviene subrayar: como σ_x^2 y σ_y son siempre no negativas (son normas, o normas al cuadrado), el *signo* de $\hat{\beta}_2$ coincide siempre con el signo de σ_{xy} , que a su vez coincide siempre con el signo de ρ_{xy} . Dicho de otro modo: *la pendiente de la recta de regresión y la correlación entre las dos variables tienen siempre el mismo signo*. Esta observación, trivial una vez vista la fórmula, es la base de la intuición habitual de que “pendiente positiva” y “correlación positiva” son, en cierto sentido, la misma cosa vista desde ángulos distintos.

Para profundizar:

- Bujosa, M. *Curso de Álgebra Lineal*, [libro online](#), cap. 11.
- Wooldridge, J. M. (2020), cap. 2, ecuación (2.19) y comentario adyacente sobre el signo de $\hat{\beta}_2$.

5 Vía alternativa: generadores ortogonales

Recordemos (lección 6): $\mathcal{L}(\mathbf{1}, \mathbf{x}) = \mathcal{L}(\mathbf{1}, \mathbf{x} - \mu_x \mathbf{1})$.

Y ahora sí: $\mathbf{1} \perp (\mathbf{x} - \mu_x \mathbf{1})$, porque $\langle (\mathbf{x} - \mu_x \mathbf{1}), \mathbf{1} \rangle_s = \mu_x - \mu_x = 0$.

Con generadores *ortogonales*, la proyección se descompone en dos componentes ortogonales, uno por generador; cada coeficiente es la razón que resulta de proyectar sobre esa dirección *por separado* (Cauchy–Schwarz, lección 3: $\alpha = \frac{\langle \mathbf{a}, \mathbf{b} \rangle}{\langle \mathbf{a}, \mathbf{a} \rangle} = \frac{\langle \mathbf{a}, \mathbf{b} \rangle}{\|\mathbf{a}\|^2}$):

$$c_1 = \frac{\langle \mathbf{y}, \mathbf{1} \rangle_s}{\|\mathbf{1}\|_s^2} = \mu_y, \quad c_2 = \frac{\langle \mathbf{y}, (\mathbf{x} - \mu_x \mathbf{1}) \rangle_s}{\|\mathbf{x} - \mu_x \mathbf{1}\|_s^2} = \frac{\sigma_{xy}}{\sigma_x^2} = \hat{\beta}_2.$$

Sin resolver ningún sistema: dos proyecciones independientes, sobre dos direcciones perpendiculares.

Vía alternativa: generadores ortogonales

Existe un segundo camino hacia el mismo resultado, más elegante que resolver un sistema de ecuaciones, y que conviene conocer porque reaparecerá (en su versión con más de un regresor) en la lección 10.

La observación de partida ya se hizo en la lección 6: el plano $\mathcal{L}(\mathbf{1}, \mathbf{x})$ —el conjunto de todas las combinaciones lineales de $\mathbf{1}$ y $\mathbf{x}$ — puede generarse igualmente con $\mathbf{1}$ y el vector en desviaciones $\mathbf{x} - \mu_x \mathbf{1}$:

$$\mathcal{L}(\mathbf{1}, \mathbf{x}) = \mathcal{L}(\mathbf{1}, \mathbf{x} - \mu_x \mathbf{1}).$$

(Esto es así porque $\mathbf{x} - \mu_x \mathbf{1}$ es una combinación lineal de $\mathbf{1}$ y $\mathbf{x}$, y viceversa: $\mathbf{x} = (\mathbf{x} - \mu_x \mathbf{1}) + \mu_x \mathbf{1}$; cambiar de generadores de este modo no cambia el conjunto de combinaciones lineales que se pueden formar, es decir, no cambia el plano.)

La razón por la que este cambio de generadores es interesante —y no un mero capricho notacional— es que, a diferencia de $\mathbf{1}$ y $\mathbf{x}$ (que en general *no* son ortogonales, salvo que $\mu_x = 0$, pues $\langle \mathbf{1}, \mathbf{x} \rangle_s = \mu_x$), los nuevos generadores $\mathbf{1}$ y $\mathbf{x} - \mu_x \mathbf{1}$ **SÍ** son ortogonales:

$$\langle (\mathbf{x} - \mu_x \mathbf{1}), \mathbf{1} \rangle_s = \langle \mathbf{x}, \mathbf{1} \rangle_s - \mu_x \langle \mathbf{1}, \mathbf{1} \rangle_s = \mu_x - \mu_x \cdot 1 = 0.$$

¿Por qué importa la ortogonalidad de los generadores? Porque cuando un plano se genera con dos vectores ortogonales, la proyección de cualquier vector sobre ese plano se descompone en *dos proyecciones independientes*, una sobre cada generador —exactamente la fórmula de proyección de la lección 3, $\alpha = \frac{\langle \mathbf{b}, \mathbf{a} \rangle}{\langle \mathbf{a}, \mathbf{a} \rangle} = \frac{\langle \mathbf{b}, \mathbf{a} \rangle}{\|\mathbf{a}\|^2}$ (la misma que sirvió para demostrar Cauchy–Schwarz), aplicada dos veces, una por generador. No hace falta resolver ningún sistema de ecuaciones simultáneas: cada coeficiente se calcula de forma completamente independiente del otro.

Aplicándolo aquí: sea $\hat{\mathbf{y}} = c_1 \mathbf{1} + c_2 (\mathbf{x} - \mu_x \mathbf{1})$ la proyección de $\mathbf{y}$ sobre el plano, escrita en esta nueva base. Por ortogonalidad de los generadores,

$$c_1 = \frac{\langle \mathbf{y}, \mathbf{1} \rangle_s}{\|\mathbf{1}\|_s^2} = \frac{\mu_y}{1} = \mu_y, \quad c_2 = \frac{\langle \mathbf{y}, (\mathbf{x} - \mu_x \mathbf{1}) \rangle_s}{\|\mathbf{x} - \mu_x \mathbf{1}\|_s^2} = \frac{\sigma_{xy}}{\sigma_x^2} = \hat{\beta}_2.$$

El mismo resultado que obtuvimos resolviendo el sistema, pero sin resolver ningún sistema.

Falta un último paso: traducir c_1, c_2 (los coeficientes en la base “centrada”) a $\hat{\beta}_1, \hat{\beta}_2$ (los coeficientes en la base original $\mathbf{1}, \mathbf{x}$). Basta sustituir c_1 y c_2 en $\hat{\mathbf{y}} = c_1 \mathbf{1} + c_2(\mathbf{x} - \mu_x \mathbf{1})$ y reagrupar:

$$\hat{\mathbf{y}} = \mu_y \mathbf{1} + \hat{\beta}_2(\mathbf{x} - \mu_x \mathbf{1}) = \hat{\beta}_2 \mathbf{x} + \underbrace{(\mu_y - \hat{\beta}_2 \mu_x)}_{\hat{\beta}_1} \mathbf{1}.$$

Reconocemos, en el término entre paréntesis, exactamente la fórmula de $\hat{\beta}_1$ que ya habíamos obtenido de la primera condición de ortogonalidad. Los dos caminos —resolver el sistema directamente, o cambiar a generadores ortogonales— llegan, como no podía ser de otro modo, al mismo resultado. El segundo camino es bastante ilustrativo: muestra con claridad que la dificultad de resolver “un sistema de dos ecuaciones” era, en el fondo, solo un problema de *elección de base*; con la base adecuada (generadores ortogonales), el sistema se desacopla por completo.

Para profundizar:

- Bujosa, M. *Curso de Álgebra Lineal*, cap. 11: proyección sobre un vector, fórmula $\alpha = \langle \mathbf{a}, \mathbf{b} \rangle / \langle \mathbf{a}, \mathbf{a} \rangle$.
- Strang, G. (2016). *Introduction to Linear Algebra*, sección sobre bases ortogonales y proyecciones independientes.

6 $\hat{\beta}_1$ y el lugar de las medias

De $\hat{\beta}_1 = \mu_y - \hat{\beta}_2 \mu_x$, reordenando:

$$\mu_y = \hat{\beta}_1 + \hat{\beta}_2 \mu_x.$$

Así, evaluando la recta ajustada $\hat{y} = \hat{\beta}_1 + \hat{\beta}_2 x$ en el punto $x = \mu_x$, el resultado es μ_y

La recta ajustada pasa exactamente por el punto de las medias (μ_x, μ_y) .

Además (primera condición): $\mu_{\hat{\mathbf{y}}} = \mu_y$, el ajuste preserva la media.

Tres cantidades que coinciden: $\mu_y = \mu_{\hat{\mathbf{y}}} = \hat{\beta}_1 + \hat{\beta}_2 \mu_x$ ([versión interactiva](#)).

Y hay más, los vectores en desviaciones de $\hat{\mathbf{y}}$ y $\hat{\beta}_2 \mathbf{x}$ son iguales:

$$\hat{\mathbf{y}} - \mu_{\hat{\mathbf{y}}} \mathbf{1} = (\hat{\beta}_1 \mathbf{1} + \hat{\beta}_2 \mathbf{x}) - (\hat{\beta}_1 + \hat{\beta}_2 \mu_x) \mathbf{1} = \hat{\beta}_2 \mathbf{x} - \hat{\beta}_2 \mu_x \mathbf{1} = \hat{\beta}_2 (\mathbf{x} - \mu_x \mathbf{1}).$$

De ahí que $\sigma_{\hat{\mathbf{y}}} = \sigma_{\hat{\beta}_2 \mathbf{x}}$. Y también que $\mathbf{y}$ forme el mismo ángulo con $\hat{\mathbf{y}}$ que con $\hat{\beta}_2 \mathbf{x}$.

$\hat{\beta}_1$ y el lugar de las medias

La fórmula $\hat{\beta}_1 = \mu_y - \hat{\beta}_2 \mu_x$, obtenida ya en la transparencia “Primera condición: el residuo tiene media cero”, admite una lectura geométrica muy bonita si la reordenamos como

$$\mu_y = \hat{\beta}_1 + \hat{\beta}_2 \mu_x.$$

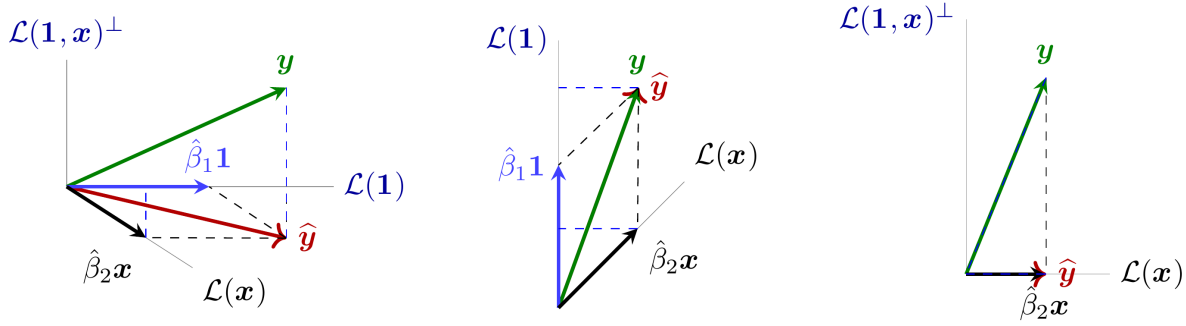


Figura 1: La proyección de la lección 6, vista desde tres ángulos, con las medias señaladas. (Izquierda) De frente al plano $(\mathcal{L}(\mathbf{1}), \mathcal{L}(\mathbf{1}, \mathbf{x})^\perp)$: las proyecciones de $\mathbf{y}$ y de $\hat{\mathbf{y}}$ sobre $\mathcal{L}(\mathbf{1})$ caen en el mismo punto. (Centro) Vista cenital: sumar la proyección de $\hat{\beta}_2\mathbf{x}$ y la de $\hat{\beta}_1\mathbf{1}$ sobre $\mathcal{L}(\mathbf{1})$ reconstruye ese mismo punto. (Derecha) Vista alineada con $\mathcal{L}(\mathbf{1})$: lo que se ve son, literalmente, los vectores en desviaciones; el ángulo entre el de $\mathbf{y}$ y el de $\hat{\mathbf{y}}$ será protagonista de la lección 8.

Esta expresión dice, literalmente, que si evaluamos la recta ajustada $\hat{y} = \hat{\beta}_1 + \hat{\beta}_2x$ en el punto $x = \mu_x$, el resultado es exactamente μ_y . Dicho de otro modo: *la recta de regresión MCO pasa siempre, sin excepción, por el punto de las medias (μ_x, μ_y)* . Este resultado, que en un curso de estadística aparece a menudo como una propiedad más entre otras, es aquí una consecuencia directa e inevitable de la primera condición de ortogonalidad: no puede dejar de cumplirse.

Esta propiedad se conecta directamente con la observación que hicimos en la transparencia “Primera condición: el residuo tiene media cero”: como $\hat{\mathbf{e}} \perp \mathbf{1}$, se cumple $\mu_{\hat{\mathbf{e}}} = 0$, y por tanto $\mu_y = \mu_{\hat{y}}$: la media del vector ajustado coincide con la media del vector original. Podemos ahora reunir ambas piezas en una única cadena de igualdades:

$$\mu_y = \mu_{\hat{y}} = \hat{\beta}_1 + \hat{\beta}_2\mu_x.$$

La primera igualdad es un hecho sobre *vectores* (dos vectores distintos, $\mathbf{y}$ y $\hat{\mathbf{y}}$, comparten la misma media). La segunda es un hecho sobre la *recta ajustada*, evaluada en un punto concreto. Ambas son, en realidad, la misma consecuencia geométrica vista desde dos ángulos.

La figura recién mostrada más arriba hace visible, con tres perspectivas complementarias, toda esta cadena de igualdades. La vista de la izquierda retoma el plano frontal $(\mathcal{L}(\mathbf{1}), \mathcal{L}(\mathbf{1}, \mathbf{x})^\perp)$ ya usado en la lección 6 (la vista central de aquella figura): en ella se aprecia directamente que $\mathbf{y}$ y $\hat{\mathbf{y}}$ tienen la misma proyección sobre $\mathcal{L}(\mathbf{1})$, es decir, $\mu_y = \mu_{\hat{y}}$. La vista central es una vista *cenital*, perpendicular al plano de los regresores: en ella se ve que sumar la proyección de $\hat{\beta}_2\mathbf{x}$ (sobre $\mathcal{L}(\mathbf{1})$) a $\hat{\beta}_1\mathbf{1}$ reconstruye exactamente el punto $\mu_y\mathbf{1}$; y que las longitudes de las componentes perpendiculares a $\mathcal{L}(\mathbf{1})$ de $\hat{\mathbf{y}}$ y de $\hat{\beta}_2\mathbf{x}$ son iguales (es decir, que $\hat{\mathbf{y}}$ y $\hat{\beta}_2\mathbf{x}$ tienen la misma desviación típica). La vista de la derecha, nueva, está tomada *alineada con $\mathcal{L}(\mathbf{1})$* . Al mirar en la dirección de $\mathbf{1}$ perdemos toda profundidad en esa dirección, y lo que vemos —de frente, sin distorsión— es el propio plano $\mathcal{L}(\mathbf{1})^\perp$. Es decir: vemos los vectores en desviaciones de $\mathbf{y}$, de $\hat{\mathbf{y}}$ y de $\hat{\beta}_2\mathbf{x}$, con su verdadera longitud y su verdadero ángulo entre sí.

Esta última vista no es un simple adorno: adelanta, sin demostrarlo aún, al protagonista de la lección 8. El ángulo que en ella se aprecia entre $\mathbf{y}$ y $\hat{\mathbf{y}}$ —medido, como siempre, entre sus vectores en desviaciones— será exactamente el ángulo cuyo coseno al cuadrado define el R^2 .

La igualdad que anuncian las vistas central y derecha —que las componentes perpendiculares de $\hat{\mathbf{y}}$ y de $\hat{\beta}_2\mathbf{x}$ coinciden— no es una casualidad geométrica ni requiere ningún argumento nuevo: es consecuencia inmediata de una propiedad ya demostrada en la lección 4. Como $\hat{\mathbf{y}} = \hat{\beta}_1\mathbf{1} + \hat{\beta}_2\mathbf{x}$ y $\hat{\beta}_1\mathbf{1}$ es un vector constante, sumarlo no cambia el vector en desviaciones (lección 4: “sumar una constante no cambia la desviación típica”; aquí usamos, más precisamente, que ni siquiera cambia el propio vector en desviaciones, no solo su norma). Por tanto:

$$\hat{\mathbf{y}} - \mu_{\hat{\mathbf{y}}}\mathbf{1} = (\hat{\beta}_1\mathbf{1} + \hat{\beta}_2\mathbf{x}) - (\hat{\beta}_1 + \hat{\beta}_2\mu_x)\mathbf{1} = \hat{\beta}_2\mathbf{x} - \hat{\beta}_2\mu_x\mathbf{1} = \hat{\beta}_2(\mathbf{x} - \mu_x\mathbf{1}).$$

El vector en desviaciones de $\hat{\mathbf{y}}$ es, *literalmente*, el mismo vector que el vector en desviaciones de $\hat{\beta}_2\mathbf{x}$ (no solo un vector de la misma longitud: es el mismísimo vector). De ahí se siguen, sin ningún paso adicional, dos consecuencias:

- Tomando normas: $\sigma_{\hat{\mathbf{y}}} = \sigma_{\hat{\beta}_2\mathbf{x}}$ (dos vectores idénticos tienen, evidentemente, la misma norma).
- Como la correlación de $\mathbf{y}$ con cualquier vector depende únicamente del vector en desviaciones de ese vector (lección 5), y $\hat{\mathbf{y}}$ y $\hat{\beta}_2\mathbf{x}$ comparten el mismo vector en desviaciones, $\mathbf{y}$ forma el mismo ángulo con ambos.

Guardamos esta identidad — $\hat{\mathbf{y}} - \mu_{\hat{\mathbf{y}}}\mathbf{1} = \hat{\beta}_2(\mathbf{x} - \mu_x\mathbf{1})$ — porque reaparecerá en la lección 8: es la pieza que permite traducir directamente entre la descomposición de varianzas $\text{STC} = \text{SEC} + \text{SRC}$ y la relación $R^2 = \rho_{xy}^2$ de la regresión simple.

Conviene subrayar que esta propiedad —la recta pasa por el punto de las medias— es exclusiva del criterio de mínimos cuadrados *con constante incluida*. Si forzáramos, por alguna razón, un ajuste sin intercepto (es decir, proyectando solo sobre $\mathcal{L}(\mathbf{x})$, sin $\mathbf{1}$), esta propiedad desaparecería, porque ya no tendríamos la primera condición de ortogonalidad ($\langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = 0$) que la garantiza. En este curso, todos los modelos de regresión incluyen constante, así que esta propiedad se cumple siempre. Es más: en opinión del autor de estos apuntes, **sin constante no puede hablarse de regresión**, porque nos salimos del marco de la estadística; no incluirla es un error conceptual. El laboratorio que sigue a la lección 8 muestra qué se pierde al omitirla.

Para profundizar:

- Wooldridge, J. M. (2020), cap. 2, comentario tras la ecuación (2.18) sobre el punto de las medias.

7 Ejemplo numérico completo

Como en la [actividad 3 \(html\)](#) del laboratorio ($\mathbf{y} = 2 + 3\mathbf{x} + \mathbf{u}$), pero con $n = 5$, calculable a mano:

x_i	u_i	$y_i = 2 + 3x_i + u_i$	$\hat{y}_i = \hat{\beta}_1 + \hat{\beta}_2 x_i$
1	1	6	5
2	-2	6	8
3	0	11	11
4	2	16	14
5	-1	16	17

$$\mu_x = 3, \quad \mu_y = 11, \quad \sigma_x^2 = 2, \quad \sigma_{xy} = 6.$$

$$\hat{\beta}_2 = \frac{6}{2} = 3, \quad \hat{\beta}_1 = 11 - 3 \cdot 3 = 2.$$

¡Recuperamos exactamente $\beta_1 = 2$, $\beta_2 = 3$! (No es lo habitual: $\mathbf{u}$ se eligió ortogonal a los regresores.)

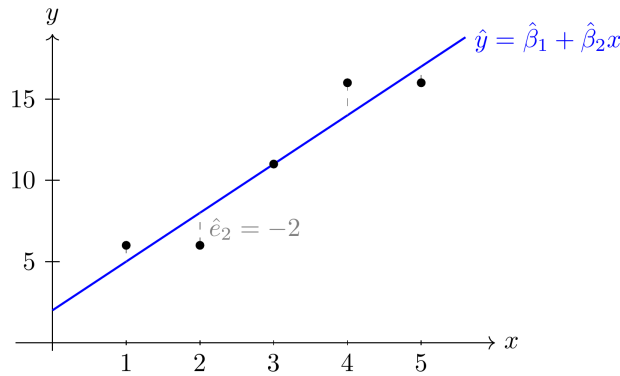


Figura 2: Diagrama de dispersión clásico del ejemplo: los cinco puntos (x_i, y_i) y la recta ajustada $\hat{y} = 2 + 3x$. Las líneas discontinuas son los residuos $\hat{e}_i = u_i$.

Versión interactiva: [cuaderno 2 en MyBinder](#), parte 3 (qué le exigimos a $\mathbf{u}$ para recuperar los coeficientes).

Ejemplo numérico completo

Verifiquemos con números todo lo dicho hasta ahora. Construimos un ejemplo pequeño, calculable enteramente a mano, en el mismo espíritu que la actividad de la práctica de laboratorio ([actividad 3 \(html\)](#)), donde se generaban datos a partir de una relación lineal conocida $\mathbf{y} = 2 + 3\mathbf{x} + \mathbf{u}$ para comprobar que MCO recupera aproximadamente esos coeficientes generadores. Aquí, con solo $n = 5$ observaciones (frente a la muestra mayor y aleatoria de la práctica), podemos seguir el cálculo paso a paso.

Tomamos $\mathbf{x} = (1, 2, 3, 4, 5)$ y un vector $\mathbf{u} = (1, -2, 0, 2, -1)$ (nótese que $\sum u_i = 0$ y que, como veremos, $\mathbf{u}$ se eligió deliberadamente también ortogonal a $\mathbf{x}$). Generamos $y_i = 2 + 3x_i + u_i$:

$$\mathbf{y} = (6, 6, 11, 16, 16).$$

Medias: $\mu_x = \frac{1+2+3+4+5}{5} = 3$, $\mu_y = \frac{6+6+11+16+16}{5} = \frac{55}{5} = 11$.

Vectores en desviaciones: $\mathbf{x} - \bar{\mathbf{x}} = (-2, -1, 0, 1, 2)$, $\mathbf{y} - \bar{\mathbf{y}} = (-5, -5, 0, 5, 5)$.

Varianza de $\mathbf{x}$:

$$\sigma_x^2 = \frac{1}{5}[(-2)^2 + (-1)^2 + 0^2 + 1^2 + 2^2] = \frac{1}{5}(4 + 1 + 0 + 1 + 4) = \frac{10}{5} = 2.$$

Covarianza:

$$\sigma_{xy} = \frac{1}{5}[(-2)(-5) + (-1)(-5) + (0)(0) + (1)(5) + (2)(5)] = \frac{1}{5}(10 + 5 + 0 + 5 + 10) = \frac{30}{5} = 6.$$

Pendiente y ordenada en el origen:

$$\hat{\beta}_2 = \frac{\sigma_{xy}}{\sigma_x^2} = \frac{6}{2} = 3, \quad \hat{\beta}_1 = \mu_y - \hat{\beta}_2 \mu_x = 11 - 3 \cdot 3 = 2.$$

Recuperamos *exactamente* los coeficientes generadores $\beta_1 = 2$, $\beta_2 = 3$. Esto no es casualidad, pero tampoco es lo que ocurre en general: sucede porque, en este ejemplo pequeño, elegimos deliberadamente $\mathbf{u}$ de manera que su covarianza con $\mathbf{x}$ fuera cero (compruébese: $\frac{1}{5} \sum (x_i - \mu_x) u_i = \frac{1}{5}[(-2)(1) + (-1)(-2) + (0)(0) + (1)(2) + (2)(-1)] = \frac{1}{5}(-2 + 2 + 0 + 2 - 2) = 0$) y su media cero. Con datos reales, o incluso con datos simulados de la forma habitual (perturbaciones generadas al azar, sin imponer esta cancelación exacta), $\hat{\beta}_2$ y $\hat{\beta}_1$ se *aproximarán* a los valores generadores, pero no los reproducirán con exactitud —como puede comprobarse en la práctica de laboratorio, con su muestra de mayor tamaño y sus perturbaciones genuinamente aleatorias—.

Podemos además comprobar, con este mismo ejemplo, que los residuos coinciden con $\mathbf{u}$: como el ajuste recupera exactamente los coeficientes generadores, $\hat{e}_i = y_i - (\hat{\beta}_1 + \hat{\beta}_2 x_i) = y_i - (2 + 3x_i) = u_i$. Y, como debía ocurrir por construcción, $\sum \hat{e}_i = \sum u_i = 0$ y $\sum x_i \hat{e}_i = \sum x_i u_i = 0$ (la ortogonalidad con los regresores que impusimos al elegir $\mathbf{u}$).

Para profundizar:

- Wooldridge, J. M. (2020), cap. 2, ejemplo numérico análogo con datos de salario/educación.

8 Recapitulación y guiño a la sesión siguiente

Condición de ortogonalidad	Consecuencia
$\langle \hat{\mathbf{e}}, \mathbf{1} \rangle_s = 0$	$\hat{\beta}_1 = \mu_y - \hat{\beta}_2 \mu_x$; $\mu_{\hat{\mathbf{y}}} = \mu_y$
$\langle \hat{\mathbf{e}}, \mathbf{x} \rangle_s = 0$	$\hat{\beta}_2 = \sigma_{xy} / \sigma_x^2 = \rho_{xy} \sigma_y / \sigma_x$

Nótese el caso límite $\hat{\beta}_2 = 0$: entonces $\hat{\beta}_1 = \mu_y$, y recuperamos exactamente el ajuste de la lección 4 —la media—. La regresión simple generaliza aquella proyección sobre una recta añadiendo una segunda dirección; no la sustituye.

Dos ortogonalidades, dos coeficientes. Nada de matrices; solo Pitágoras y bilinealidad, aplicados dos veces.

Próxima sesión — Lección 8: las mismas dos ortogonalidades, combinadas con el teorema de Pitágoras, dan $\sigma_y^2 = \sigma_{\hat{\mathbf{y}}}^2 + \sigma_e^2$. De ahí saldrá el $R^2 = \cos^2 \theta$: ¿cómo de bueno es nuestro ajuste?

Recapitulación y guiño a la sesión siguiente

El recorrido de hoy ha sido, en el fondo, muy corto: partiendo de las dos condiciones de ortogonalidad que dejamos planteadas en la lección 6, y usando únicamente álgebra elemental (despejar, sustituir, dividir por n), hemos llegado a fórmulas cerradas para $\hat{\beta}_1$ y $\hat{\beta}_2$. Ninguna de las piezas era nueva: la media, la varianza, la covarianza y la correlación ya estaban completamente definidas y demostradas desde las lecciones 4 y 5. Lo único que hemos hecho hoy es plegar esas piezas sobre un problema de dos incógnitas.

Merece la pena señalar, aunque no lo desarrollemos aquí, que estas mismas fórmulas podrían obtenerse igualmente por la vía puramente analítica: minimizando $\|\hat{\mathbf{e}}\|_s^2 = \frac{1}{n} \sum_i (y_i - \beta_1 - \beta_2 x_i)^2$ respecto a β_1, β_2 e igualando a cero las dos derivadas parciales. Esas dos condiciones de primer orden son, de hecho, exactamente las dos condiciones de ortogonalidad que hemos resuelto en esta lección —el mismo fenómeno que ya vimos, con un solo coeficiente, en la lección 4—:

1. La primera condición de ortogonalidad ($\hat{\mathbf{e}} \perp \mathbf{1}$) fija el intercepto en función de la pendiente, $\hat{\beta}_1 = \mu_y - \hat{\beta}_2 \mu_x$, y garantiza que la recta ajustada pasa por el punto de las medias (μ_x, μ_y) , así como que $\mu_{\hat{\mathbf{y}}} = \mu_y$.
2. La segunda condición ($\hat{\mathbf{e}} \perp \mathbf{x}$) fija la pendiente como el cociente entre la covarianza y la varianza, $\hat{\beta}_2 = \sigma_{xy}/\sigma_x^2$, que puede reescribirse como la correlación reescalada por el cociente de las desviaciones típicas, $\hat{\beta}_2 = \rho_{xy} \sigma_y/\sigma_x$.

También vimos un camino alternativo —cambiar a generadores ortogonales, $\mathbf{1}$ y $\mathbf{x} - \mu_x \mathbf{1}$ — que llega al mismo resultado sin resolver ningún sistema, proyectando cada coeficiente de forma independiente. Este segundo camino es el que se generalizará, con más de un regresor, en la lección 10 (regresión múltiple). Allí ya no bastarán dos condiciones de ortogonalidad, sino tantas como regresores tengamos; y la notación matricial —anunciada en la lección 5, pero todavía no utilizada— empezará a ganarse su lugar como forma compacta de escribir esto mismo.

La sesión de laboratorio que sigue a esta lección (regresión simple con datos reales) comprobará numéricamente, con un dataset completo, que efectivamente $\sum \hat{\mathbf{e}}_i = 0$ y $\sum x_i \hat{\mathbf{e}}_i = 0$ —las dos condiciones de ortogonalidad de hoy—, y que las fórmulas de $\hat{\beta}_1, \hat{\beta}_2$ coinciden con las que reporta Gretl automáticamente al pedir una regresión.

Y la próxima sesión teórica, la lección 8, dará el siguiente paso natural. Las mismas dos condiciones de ortogonalidad, vistas bajo la óptica del teorema de Pitágoras —con $\mathbf{y} = \hat{\mathbf{y}} + \hat{\mathbf{e}}$ como suma de dos vectores ortogonales—, implican una descomposición de varianzas, $\sigma_y^2 = \sigma_{\hat{\mathbf{y}}}^2 + \sigma_{\hat{\mathbf{e}}}^2$. De ella surgirá la medida de bondad de ajuste R^2 , expresada geométricamente como el coseno al cuadrado de un ángulo. El ángulo de la lección 5 —el que definía la correlación— no ha desaparecido: solo cambia de protagonistas.

Para profundizar:

- Wooldridge, J. M. (2020), cap. 2, secciones 2.2-2.3: resumen algebraico completo de $\hat{\beta}_1, \hat{\beta}_2$ y sus propiedades.
- Bujosa, M. *Curso de Álgebra Lineal*, cap. 11: proyección ortogonal sobre subespacios generados por vectores ortogonales entre sí.