

Modelo Lineal

Graciela Boente¹

¹Universidad de Buenos Aires and CONICET, Argentina

Modelización y Predicción

Algunas de los métodos estadísticos más extendidos se ocupan de la modelización de datos y de la predicción.

Muchas de estas técnicas estadísticas se encuadran en lo que hoy se conoce como **aprendizaje estadístico** (AE).

Modelización y Predicción

Algunas de los métodos estadísticos más extendidos se ocupan de la modelización de datos y de la predicción.

Muchas de estas técnicas estadísticas se encuadran en lo que hoy se conoce como **aprendizaje estadístico** (AE).

El AE abarca una vasta cantidad de técnicas que ayudan a comprender los datos cuando se analizan varias variables al mismo tiempo, ya sea postulando modelos o encontrando relaciones entre las variables o estructuras que ayudan a su comprensión.

Aprendizaje Estadístico

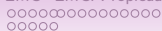
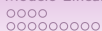
Algunos ejemplos de aprendizaje estadístico:

Vamos a considerar el último ejemplo:

En 97 pacientes que van a ser tratados con una prostatectomía radical se miden las siguientes variables:

- $x_1 = \text{lwight}$: log del peso de la próstata
- $x_2 = \text{age}$: edad
- $x_3 = \text{lbph}$: log de la cantidad de hiperplasia prostática benigna
- $x_4 = \text{svi}$: invasión seminal (si o no)
- $x_5 = \text{lcp}$: logaritmo de la penetración capsular
- $x_6 = \text{gleason}$: score de Gleason
- $x_7 = \text{pgg46}$: porcentaje de scores de Gleason 4 or 5.
- $x_8 = \text{lpsa}$: log de PSA

El objetivo es poder predecir el logaritmo del volumen del tumor ($y = \text{lcavol}$).



Diagramas de Dispersión

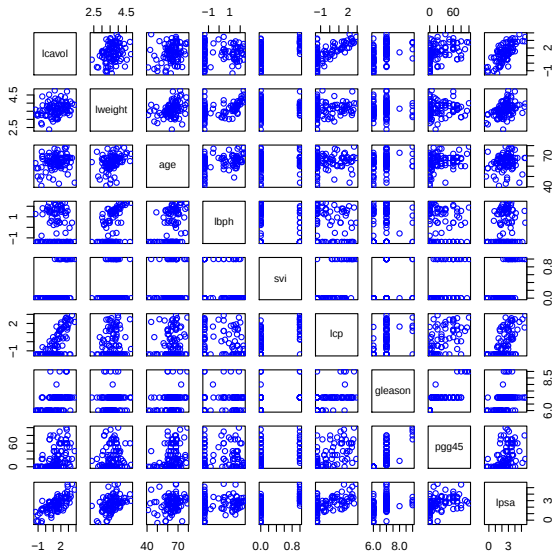
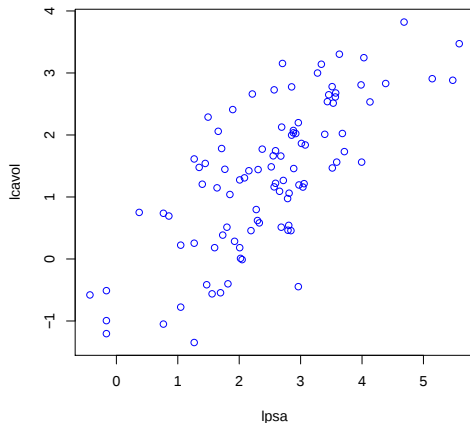


Diagrama de Dispersión



Modelos de regresión

En general, identificamos a

- la variable dependiente (o variable respuesta) con la letra y .
- las variables predictoras o covariables con la letra x .

Modelos de regresión

En general, identificamos a

- la variable dependiente (o variable respuesta) con la letra y .
- las variables predictoras o covariables con la letra x .

El modelo pretende describir cómo varía el comportamiento de $\mathbb{E}(\mathbf{y})$ bajo condiciones cambiantes del vector de p covariables $\mathbf{x}$.

Modelos de regresión

En general, identificamos a

- la variable dependiente (o variable respuesta) con la letra y .
- las variables predictoras o covariables con la letra x .

El modelo pretende describir cómo varía el comportamiento de $E(y)$ bajo condiciones cambiantes del vector de p covariables x .

En un modelo de regresión se postularía:

$$y = \eta(x_1, x_2, x_3, \dots, x_p) + \epsilon = \eta(\mathbf{x}) + \epsilon$$

Modelos de regresión

$$y = \eta(\mathbf{x}) + \epsilon$$

Las posibles funciones de regresión η pertenecen a una clase $\mathcal{G}$ tan grande que es frecuente que se simplifique el problema suponiendo cierta forma o ciertas propiedades de la función de regresión $\eta(\mathbf{x})$.

Una forma de simplificar el problema es suponer que la familia $\mathcal{G}$ puede expresarse en función de un número finito de constantes desconocidas, a estimar, llamadas **parámetros**, que controlan el comportamiento del modelo.

$$\mathcal{G} = \{\eta(\mathbf{x}) = g(\mathbf{x}, \boldsymbol{\theta}) : \boldsymbol{\theta} \in \Theta \subset \mathbb{R}^s\}$$

En este caso, la función $g(\cdot, \boldsymbol{\theta})$ está completamente determinada por $\boldsymbol{\theta}$ y por eso, diremos que el **modelo de regresión es paramétrico**.

Modelos de regresión

Para simplificar, pensemos que tenemos dos covariables: x_1 y x_2 .

Algunos ejemplos de modelos no paramétricos, semi-paramétricos y paramétricos cuando hay dos covariables son

Modelos de regresión

Para simplificar, pensemos que tenemos dos covariables: x_1 y x_2 .

Algunos ejemplos de modelos no paramétricos, semi-paramétricos y paramétricos cuando hay dos covariables son

Modelos no paramétricos

- (i) $y = \eta(x_1, x_2) + \varepsilon$ donde $\eta(x_1, x_2)$ es una función continua.
- (ii) $y = \eta(x_1, x_2) + \varepsilon$ donde $\eta(x_1, x_2)$ es una función continua y derivable.

Modelos de regresión

Para simplificar, pensemos que tenemos dos covariables: x_1 y x_2 .

Algunos ejemplos de modelos no paramétricos, semi-paramétricos y paramétricos cuando hay dos covariables son

Modelos no paramétricos

- (i) $y = \eta(x_1, x_2) + \varepsilon$ donde $\eta(x_1, x_2)$ es una función continua.
- (ii) $y = \eta(x_1, x_2) + \varepsilon$ donde $\eta(x_1, x_2)$ es una función continua y derivable.
- (iii) $y = \eta_1(x_1) + \eta_2(x_2) + \varepsilon$, con η_i funciones continuas. Modelo aditivo

Modelos de regresión

Modelos semi-paramétricos

- (i) $y = \eta(\boldsymbol{\theta}^T \mathbf{x}) + \varepsilon$ donde $\eta : \mathbb{R} \rightarrow \mathbb{R}$ es una función continua y $\boldsymbol{\theta} = (\theta_1, \theta_2)^T$ es tal que $\|\boldsymbol{\theta}\| = 1$. Modelo de índice simple

Modelos de regresión

Modelos semi-paramétricos

- (i) $y = \eta(\boldsymbol{\theta}^T \mathbf{x}) + \varepsilon$ donde $\eta : \mathbb{R} \rightarrow \mathbb{R}$ es una función continua y $\boldsymbol{\theta} = (\theta_1, \theta_2)^T$ es tal que $\|\boldsymbol{\theta}\| = 1$. Modelo de índice simple
- (ii) $y = \eta(x_1) + \theta_2 x_2 + \varepsilon$ donde $\eta : \mathbb{R} \rightarrow \mathbb{R}$ es una función continua y $\theta_2 \in \mathbb{R}$. Modelo parcialmente lineal

Modelos de regresión

Modelos semi-paramétricos

- (i) $y = \eta(\boldsymbol{\theta}^T \mathbf{x}) + \varepsilon$ donde $\eta : \mathbb{R} \rightarrow \mathbb{R}$ es una función continua y $\boldsymbol{\theta} = (\theta_1, \theta_2)^T$ es tal que $\|\boldsymbol{\theta}\| = 1$. Modelo de índice simple
- (ii) $y = \eta(x_1) + \theta_2 x_2 + \varepsilon$ donde $\eta : \mathbb{R} \rightarrow \mathbb{R}$ es una función continua y $\theta_2 \in \mathbb{R}$. Modelo parcialmente lineal
- (iii) $y = \theta_1 + x_2 \eta(x_1) + \varepsilon$, donde $\eta : \mathbb{R} \rightarrow \mathbb{R}$ es una función continua y $\theta_1 \in \mathbb{R}$. Modelo de coeficientes variables

Modelo Lineal

Uno de los modelos más sencillos es el **modelo lineal**, en el que los parámetros intervienen como simples coeficientes de las variables independientes o de funciones de éstas.

Modelo Lineal

Algunos ejemplos sencillos de modelos lineales dependientes de una sola variable son:

$$g(x, \theta) = \alpha + \beta x$$

$$\theta = (\alpha, \beta)^T$$

$$g(x, \theta) = \alpha + \beta x + \gamma x^2$$

$$\theta = (\alpha, \beta, \gamma)^T$$

$$g(x, \theta) = \alpha + \beta \log(x)$$

$$\theta = (\alpha, \beta)^T$$

Hipótesis del modelo

Consideremos una muestra $(y_i, \mathbf{x}_i)$ con $\mathbf{x}_i = (x_{i1}, \dots, x_{ip})^T$, $1 \leq i \leq n$ que cumplen el modelo lineal

$$y_i = \theta_1 x_{i1} + \dots + \theta_p x_{ip} + \epsilon_i = \boldsymbol{\theta}^T \mathbf{x}_i + \epsilon_i.$$

Supondremos x_{ij} , $1 \leq j \leq p$ son determinísticas, es decir, no son variables aleatorias.

Hipótesis del modelo

Las hipótesis que usaremos son

$$\text{H1 } \mathbb{E}(\epsilon_i) = 0, \text{ para } 1 \leq i \leq n$$

$$\text{H2 } \text{VAR}(\epsilon_i) = \sigma^2, \text{ para } 1 \leq i \leq n$$

$$\text{COV}(\epsilon_i, \epsilon_\ell) = 0, \text{ para } 1 \leq i \neq \ell \leq n$$

Hipótesis del modelo

Las hipótesis que usaremos son

$$\text{H1 } \mathbb{E}(\epsilon_i) = 0, \text{ para } 1 \leq i \leq n$$

$$\text{H2 } \text{VAR}(\epsilon_i) = \sigma^2, \text{ para } 1 \leq i \leq n$$

$$\text{COV}(\epsilon_i, \epsilon_\ell) = 0, \text{ para } 1 \leq i \neq \ell \leq n$$

$$\text{H3 } \epsilon_1, \dots, \epsilon_n \text{ son i.i.d. } N(0, \sigma^2)$$

Enfoque matricial

Queremos escribir el modelo

$$y_i = \theta_1 x_{i1} + \cdots + \theta_p x_{ip} + \epsilon_i = \boldsymbol{\theta}^T \mathbf{x}_i + \epsilon_i .$$

en forma matricial.

Enfoque matricial

Queremos escribir el modelo

$$y_i = \theta_1 x_{i1} + \cdots + \theta_p x_{ip} + \epsilon_i = \boldsymbol{\theta}^T \mathbf{x}_i + \epsilon_i.$$

en forma matricial. Sean

$$\mathbf{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \quad \mathbf{X} = \begin{pmatrix} \mathbf{x}_1^T \\ \vdots \\ \mathbf{x}_n^T \end{pmatrix} \quad \boldsymbol{\epsilon} = \begin{pmatrix} \epsilon_1 \\ \vdots \\ \epsilon_n \end{pmatrix}$$



$$\Omega : \quad \mathbf{y} = \mathbf{X} \boldsymbol{\theta} + \boldsymbol{\epsilon}$$

$n \times 1 \qquad n \times p \quad p \times 1 \qquad n \times 1$

20

Algunos ejemplos: Modelo de regresión simple

En el caso más sencillo de regresión simple, tendríamos

$$y_i = \theta_1 + \theta_2 x_i + \epsilon_i \quad 1 \leq i \leq n$$

$$p = 2 \quad \mathbf{X} = \begin{pmatrix} 1 & x_1 \\ 1 & x_2 \\ \vdots & \vdots \\ 1 & x_n \end{pmatrix} \quad \boldsymbol{\theta} = \begin{pmatrix} \theta_1 \\ \theta_2 \end{pmatrix}$$

Modelo de regresión a través del origen

Es un modelo de regresión simple donde se supone que la ordenada al origen es 0

$$y_i = \theta_1 x_i + \epsilon_i \quad 1 \leq i \leq n$$

$$p = 1 \quad \mathbf{X} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \quad \boldsymbol{\theta} = \theta_1$$

Modelo de posición

$$y_i = \mu + \epsilon_i \quad 1 \leq i \leq n$$

$$p = 1 \quad \mathbf{X} = \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix} \quad \theta = \mu$$

Propiedades

Sean $\mathbf{V} = (V_{ij}) \in \mathbb{R}^{\ell \times m}$ una matriz de variables aleatorias, $\mathbf{x} \in \mathbb{R}^p$ e $\mathbf{y} \in \mathbb{R}^q$ dos vectores aleatorios

- $\mathbb{E}(\mathbf{A}\mathbf{V}\mathbf{B} + \mathbf{C}) = \mathbf{A} \mathbb{E}(\mathbf{V}) \mathbf{B} + \mathbf{C}$, con $\mathbf{A} \in \mathbb{R}^{s \times \ell}$, $\mathbf{B} \in \mathbb{R}^{m \times k}$ y $\mathbf{C} \in \mathbb{R}^{s \times k}$ matrices fijas.
- Si $\mathbf{x}$ e $\mathbf{y}$ son independientes, entonces $\text{Cov}(\mathbf{x}, \mathbf{y}) = \mathbf{0}$.
- $\text{Cov}(\mathbf{A}\mathbf{x}, \mathbf{B}\mathbf{y}) = \mathbf{A} \text{Cov}(\mathbf{x}, \mathbf{y}) \mathbf{B}^T$, siendo $\mathbf{A} \in \mathbb{R}^{s \times p}$ y $\mathbf{B} \in \mathbb{R}^{\ell \times q}$ matrices fijas.
- $\text{VAR}(\mathbf{A}\mathbf{x} + \mathbf{d}) = \mathbf{A} \text{VAR}(\mathbf{x}) \mathbf{A}^T$ siendo $\mathbf{A} \in \mathbb{R}^{s \times p}$ una matriz fija y $\mathbf{d} \in \mathbb{R}^s$ un vector fijo.

El modelo lineal

El modelo lineal puede entonces escribirse como:

$$\Omega : \mathbf{y} = \mathbf{X}\boldsymbol{\theta} + \boldsymbol{\epsilon} \quad \text{con} \quad \mathbb{E}(\boldsymbol{\epsilon}) = \mathbf{0} \quad \boldsymbol{\Sigma}_{\boldsymbol{\epsilon}} = \sigma^2 \mathbf{I}$$

El modelo lineal: Interpretación

Supondremos que $n \geq p$,

$$\mathbf{X} = \begin{pmatrix} \mathbf{x}_1^T \\ \vdots \\ \mathbf{x}_n^T \end{pmatrix} = (\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(p)})$$

Sea $\mathcal{V}$ el subespacio de $\mathbb{R}^n$ generado por las columnas de $\mathbf{X}$.

$$\dim(\mathcal{V}) = \text{rango}(\mathbf{X})$$

El modelo lineal: Interpretación

Supondremos que $n \geq p$,

$$\mathbf{X} = \begin{pmatrix} \mathbf{x}_1^T \\ \vdots \\ \mathbf{x}_n^T \end{pmatrix} = (\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(p)})$$

Sea $\mathcal{V}$ el subespacio de $\mathbb{R}^n$ generado por las columnas de $\mathbf{X}$.

$$\dim(\mathcal{V}) = \text{rango}(\mathbf{X})$$

La hipótesis H1 dice que

$$\mathbb{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\theta} = \sum_{j=1}^p \theta_j \mathbf{x}^{(j)} \in \mathcal{V}$$

El modelo lineal: Interpretación

Supondremos que $n \geq p$,

$$\mathbf{X} = \begin{pmatrix} \mathbf{x}_1^T \\ \vdots \\ \mathbf{x}_n^T \end{pmatrix} = (\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(p)})$$

Sea $\mathcal{V}$ el subespacio de $\mathbb{R}^n$ generado por las columnas de $\mathbf{X}$.

$$\dim(\mathcal{V}) = \text{rango}(\mathbf{X})$$

La hipótesis H1 dice que

$$\mathbb{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\theta} = \sum_{j=1}^p \theta_j \mathbf{x}^{(j)} \in \mathcal{V}$$

Luego, H1 es equivalente a decir que $\mathbb{E}(\mathbf{y}) \in \mathcal{V}$

El modelo lineal: Interpretación

H1 es equivalente a decir que $\mathbb{E}(\mathbf{y}) \in \mathcal{V}$

El modelo lineal: Interpretación

H1 es equivalente a decir que $\mathbb{E}(\mathbf{y}) \in \mathcal{V}$

- Si $\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(p)}$ son linealmente independientes o sea, si $\dim(\mathcal{V}) = \text{rango}(\mathbf{X}) = p$, existe un **único** $\boldsymbol{\theta}$ tal que $\mathbb{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\theta}$.

El modelo lineal: Interpretación

H1 es equivalente a decir que $\mathbb{E}(\mathbf{y}) \in \mathcal{V}$

- Si $\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(p)}$ son linealmente independientes o sea, si $\dim(\mathcal{V}) = \text{rango}(\mathbf{X}) = p$, existe un **único** $\boldsymbol{\theta}$ tal que $\mathbb{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\theta}$.
- Si $\dim(\mathcal{V}) = q < p$, $\boldsymbol{\theta}$ **no está determinado** aunque $\mathbb{E}(\mathbf{y})$ sí lo está. En este caso, tenemos dos opciones

El modelo lineal: Interpretación

H1 es equivalente a decir que $\mathbb{E}(\mathbf{y}) \in \mathcal{V}$

- Si $\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(p)}$ son linealmente independientes o sea, si $\dim(\mathcal{V}) = \text{rango}(\mathbf{X}) = p$, existe un **único** $\boldsymbol{\theta}$ tal que $\mathbb{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\theta}$.
- Si $\dim(\mathcal{V}) = q < p$, $\boldsymbol{\theta}$ **no está determinado** aunque $\mathbb{E}(\mathbf{y})$ sí lo está. En este caso, tenemos dos opciones
 - Si $\dim(\mathcal{V}) = q$ y $\mathbf{x}^{(j_1)}, \dots, \mathbf{x}^{(j_q)}$ generan $\mathcal{V}$, tenemos que

$$\mathbb{E}(\mathbf{y}) = \sum_{s=1}^q \tilde{\theta}_s \mathbf{x}^{(j_s)}$$

y el vector $\tilde{\theta} = (\tilde{\theta}_1, \dots, \tilde{\theta}_g)^T$ está bien determinado.

El modelo lineal: Interpretación

H1 es equivalente a decir que $\mathbb{E}(\mathbf{y}) \in \mathcal{V}$

- Si $\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(p)}$ son linealmente independientes o sea, si $\dim(\mathcal{V}) = \text{rango}(\mathbf{X}) = p$, existe un **único** $\boldsymbol{\theta}$ tal que $\mathbb{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\theta}$.
- Si $\dim(\mathcal{V}) = q < p$, $\boldsymbol{\theta}$ **no está determinado** aunque $\mathbb{E}(\mathbf{y})$ sí lo está. En este caso, tenemos dos opciones
 - Si $\dim(\mathcal{V}) = q$ y $\mathbf{x}^{(j_1)}, \dots, \mathbf{x}^{(j_q)}$ generan $\mathcal{V}$, tenemos que

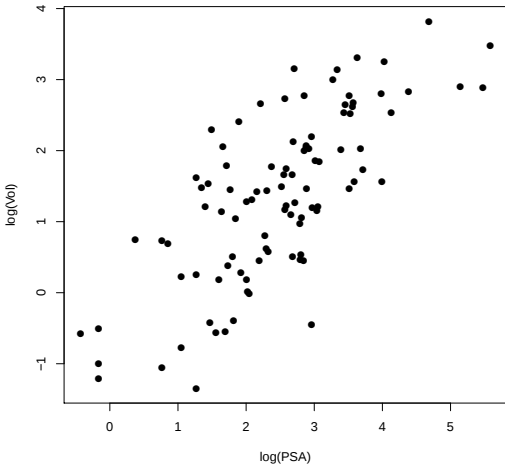
$$\mathbb{E}(\mathbf{y}) = \sum_{s=1}^q \tilde{\theta}_s \mathbf{x}^{(j_s)}$$

y el vector $\tilde{\boldsymbol{\theta}} = (\tilde{\theta}_1, \dots, \tilde{\theta}_q)^T$ está bien determinado.

- La otra opción es **imponer restricciones lineales a $\boldsymbol{\theta}$** mediante

$$\mathbf{A}\boldsymbol{\theta} = \mathbf{0}$$

con $\mathbf{A}$ una matriz elegida para que el vector quede unívocamente determinado.



¿Cómo estimamos los parámetros?

Mínimos Cuadrados

Si los puntos en un gráfico parecen seguir una recta como en la figura anterior, el problema es elegir la recta que mejor ajusta los puntos.

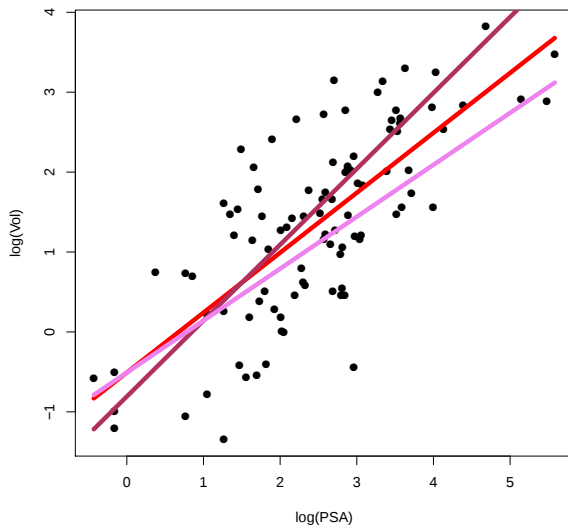
¿Cómo estimamos los parámetros?

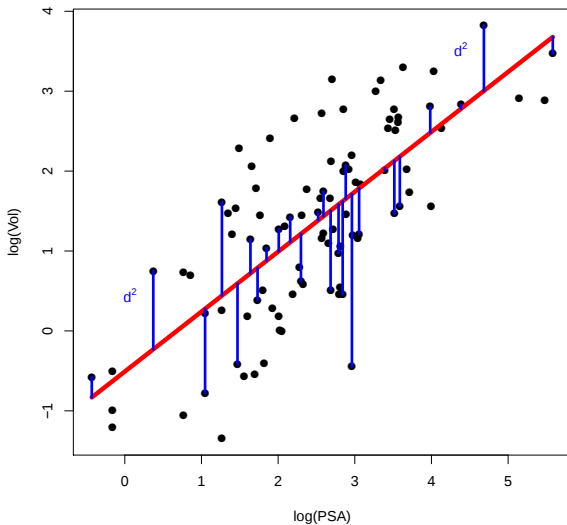
Mínimos Cuadrados

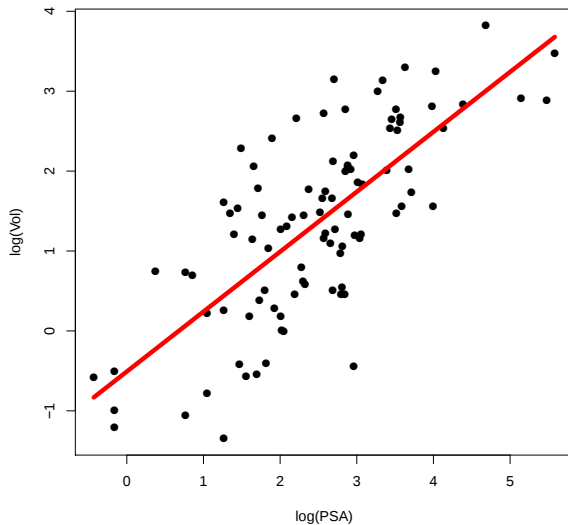
Si los puntos en un gráfico parecen seguir una recta como en la figura anterior, el problema es elegir la recta que mejor ajusta los puntos.

Tendremos en cuenta:

- tomar una distancia promedio de la recta a todos los puntos
- mover la recta hasta que esta distancia promedio sea la menor posible.







Repaso: Mínimos Cuadrados

En el caso de la regresión simple en el cual $g(x, \theta_1, \theta_2) = \theta_1 + \theta_2 x$,

$$\hat{\theta} = \underset{\mathbf{b} \in \mathbb{R}^2}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n [y_i - (b_1 + b_2 x_i)]^2.$$

Esta suma se llama **la suma de cuadrados residual para la recta**

La recta resultante **recta de cuadrados mínimos**.

Caso Modelo Lineal

El estimador de mínimos cuadrados de θ , es un conjunto de funciones de $\mathbf{y}$,

$$\hat{\theta}_1 = \hat{\theta}_1(\mathbf{y}), \quad \hat{\theta}_2 = \hat{\theta}_2(\mathbf{y}), \dots \quad \hat{\theta}_p = \hat{\theta}_p(\mathbf{y})$$

que minimiza $\mathcal{S}(\mathbf{b})$.

$$\mathcal{S}(\hat{\theta}) = \underset{\mathbf{b} \in \mathbb{R}^p}{\operatorname{argmin}} \mathcal{S}(\mathbf{b})$$

Caso Modelo Lineal

El estimador de mínimos cuadrados de θ , es un conjunto de funciones de $\mathbf{y}$,

$$\hat{\theta}_1 = \hat{\theta}_1(\mathbf{y}), \quad \hat{\theta}_2 = \hat{\theta}_2(\mathbf{y}), \dots \quad \hat{\theta}_p = \hat{\theta}_p(\mathbf{y})$$

que minimiza $\mathcal{S}(\mathbf{b})$.

$$\mathcal{S}(\hat{\theta}) = \underset{\mathbf{b} \in \mathbb{R}^p}{\operatorname{argmin}} \mathcal{S}(\mathbf{b})$$

Veremos que el EMC siempre existe, aunque no siempre es único.

Caso de errores normales

Bajo H3, o sea si $\epsilon \sim N(0, \sigma^2 \mathbf{I})$ se tiene que

- $\hat{\theta}$ es el EMV de θ si y sólo si

$$\mathcal{S}(\hat{\theta}) = \underset{\mathbf{b} \in \mathbb{R}^p}{\operatorname{argmin}} \mathcal{S}(\mathbf{b})$$

- El EMV de σ^2 es

$$\hat{\sigma}^2 = \frac{1}{n} \mathcal{S}(\hat{\theta})$$

Ecuaciones Normales

$\hat{\theta}$ minimiza $\mathcal{S}(\mathbf{b})$, por lo tanto, derivando e igualando a 0 obtenemos las **ecuaciones normales**.

Los estimadores de mínimos cuadrados $\hat{\theta}_1, \dots, \hat{\theta}_p$ cumplen:

$$\frac{\partial \mathcal{S}(\mathbf{b})}{\partial b_k} \Big|_{\mathbf{b}=\hat{\theta}} = -2 \sum_{i=1}^n \left(y_i - \sum_{j=1}^p x_{ij} b_j \right) x_{ik} \Big|_{\mathbf{b}=\hat{\theta}} = 0$$

Por lo tanto, para $1 \leq k \leq p$

$$\sum_{i=1}^n y_i x_{ik} = \sum_{i=1}^n \sum_{j=1}^p x_{ij} x_{ik} \hat{\theta}_j$$

reordenando la suma

$$\sum_{i=1}^n y_i x_{ik} = \sum_{j=1}^p \hat{\theta}_j \sum_{i=1}^n x_{ij} x_{ik}$$

Ecuaciones Normales

Como

$$\sum_{i=1}^n y_i x_{ik} = \sum_{j=1}^p \hat{\theta}_j \sum_{i=1}^n x_{ij} x_{ik} \quad 1 \leq k \leq p$$

Ecuaciones Normales

Como

$$\sum_{i=1}^n y_i x_{ik} = \sum_{j=1}^p \hat{\theta}_j \sum_{i=1}^n x_{ij} x_{ik} \quad 1 \leq k \leq p$$

El estimador de mínimos cuadrados cumple estas p ecuaciones, entonces

$$\mathbf{X}^T \mathbf{X} \hat{\boldsymbol{\theta}} = \mathbf{X}^T \mathbf{y},$$

que se conocen como **ecuaciones normales**.

Ecuaciones Normales

Como

$$\sum_{i=1}^n y_i x_{ik} = \sum_{j=1}^p \hat{\theta}_j \sum_{i=1}^n x_{ij} x_{ik} \quad 1 \leq k \leq p$$

El estimador de mínimos cuadrados cumple estas p ecuaciones, entonces

$$\mathbf{X}^T \mathbf{X} \hat{\boldsymbol{\theta}} = \mathbf{X}^T \mathbf{y},$$

que se conocen como **ecuaciones normales**.

Si $\text{rango}(\mathbf{X}) = p$, $\mathbf{X}^T \mathbf{X}$ es no singular, la solución es única y resulta

$$\hat{\boldsymbol{\theta}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}.$$

Ejemplo

En el caso de regresión simple tendríamos

$$\mathbf{X}^T \mathbf{X} = \begin{pmatrix} 1 & 1 & 1 & \dots & 1 \\ x_1 & x_2 & x_3 & \dots & x_n \end{pmatrix} \begin{pmatrix} 1 & x_1 \\ 1 & x_2 \\ \cdot & \cdot \\ \cdot & \cdot \\ 1 & x_n \end{pmatrix}$$

El sistema sería

$$\begin{pmatrix} n & \sum_{i=1}^n x_i \\ \sum_{i=1}^n x_i & \sum_{i=1}^n x_i^2 \end{pmatrix} \begin{pmatrix} \hat{\theta}_1 \\ \hat{\theta}_2 \end{pmatrix} = \begin{pmatrix} \sum_{i=1}^n y_i \\ \sum_{i=1}^n x_i y_i \end{pmatrix}$$

Ejemplo

La inversa resulta

$$(\mathbf{X}^T \mathbf{X})^{-1} = \frac{1}{n \sum_{i=1}^n x_i^2 - n^2 \bar{x}^2} \begin{pmatrix} \sum_{i=1}^n x_i^2 & -\sum_{i=1}^n x_i \\ -\sum_{i=1}^n x_i & n \end{pmatrix}$$

y además

$$\mathbf{X}^T \mathbf{y} = \begin{pmatrix} \sum_{i=1}^n y_i \\ \sum_{i=1}^n x_i y_i \end{pmatrix}$$

y por lo tanto

$$\hat{\theta} = \begin{pmatrix} \hat{\theta}_1 \\ \hat{\theta}_2 \end{pmatrix} = \frac{1}{n \sum_{i=1}^n (x_i - \bar{x})^2} \begin{pmatrix} (\sum_{i=1}^n y_i)(\sum_{i=1}^n x_i^2) - (\sum_{i=1}^n x_i)(\sum_{i=1}^n x_i y_i) \\ n \sum_{i=1}^n x_i y_i - (\sum_{i=1}^n y_i)(\sum_{i=1}^n x_i) \end{pmatrix}$$

Ejemplo

Entonces

$$\hat{\theta}_1 = \bar{y} - \bar{x} \hat{\theta}_2$$

y por otro lado

$$\hat{\theta}_2 = \frac{\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}}{\sum_{i=1}^n x_i^2 - n \bar{x}^2} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

Interpretación Geométrica del EMC

Nuestro modelo plantea

$$\Omega : \quad \begin{aligned} \mathbb{E}(\mathbf{y}) &= \mathbf{X}\boldsymbol{\theta} \\ \boldsymbol{\Sigma}_{\mathbf{y}} &= \sigma^2 \mathbf{I} \end{aligned}$$

Luego, si

$$\boldsymbol{\eta} = \mathbb{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\theta}$$

Tenemos

$$\boldsymbol{\eta} = \theta_1 \mathbf{x}^{(1)} + \cdots + \theta_p \mathbf{x}^{(p)} \in \mathcal{V}$$

Interpretación Geométrica del EMC

$$\mathcal{S}(\hat{\theta}) = \|\mathbf{y} - \mathbf{X}\hat{\theta}\|^2 = \min_{\mathbf{b} \in \mathbb{R}^p} \|\mathbf{y} - \mathbf{X}\mathbf{b}\|^2 = \min_{\mathbf{z} \in \mathcal{V}} \|\mathbf{y} - \mathbf{z}\|^2$$

Interpretación Geométrica del EMC

$$\mathcal{S}(\hat{\theta}) = \|\mathbf{y} - \mathbf{X}\hat{\theta}\|^2 = \min_{\mathbf{b} \in \mathbb{R}^p} \|\mathbf{y} - \mathbf{X}\mathbf{b}\|^2 = \min_{\mathbf{z} \in \mathcal{V}} \|\mathbf{y} - \mathbf{z}\|^2$$

Es decir, que

$$\hat{\eta} = \mathbf{X}\hat{\theta}$$

es la proyección ortogonal de $\mathbf{y}$ sobre $\mathcal{V}$.

Interpretación Geométrica del EMC

Por lo tanto, para todo $\mathbf{z} \in \mathcal{V}$

$$\mathbf{z}^T(\mathbf{y} - \hat{\boldsymbol{\eta}}) = 0$$

Interpretación Geométrica del EMC

Por lo tanto, para todo $\mathbf{z} \in \mathcal{V}$

$$\mathbf{z}^T(\mathbf{y} - \hat{\boldsymbol{\eta}}) = 0$$

en particular,

$$(\mathbf{y} - \hat{\boldsymbol{\eta}})^T \mathbf{x}^{(j)} = 0 \quad 1 \leq j \leq p \quad \Leftrightarrow \quad \mathbf{X}^T(\mathbf{y} - \hat{\boldsymbol{\eta}}) = \mathbf{0}$$

Interpretación Geométrica del EMC

Por lo tanto, para todo $\mathbf{z} \in \mathcal{V}$

$$\mathbf{z}^T(\mathbf{y} - \hat{\eta}) = 0$$

en particular,

$$(\mathbf{y} - \hat{\boldsymbol{\eta}})^T \mathbf{x}^{(j)} = 0 \quad 1 \leq j \leq p \quad \Leftrightarrow \quad \mathbf{X}^T (\mathbf{y} - \hat{\boldsymbol{\eta}}) = \mathbf{0}$$

En términos de las ecuaciones normales tenemos que:

$$\mathbf{X}^T \hat{\boldsymbol{\eta}} = \mathbf{X}^T \mathbf{y}$$

Interpretación Geométrica del EMC

Por lo tanto, para todo $\mathbf{z} \in \mathcal{V}$

$$\mathbf{z}^T(\mathbf{y} - \hat{\eta}) = 0$$

en particular,

$$(\mathbf{y} - \hat{\boldsymbol{\eta}})^T \mathbf{x}^{(j)} = 0 \quad 1 \leq j \leq p \quad \Leftrightarrow \quad \mathbf{X}^T (\mathbf{y} - \hat{\boldsymbol{\eta}}) = \mathbf{0}$$

En términos de las ecuaciones normales tenemos que:

$$\mathbf{X}^T \hat{\boldsymbol{\eta}} = \mathbf{X}^T \mathbf{y}$$

$$\mathbf{X}^T \mathbf{X} \hat{\boldsymbol{\theta}} = \mathbf{X}^T \mathbf{y}$$

Interpretación Geométrica del EMC

Es decir, tenemos que

Proposición 1.

- El EMC de θ existe y es cualquier solución de

$$\hat{\eta} = \mathbf{X}\hat{\theta}$$

- $\hat{\theta}$ es EMC si y sólo si

$$\mathbf{X}^T \mathbf{X} \hat{\theta} = \mathbf{X}^T \mathbf{y}$$

- El EMC es único si $\text{rango}(\mathbf{X}) = p$. Si $\text{rango}(\mathbf{X}) = p$

$$\hat{\theta} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

Interpretación Geométrica del EMC

Proposición 2.

- El valor mínimo de $\mathcal{S}(\mathbf{b})$ es igual a

$$\mathcal{S}(\hat{\theta}) = \|\mathbf{y}\|^2 - \|\mathbf{X}\hat{\theta}\|^2 = \mathbf{y}^T \mathbf{y} - \mathbf{y}^T \mathbf{X}\hat{\theta}$$

- Si $\text{rango}(\mathbf{X}) = p$

$$\mathcal{S}(\hat{\theta}) = \mathbf{y}^T \mathbf{y} - \mathbf{y}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

De ahora en más $\text{rango}(\mathbf{X}) = p$

Por lo tanto,

$$\hat{\theta} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

$$\hat{\theta} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$
$$\hat{\eta} = \mathbf{X} \hat{\theta} = \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} = \mathbf{P} \mathbf{y}$$

$$\text{rango}(\mathbf{X}) = p$$

El vector de residuos se escribe como

$$\begin{aligned}\mathbf{r}(\hat{\theta}) &= \mathbf{y} - \hat{\mathbf{y}} = \mathbf{y} - \mathbf{X}\hat{\theta} \\ &= \mathbf{y} - \mathbf{X}(\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T\mathbf{y}\end{aligned}$$

$$\text{rango}(\mathbf{X}) = p$$

El vector de residuos se escribe como

$$\begin{aligned} \mathbf{r}(\hat{\theta}) &= \mathbf{y} - \hat{\mathbf{y}} = \mathbf{y} - \mathbf{X}\hat{\theta} \\ &= \mathbf{y} - \mathbf{X}(\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T\mathbf{y} \\ &= \mathbf{y} - \mathbf{P}\mathbf{y} = (\mathbf{I} - \mathbf{P})\mathbf{y} \end{aligned}$$

donde

$$\mathbf{P} = \mathbf{X}(\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T \in \mathbb{R}^{n \times n}$$

Suma de Cuadrados

Tenemos que

$$\sum_{i=1}^n (y_i - \hat{y}_i)^2 = \|\mathbf{y} - \mathbf{P}\mathbf{y}\|^2$$

Propiedades del Estimador de Mínimos Cuadrados

Usando la notación matricial escribimos el modelo como

$$\begin{aligned}\Omega : \quad \mathbf{y} &= \mathbf{X}\boldsymbol{\theta} + \boldsymbol{\epsilon} \\ \mathbb{E}(\boldsymbol{\epsilon}) &= \mathbf{0} \\ \boldsymbol{\Sigma}_{\boldsymbol{\epsilon}} &= \sigma^2 \mathbf{I}\end{aligned}$$

Sea $\mathbf{a} \in \mathbb{R}^p$ y $\varphi = \mathbf{a}^T \boldsymbol{\theta}$.

Como $\text{rango}(\mathbf{X}) = p$, definimos el EMC de φ como

$$\hat{\varphi} = \mathbf{a}^T \hat{\boldsymbol{\theta}}$$

Estimación de σ^2 : $\text{rango}(\mathbf{X}) = p$

Veremos que las varianzas de los estimadores dependen del diseño y σ^2 , que es desconocida.

Propiedades

Proposición 3 (continuación).

- Si se cumplen H1 y H2,

$$\Sigma_{\hat{\theta}} = \text{VAR}(\hat{\theta}) = \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}$$

$$\sigma_{\hat{\varphi}}^2 = \text{VAR}(\hat{\varphi}) = \sigma^2 \mathbf{a}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{a}$$

Además, podemos obtener estimadores insesgados de $\Sigma_{\hat{\theta}}$ y de $\sigma_{\hat{\varphi}}^2$ como

$$\hat{\Sigma}_{\hat{\theta}} = \text{VAR}(\hat{\theta}) = s^2 (\mathbf{X}^T \mathbf{X})^{-1}$$

$$\hat{\sigma}_{\hat{\varphi}}^2 = s^2 \mathbf{a}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{a}$$

Propiedades

Sea

- $\mathbf{v}_1, \dots, \mathbf{v}_p \in \mathbb{R}^n$ una base ortonormal de $\mathcal{V}$

Propiedades

Sea

- $\mathbf{v}_1, \dots, \mathbf{v}_p \in \mathbb{R}^n$ una base ortonormal de $\mathcal{V}$
- $\mathbf{v}_{p+1}, \dots, \mathbf{v}_n \in \mathbb{R}^n$ tales que $\{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ es b.o. de $\mathbb{R}^n$
- Definamos $\mathbf{B} = (\mathbf{v}_1, \dots, \mathbf{v}_n)$ $\mathbf{z} = \mathbf{B}^T \mathbf{y}$.

Propiedades

Luego,

- $\mathcal{S}(\hat{\theta}) = \sum_{i=p+1}^n z_i^2$

Propiedades

Luego,

- $\mathcal{S}(\hat{\boldsymbol{\theta}}) = \sum_{i=p+1}^n z_i^2$
- $\mathbb{E}(\mathbf{z}) = \boldsymbol{\gamma} = (\gamma_1, \dots, \gamma_p, 0, \dots, 0)^T$ con $\gamma_j = \mathbf{v}_j^T \boldsymbol{\eta}$.

Propiedades

Luego,

- $\mathcal{S}(\hat{\boldsymbol{\theta}}) = \sum_{i=p+1}^n z_i^2$
- $\mathbb{E}(\mathbf{z}) = \boldsymbol{\gamma} = (\gamma_1, \dots, \gamma_p, 0, \dots, 0)^T$ con $\gamma_j = \mathbf{v}_j^T \boldsymbol{\eta}$.
- $\text{VAR}(\mathbf{z}) = \sigma^2 \mathbf{I}$

Errores Normales

Si los errores además de ser independientes con media 0 y varianza σ^2 , son normales, es decir

$$\epsilon_i \sim N(0, \sigma^2) \quad i = 1, \dots, n$$

el modelo Ω puede formularse de la siguiente forma:

$$\Omega : \quad \mathbf{y} \sim N(\mathbf{X}\boldsymbol{\theta}, \sigma^2 \mathbf{I})$$

Vimos que el estimador de máxima verosimilitud de $\boldsymbol{\theta}$ coincide con el de mínimos cuadrados.

Propiedad de EMC

Teorema 2. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\boldsymbol{\theta} \in \mathbb{R}^p$ y vale H3. Entonces,

- i) $\hat{\theta} \sim N(\theta, \sigma^2(\mathbf{X}^T \mathbf{X})^{-1})$

Propiedad de EMC

Teorema 2. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\theta \in \mathbb{R}^p$ y vale H3. Entonces,

i) $\hat{\theta} \sim N(\theta, \sigma^2(\mathbf{X}^T \mathbf{X})^{-1})$

ii) $\frac{(\hat{\theta} - \theta)^T (\mathbf{X}^T \mathbf{X}) (\hat{\theta} - \theta)}{\sigma^2} \sim \chi_p^2$

iii) $\frac{(n-p)s^2}{\sigma^2} \sim \chi_{n-p}^2$

iv) $\hat{\theta}$ y s^2 son independientes

v) $\frac{(\hat{\theta} - \theta)^T (\mathbf{X}^T \mathbf{X}) (\hat{\theta} - \theta)}{p s^2} \sim \mathcal{F}_{p, n-p}$

Propiedad de EMC

Teorema 3. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\boldsymbol{\theta} \in \mathbb{R}^p$ y vale H3. Entonces,

- i) $\hat{\varphi} = \mathbf{a}^T \hat{\boldsymbol{\theta}}$ es un estimador IMVU para $\varphi = \mathbf{a}^T \boldsymbol{\theta}$
- ii) $\text{VAR}(\hat{\varphi}) = \sigma^2 \mathbf{a}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{a} = \sigma_{\hat{\varphi}}^2$
- iii) $\hat{\varphi} \sim N(\varphi, \sigma^2 \mathbf{a}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{a}) = N(\varphi, \sigma_{\hat{\varphi}}^2)$
- iv) Sea $\hat{\sigma}_{\hat{\varphi}}^2 = s^2 \mathbf{a}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{a}$ entonces

$$\frac{(\hat{\varphi} - \varphi)}{\hat{\sigma}_{\hat{\varphi}}} \sim \mathcal{T}_{n-p}$$

Propiedad de EMC

Teorema 4. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\boldsymbol{\theta} \in \mathbb{R}^p$ y vale H3.

Sea $\mathbf{A} \in \mathbb{R}^{k \times p}$ con $\text{rango}(\mathbf{A}) = k \leq p$ y definamos

$$\boldsymbol{\psi} = \mathbf{A}\boldsymbol{\theta} \quad \hat{\boldsymbol{\psi}} = \mathbf{A}\hat{\boldsymbol{\theta}}$$

Entonces,

- i) $\hat{\boldsymbol{\psi}}$ es un estimador IMVU para $\boldsymbol{\psi}$

Propiedad de EMC

Teorema 4. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\boldsymbol{\theta} \in \mathbb{R}^p$ y vale H3.

Sea $\mathbf{A} \in \mathbb{R}^{k \times p}$ con $\text{rango}(\mathbf{A}) = k \leq p$ y definamos

$$\boldsymbol{\psi} = \mathbf{A}\boldsymbol{\theta} \quad \hat{\boldsymbol{\psi}} = \mathbf{A}\hat{\boldsymbol{\theta}}$$

Entonces,

- i) $\hat{\boldsymbol{\psi}}$ es un estimador IMVU para $\boldsymbol{\psi}$
- ii) $\text{VAR}(\hat{\boldsymbol{\psi}}) = \sigma^2 \mathbf{A}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{A} = \boldsymbol{\Sigma}_{\hat{\boldsymbol{\psi}}}$

Propiedad de EMC

Teorema 4. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\boldsymbol{\theta} \in \mathbb{R}^p$ y vale H3.

Sea $\mathbf{A} \in \mathbb{R}^{k \times p}$ con $\text{rango}(\mathbf{A}) = k \leq p$ y definamos

$$\boldsymbol{\psi} = \mathbf{A}\boldsymbol{\theta} \quad \hat{\boldsymbol{\psi}} = \mathbf{A}\hat{\boldsymbol{\theta}}$$

Entonces,

- i) $\hat{\boldsymbol{\psi}}$ es un estimador IMVU para $\boldsymbol{\psi}$
- ii) $\text{VAR}(\hat{\boldsymbol{\psi}}) = \sigma^2 \mathbf{A}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{A} = \boldsymbol{\Sigma}_{\hat{\boldsymbol{\psi}}}$
- iii) $\hat{\boldsymbol{\psi}} \sim N(\boldsymbol{\psi}, \sigma^2 \mathbf{A}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{A}) = N(\boldsymbol{\psi}, \boldsymbol{\Sigma}_{\hat{\boldsymbol{\psi}}})$

Test y Regiones de Confianza

Sea $\mathbf{A} \in \mathbb{R}^{k \times p}$ con $\text{rango}(\mathbf{A}) = k \leq p$ y

$$\psi = \mathbf{A}\theta \qquad \hat{\psi} = \mathbf{\hat{A}}\hat{\theta}$$

Test y Regiones de Confianza

Sea $\mathbf{A} \in \mathbb{R}^{k \times p}$ con $\text{rango}(\mathbf{A}) = k \leq p$ y

$$\boldsymbol{\psi} = \mathbf{A}\boldsymbol{\theta} \quad \hat{\boldsymbol{\psi}} = \mathbf{A}\hat{\boldsymbol{\theta}}$$

Del Teorema 4 deducimos

Corolario 1. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\boldsymbol{\theta} \in \mathbb{R}^p$ y vale H3. Entonces,

i) Una región de confianza de nivel $1 - \alpha$ para $\boldsymbol{\psi} = \mathbf{A}\boldsymbol{\theta}$ es

$$\{\boldsymbol{\psi} \in \mathbb{R}^k : (\hat{\boldsymbol{\psi}} - \boldsymbol{\psi})^T \hat{\boldsymbol{\Sigma}}_{\hat{\boldsymbol{\psi}}}^{-1} (\hat{\boldsymbol{\psi}} - \boldsymbol{\psi}) \leq k f_{k, n-p, \alpha}\}$$

Test e Intervalos de Confianza

En particular, si $\mathbf{a} \in \mathbb{R}^p$ y

$$\varphi = \mathbf{a}^T \boldsymbol{\theta} \qquad \hat{\varphi} = \mathbf{a}^T \hat{\boldsymbol{\theta}}$$

Test e Intervalos de Confianza

En particular, si $\mathbf{a} \in \mathbb{R}^p$ y

$$\varphi = \mathbf{a}^T \boldsymbol{\theta} \quad \hat{\varphi} = \mathbf{a}^T \hat{\boldsymbol{\theta}}$$

Del Teorema 3 (o del Corolario 1) deducimos

Corolario 2. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\boldsymbol{\theta} \in \mathbb{R}^p$ y vale H3. Sea $\hat{\sigma}_{\hat{\varphi}}^2 = s^2 \mathbf{a}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{a}$ Entonces,

i) Un intervalo de confianza de nivel $1 - \alpha$ para $\varphi = \mathbf{a}^T \boldsymbol{\theta}$ es

$$[\hat{\varphi} - t_{n-p, \frac{\alpha}{2}} \hat{\sigma}_{\hat{\varphi}}; \hat{\varphi} + t_{n-p, \frac{\alpha}{2}} \hat{\sigma}_{\hat{\varphi}}]$$

Test e Intervalos de Confianza

En particular, si $\mathbf{a} \in \mathbb{R}^p$ y

$$\varphi = \mathbf{a}^T \boldsymbol{\theta} \quad \hat{\varphi} = \mathbf{a}^T \hat{\boldsymbol{\theta}}$$

Del Teorema 3 (o del Corolario 1) deducimos

Corolario 2. Supongamos que $\text{rango}(\mathbf{X}) = p$, $\boldsymbol{\theta} \in \mathbb{R}^p$ y vale H3. Sea $\hat{\sigma}_{\hat{\varphi}}^2 = s^2 \mathbf{a}^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{a}$ Entonces,

i) Un intervalo de confianza de nivel $1 - \alpha$ para $\varphi = \mathbf{a}^T \boldsymbol{\theta}$ es

$$[\hat{\varphi} - t_{n-p, \frac{\alpha}{2}} \hat{\sigma}_{\hat{\varphi}}; \hat{\varphi} + t_{n-p, \frac{\alpha}{2}} \hat{\sigma}_{\hat{\varphi}}]$$

ii) Un test de hipótesis para $H_0 : \varphi = \varphi_0$ versus $H_1 : \varphi \neq \varphi_0$ está dado

$$\phi(\mathbf{y}) = \begin{cases} 1 & \frac{|\hat{\varphi} - \varphi_0|}{\hat{\sigma}_{\hat{\varphi}}} > t_{n-p, \frac{\alpha}{2}} \\ 0 & \frac{|\hat{\varphi} - \varphi_0|}{\hat{\sigma}_{\hat{\varphi}}} \leq t_{n-p, \frac{\alpha}{2}} \end{cases}$$

Test e Intervalos de Confianza para θ_i

Estos resultados nos permiten deducir intervalos de confianza o tests para cada uno de los coeficientes del modelo lineal:

Como $\hat{\boldsymbol{\theta}} \sim N_p(\boldsymbol{\theta}, \sigma^2(\mathbf{X}^T\mathbf{X})^{-1})$, entonces

$$\hat{\theta}_i = \mathbf{e}_i^T \hat{\boldsymbol{\theta}} \sim N(\theta_i, \sigma^2 \mathbf{e}_i^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{e}_i)$$

Si llamamos $\boldsymbol{\Sigma}_{\hat{\boldsymbol{\theta}}} = \sigma^2 \mathbf{D}$

$$\hat{\theta}_i \sim N(\theta_i, \sigma^2 d_{ii})$$

siendo d_{ii} el i -ésimo elemento diagonal de $\mathbf{D}$.

○○○○○
○○○○○○○○○
○○○○○○○○

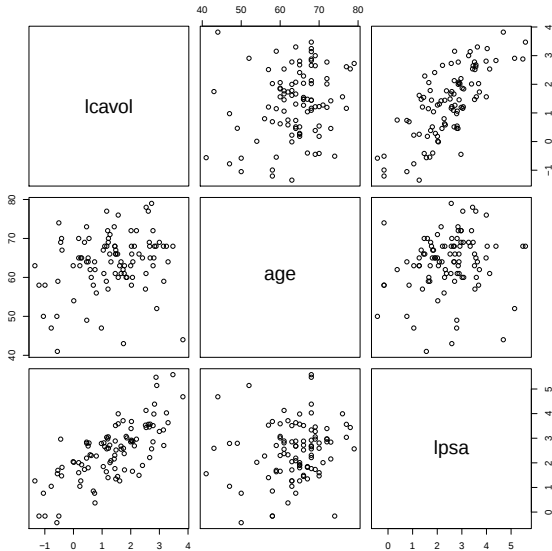
○○○○○○○

○○○○○○○
○○○○○○○○○○○○○○○○○○
○○○○○○○○○○
○○○○●○○○○○○○

○○○○○○○

○○○○
○○○○○

Ejemplo



$$lcavol_i = \theta_0 + \theta_1 age_i + \theta_2 lpsa_i + \varepsilon_i$$

```
salida.2 <- lm(lcavol ~ age + lpsa)
summary(salida.2)
```

$$lcavol_i = \theta_0 + \theta_1 age_i + \theta_2 lpsa_i + \varepsilon_i$$

```
salida.2 <- lm(lcavol ~ age + lpsa)
summary(salida.2)
```

Call:

```
lm(formula = lcavol ~ age + lpsa) Residuals:
```

Min	1Q	Median	3Q	Max
-2.23487	-0.62467	0.02114	0.54422	1.71757

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)	
(Intercept)	-1.50978	0.70670	-2.136	0.0352	*
age	0.01637	0.01112	1.473	0.1442	
lpsa	0.73201	0.07170	10.210	<2e-16	***

Residual standard error: 0.7992 on 94 degrees of freedom

Multiple R-squared: 0.5498, Adjusted R-squared: 0.5402

F-statistic: 57.4 on 2 and 94 DF, p-value: < 2.2e-16

Ejemplo

Vamos a reproducir el cálculo del estadístico F reportado en el summary en el que se testa la significación de la regresión, o sea

$$H_0 : \begin{pmatrix} \theta_1 \\ \theta_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

El vector reducido (es decir sin la primera componente)

$\boldsymbol{\theta}^* = (\theta_1, \theta_2)^T$ es el resultado de multiplicar a $\boldsymbol{\theta} = (\theta_0, \theta_1, \theta_2)^T$ por la matriz

$$\mathbf{A} = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Ejemplo

Usando las propiedades de vectores aleatorios y de la distribución normal tenemos que

$$\hat{\boldsymbol{\theta}}^* = \begin{pmatrix} \hat{\theta}_1 \\ \hat{\theta}_2 \end{pmatrix} \sim N_2 \left(\boldsymbol{\theta}^*, \sigma^2 \mathbf{A} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{A}^T \right)$$

Ejemplo

```
tita.est <- salida.2$coefficients
```

```
ese<- summary(salida.2)$sigma
```

```
XtX = t(model.matrix(salida.2))%*%model.matrix(salida.2)
```

```
XtX.inv = solve(XtX)
```

```
XtX.inv
```

	(Intercept)	age	lpsa
(Intercept)	0.781970977	-0.0118329283	-0.0064316618
age	-0.011832928	0.0001934906	-0.0002116456
lpsa	-0.006431662	-0.0002116456	0.0080490331

Ejemplo

```
tita.est <- salida.2$coefficients
```

```
ese<- summary(salida.2)$sigma
```

```
XtX = t(model.matrix(salida.2))%*%model.matrix(salida.2)
```

```
XtX.inv = solve(XtX)
```

```
XtX.inv
```

	(Intercept)	age	lpsa
(Intercept)	0.781970977	-0.0118329283	-0.0064316618
age	-0.011832928	0.0001934906	-0.0002116456
lpsa	-0.006431662	-0.0002116456	0.0080490331

```
A<-matrix(rep(0,6),ncol = 3,nrow = 2)
```

```
A[1,2]=A[2,3]=1
```

```
A
```

	[, 1]	[, 2]	[, 3]
[1,]	0	1	0
[2,]	0	0	1

Ejemplo

Calculemos $\hat{\theta}^*$

```
theta.star=A%*%tita.est
```

```
theta.star
```

Ejemplo

Calculemos $\hat{\theta}^*$

```
theta.star=A%*%tita.est
```

```
theta.star
```

```
      [,1]
[1,] 0.01637148
[2,] 0.73201132
```

Ejemplo

```
# Calculemos  $\hat{\theta}^*$ 
```

```
theta.star=A%*%tita.est
```

```
theta.star
```

```
      [,1]
[1,] 0.01637148
[2,] 0.73201132
```

```
CovMatj- (ese^2)*A%*%solve(XtX)%*%t(A)
```

```
CovMat
```


Ejemplo

Dado que

$$V_1 = \frac{(\hat{\theta}^* - \theta^*)^T \left(\mathbf{A} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{A}^T \right)^{-1} (\hat{\theta}^* - \theta^*)}{\sigma^2} \sim \chi_2^2$$

Ejemplo

Dado que

$$V_1 = \frac{(\hat{\theta}^* - \theta^*)^T \left(\mathbf{A} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{A}^T \right)^{-1} (\hat{\theta}^* - \theta^*)}{\sigma^2} \sim \chi_2^2$$

es independiente de

$$V_2 = \frac{(n-3)s^2}{\sigma^2} \sim \chi_{n-3}^2$$

Ejemplo

Dado que

$$V_1 = \frac{(\hat{\theta}^* - \theta^*)^T \left(\mathbf{A} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{A}^T \right)^{-1} (\hat{\theta}^* - \theta^*)}{\sigma^2} \sim \chi_2^2$$

es independiente de

$$V_2 = \frac{(n-3)s^2}{\sigma^2} \sim \chi_{n-3}^2$$

resulta que el cociente $(V_1/2)/(V_2/(n-3))$ tiene distribución $F_{2,n-3}$

$$\frac{(\hat{\theta}^* - \theta^*)^T \left(\mathbf{A} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{A}^T \right)^{-1} (\hat{\theta}^* - \theta^*)}{2 s^2} \sim F_{2,n-3}$$

Ejemplo

```
t(theta.star)%*%(solve(CovMat))%*%(theta.star)/2
      [,1]
[1,] 57.40246
```


80

Funciones Lineales de los Parámetros

En muchas aplicaciones estamos más interesados en estimar funciones lineales de θ que en estimar θ en sí mismo:

$$a^T \theta \quad \text{o} \quad A\theta$$

Un caso de especial interés:

En un futuro valor $\mathbf{x}_0$, ¿qué esperamos observar? $\mathbb{E}(y_0) = ?$

$$\mathbb{E}(y_0) = \mathbf{x}_0^T \theta \Rightarrow \widehat{\mathbb{E}(y_0)} = \mathbf{x}_0^T \hat{\theta}$$

Observemos que si $\hat{\theta}$ es insesgado entonces el estimador propuesto sería insesgado de $\mathbb{E}(y_0)$.

80

Funciones Estimables

Definición 1. Una **función paramétrica** φ se dice que es una **función lineal** de los parámetros $\{\theta_1, \dots, \theta_p\}$ si existen $\{a_1, \dots, a_p\}$ constantes conocidas tal que

$$\varphi = \mathbf{a}^T \boldsymbol{\theta} = \sum_{j=1}^p a_j \theta_j$$

donde $\mathbf{a} = (a_1, \dots, a_p)^T$.

Funciones Estimables

Definición 1. Una **función paramétrica** φ se dice que es una **función lineal** de los parámetros $\{\theta_1, \dots, \theta_p\}$ si existen $\{a_1, \dots, a_p\}$ constantes conocidas tal que

$$\varphi = \mathbf{a}^T \boldsymbol{\theta} = \sum_{j=1}^p a_j \theta_j$$

donde $\mathbf{a} = (a_1, \dots, a_p)^T$.

Definición 2. Decimos que una función paramétrica $\varphi = \mathbf{a}^T \boldsymbol{\theta}$ es **estimable** si tiene un estimador lineal (en $\mathbf{y}$) insesgado, es decir si existe $\mathbf{c} \in \mathbb{R}^n$ tal que

$$\mathbb{E}(\mathbf{c}^T \mathbf{y}) = \varphi = \mathbf{a}^T \boldsymbol{\theta} \quad \forall \boldsymbol{\theta} \in \mathbb{R}^p$$

Funciones Estimables: Caracterización

El siguiente resultado caracteriza las funciones paramétricas estimables suponiendo el modelo

$$\Omega : \mathbb{E}(\mathbf{y}) = \mathbf{X}\boldsymbol{\theta} \quad \boldsymbol{\Sigma}_y = \sigma^2 \mathbf{I}$$

Teorema 5. La función paramétrica $\varphi = \mathbf{a}^T \boldsymbol{\theta}$ es estimable si y sólo si $\mathbf{a}$ es una combinación lineal de las filas de $\mathbf{X}$, o sea si existe $\mathbf{c} \in \mathbb{R}^n$ tal que

$$\mathbf{a} = \mathbf{X}^T \mathbf{c}$$

Funciones Estimables: Caracterización

El siguiente resultado caracteriza las funciones paramétricas estimables suponiendo el modelo

$$\Omega : \mathbb{E}(\mathbf{y}) = \mathbf{X}\theta \quad \Sigma_{\mathbf{y}} = \sigma^2 \mathbf{I}$$

Teorema 5. La función paramétrica $\varphi = \mathbf{a}^T \boldsymbol{\theta}$ es estimable si y sólo si $\mathbf{a}$ es una combinación lineal de las filas de $\mathbf{X}$, o sea si existe $\mathbf{c} \in \mathbb{R}^n$ tal que

$$\mathbf{a} = \mathbf{X}^T \mathbf{c}$$

Notemos que

- Si $\text{rango}(\mathbf{X}) = p$, toda función $\varphi = \mathbf{a}^T \boldsymbol{\theta}$ es estimable.

Ejemplo

Supongamos que

$$y_{i,j} = \mu + \alpha_j + \epsilon_{ij} \quad 1 \leq j \leq 2, \quad 1 \leq i \leq n_j$$

entonces

$$p = 3 \quad \mathbf{y} = \begin{pmatrix} y_{1,1} \\ y_{2,1} \\ \vdots \\ y_{n_1,1} \\ y_{1,2} \\ y_{2,2} \\ \vdots \\ y_{n_2,2} \end{pmatrix} \quad \mathbf{X} = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 1 & 0 \\ \vdots & \vdots & \vdots \\ 1 & 1 & 0 \\ 1 & 0 & 1 \\ 1 & 0 & 1 \\ \vdots & \vdots & \vdots \\ 1 & 0 & 1 \end{pmatrix} \quad \boldsymbol{\theta} = \begin{pmatrix} \mu \\ \alpha_1 \\ \alpha_2 \end{pmatrix}$$

Ejemplo

Puedo estimar las combinaciones lineales de las componentes de

$$\eta = \begin{pmatrix} \mu + \alpha_1 \\ \vdots \\ \mu + \alpha_1 \\ \mu + \alpha_2 \\ \vdots \\ \mu + \alpha_2 \end{pmatrix}$$

es decir, puedo estimar $\mu + \alpha_1$ y $\alpha_2 - \alpha_1$, pero no puedo estimar ni μ ni α_1 .

Mejor Estimador Lineal Insesgado (BLUE)

Lema 3. Supongamos que vale el modelo Ω . Sean

- $\varphi = \mathbf{a}^T \boldsymbol{\theta}$ una función estimable
- $\mathcal{V}$ el espacio generado por las columnas de $\mathbf{X}$ ($q = \text{rango}(\mathbf{X}) \leq p$).

Luego,

- existe un único $\mathbf{c}_0 \in \mathcal{V}$ tal que el estimador lineal

$$\hat{\varphi} = \mathbf{c}_0^T \mathbf{y}$$

es un estimador insesgado de φ .

Mejor Estimador Lineal Insesgado (BLUE)

Lema 3. Supongamos que vale el modelo Ω . Sean

- $\varphi = \mathbf{a}^T \boldsymbol{\theta}$ una función estimable
- $\mathcal{V}$ el espacio generado por las columnas de $\mathbf{X}$ ($q = \text{rango}(\mathbf{X}) \leq p$).

Luego,

- existe un único $\mathbf{c}_0 \in \mathcal{V}$ tal que el estimador lineal

$$\hat{\varphi} = \mathbf{c}_0^T \mathbf{y}$$

es un estimador insesgado de φ .

- Más aún, si $\mathbf{c}^T \mathbf{y}$ es un estimador insesgado de φ , $\mathbf{c}_0$ es la proyección ortogonal de $\mathbf{c}$ sobre $\mathcal{V}$ y

$$\text{VAR}(\mathbf{c}_0^T \mathbf{y}) \leq \text{VAR}(\mathbf{c}^T \mathbf{y})$$

y vale el igual sólo si $\mathbf{c} = \mathbf{c}_0$.

Notación

El modelo se escribe como

$$\Omega: \quad \mathbf{y} = \mathbf{X}\boldsymbol{\theta} + \epsilon$$

$$\mathbb{E}(\epsilon) = 0$$

$$\mathbf{\Sigma}_{\epsilon} = \sigma^2 \mathbf{I}_n$$

Notación

El modelo se escribe como

$$\Omega : \mathbf{y} = \mathbf{X}\boldsymbol{\theta} + \boldsymbol{\epsilon}$$

$$\mathbb{E}(\boldsymbol{\epsilon}) = \mathbf{0}$$

$$\boldsymbol{\Sigma}_{\boldsymbol{\epsilon}} = \sigma^2 \mathbf{I}_n$$

$$\text{rango}(\mathbf{X}) = p$$

Notación

El modelo se escribe como

$$\Omega : \quad \mathbf{y} = \mathbf{X}\boldsymbol{\theta} + \boldsymbol{\epsilon}$$

$$\mathbb{E}(\boldsymbol{\epsilon}) = \mathbf{0}$$

$$\boldsymbol{\Sigma}_{\boldsymbol{\epsilon}} = \sigma^2 \mathbf{I}_n$$

$$\text{rango}(\mathbf{X}) = p$$

$$\hat{\boldsymbol{\theta}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

Notación

El modelo se escribe como

$$\begin{aligned}\Omega : \quad \mathbf{y} &= \mathbf{X}\boldsymbol{\theta} + \boldsymbol{\epsilon} \\ \mathbb{E}(\boldsymbol{\epsilon}) &= \mathbf{0} \\ \boldsymbol{\Sigma}_{\boldsymbol{\epsilon}} &= \sigma^2 \mathbf{I}_n\end{aligned}$$

$$\text{rango}(\mathbf{X}) = p$$

$$\hat{\boldsymbol{\theta}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

Dada una matriz $\mathbf{A} \in \mathbb{R}^{p \times p}$ llamaremos

$$\lambda_1(\mathbf{A}) \leq \lambda_2(\mathbf{A}) \leq \cdots \leq \lambda_p(\mathbf{A})$$

a los autovalores de $\mathbf{A}$

Consistencia

a) Si $\lambda_1(\mathbf{X}^T\mathbf{X}) \rightarrow \infty$ entonces $\hat{\theta} \xrightarrow{p} \theta$.

b) Supongamos que se cumple H3, entonces la condición $\lambda_1(\mathbf{X}^T\mathbf{X}) \rightarrow \infty$ también es necesaria.

Ejemplos

a) Regresión por el origen: $y_i = \theta_1 x_i + \epsilon_i \quad 1 \leq i \leq n$

$$\mathbf{X} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \quad \mathbf{X}^T \mathbf{X} = \sum_{i=1}^n x_i^2$$

La condición $\lambda_1(\mathbf{X}^T \mathbf{X}) \rightarrow \infty$ dice que $\sum_{i=1}^n x_i^2 \rightarrow \infty$

Ejemplos

b) Regresión simple: $y_i = \theta_1 + \theta_2 x_i + \epsilon_i \quad 1 \leq i \leq n$

$$\mathbf{X}^T \mathbf{X} = \begin{pmatrix} n & \sum_{i=1}^n x_i \\ \sum_{i=1}^n x_i & \sum_{i=1}^n x_i^2 \end{pmatrix}$$

Ejemplos

b) Regresión simple: $y_i = \theta_1 + \theta_2 x_i + \epsilon_i \quad 1 \leq i \leq n$

$$\mathbf{X}^T \mathbf{X} = \begin{pmatrix} n & \sum_{i=1}^n x_i \\ \sum_{i=1}^n x_i & \sum_{i=1}^n x_i^2 \end{pmatrix}$$

$$(\mathbf{X}^T \mathbf{X})^{-1} = \frac{1}{\sum_{i=1}^n (x_i - \bar{x})^2} \begin{pmatrix} \frac{1}{n} \sum_{i=1}^n x_i^2 & -\frac{1}{n} \sum_{i=1}^n x_i \\ -\frac{1}{n} \sum_{i=1}^n x_i & 1 \end{pmatrix}$$

La condición $\lambda_1(\mathbf{X}^T \mathbf{X}) \rightarrow \infty$ es equivalente a $\sum_{i=1}^n (x_i - \bar{x})^2 \rightarrow 0$.

Distribución asintótica

Bajo H3,

$$\hat{\theta} \sim N(\theta, \sigma^2(\mathbf{X}^T \mathbf{X})^{-1})$$

Luego si $\mathbf{S} \in \mathbb{R}^{p \times p}$ es tal que

$$\mathbf{S}\mathbf{S}^T = \mathbf{X}^T\mathbf{X}$$

Tenemos que

$$\mathbf{S}^T(\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}) \sim N(\mathbf{0}, \sigma^2 \mathbf{I}_p)$$

Distribución asintótica

$$\mathbf{z}_i = \mathbf{S}^{-1} \mathbf{x}_i \quad \boldsymbol{\gamma} = \mathbf{S}^T \boldsymbol{\theta}$$

El modelo se reescribe

$$y_i = \mathbf{x}_i^T \boldsymbol{\theta} + \epsilon_i = (\mathbf{S} \mathbf{z}_i)^T (\mathbf{S}^T)^{-1} \boldsymbol{\gamma} + \epsilon_i = \mathbf{z}_i^T \boldsymbol{\gamma} + \epsilon_i$$

El EMC de $\boldsymbol{\gamma}$ es

$$\hat{\boldsymbol{\gamma}} = \mathbf{S}^T \hat{\boldsymbol{\theta}}$$

Basta ver que

$$\hat{\boldsymbol{\gamma}} - \boldsymbol{\gamma} \xrightarrow{D} N(\mathbf{0}, \sigma^2 \mathbf{I}_p)$$

Distribución asintótica

Propiedad. Bajo H_1 y H_2 ,

$$\frac{\hat{\gamma} - \gamma}{p} = O_{\mathbb{P}}(1)$$

Lema. Sean $\{V_j\}_{j \geq 1}$ i.i.d $\mathbb{E}(V_j) = 0$, $\text{VAR}(V_j) = \sigma^2$. Definamos

$$W_n = \sum_{j=1}^n a_{j,n} V_j \text{ donde } \sum_{j=1}^n a_{j,n}^2 = 1$$

Si

$$\max_{1 \leq j \leq n} a_{j,n}^2 \rightarrow 0$$

entonces $W_n \xrightarrow{D} N(0, \sigma^2)$.

Observación

Observemos que

$$\mathbf{x}_j^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{x}_j$$

es el j —ésimo elemento diagonal de la matriz de proyección

$$\mathbf{P} = \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$$