

Apuntes Econometría Teoría_T1_al_T6 - EL CORÁN

(Ejercicios propuestos y resueltos en clase en otro archivo)

Tema 1: Introducción a la econometría

¿Qué es el método econométrico?

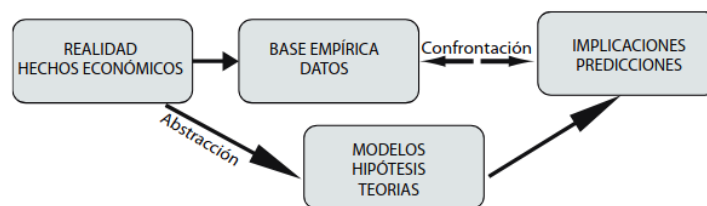
El método econométrico está construido sobre un proceso de razonamiento lógico utilizando una información inicial y un conocimiento empírico. Compuesto y estructurado por:

Realidad: Estudio de los hechos económicos que surgen de la misma sociedad.

Datos: Recopilación de los datos obtenidos del conocimiento que el investigador tiene sobre la propia realidad.

Modelo: A partir de las bases de datos, el investigador formula hipótesis dando lugar a modelos, leyes y teorías.

Predicciones: Los modelos permiten al investigador predecir determinados comportamientos.



¿Qué es la econometría?

La Econometría es la ciencia que utiliza la Teoría Económica, las Matemáticas y la Inferencia Estadística de forma conjunta utilizando como herramienta la Informática.

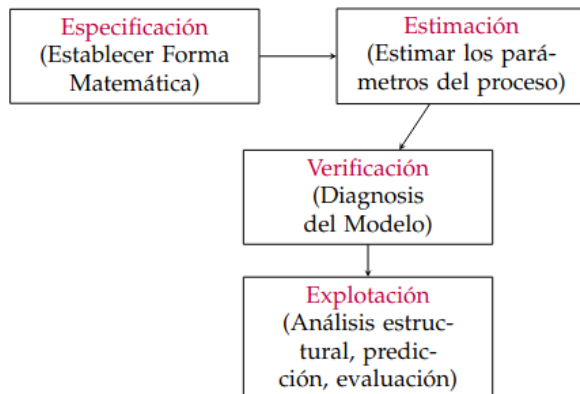
¿Qué son los Modelos Económicos y los Modelos Econométricos?

Los modelos económicos son representaciones algebraicas simplificadas de la realidad económica, usando ecuaciones que relacionan diferentes variables.

Los modelos econométricos son una aplicación práctica de estos modelos (económicos), en los que se describe matemáticamente la relación entre variables explicadas y explicativas, incluyendo un término aleatorio o perturbación.

Fases del Método Econométrico

1. **Especificación:** Se establece matemáticamente la relación entre las variables.
2. **Estimación:** Se calculan los parámetros del modelo. (Estimar los parámetros)
3. **Verificación:** Se comprueba el modelo para detectar errores y, si es necesario, reformularlo (es decir, si no cumple la verificación, volveremos al paso 1).
4. **Explotación:** El modelo, una vez validado, se usa para análisis, predicción y evaluación de políticas económicas.



(Este esquema se corresponde a las fases del método econométrico)

VARIABLES

En un modelo econométrico podemos diferenciar las variables en dos grandes grupos:

No observables: Son aquellas de las que NO disponemos de los datos muestrales. Y por ello, tienen algo llamado **perturbación (u_t)**.

Observables: Son aquellas de las que disponemos los datos muestrales:

Las **variables endógenas corrientes** son las **variables explicadas** o regresandos, que son determinadas/explicadas por el modelo econométrico. Las **variables predeterminadas o explicativas**, también llamadas regresores, describen el comportamiento de las variables endógenas/explicadas. Estas se subdividen en:

Corriente: Momento actual // Retardada: Momento anterior

- **Endógenas retardadas.**
- **Exógenas corrientes.**
- **Exógenas retardadas.**

$$\begin{aligned}
 c_t &= \alpha_0 + \alpha_1(1 - \tau)y_t + \alpha_2 r_t + u_t, \\
 i_t &= \beta_0 + \beta_1(y_{t-1} - y_{t-2}) + \beta_2 r_{t-1} + v_t, \\
 y_t &\equiv c_t + i_t + g_t
 \end{aligned}$$

Handwritten notes: "Explicado" points to c_t and i_t . "explicativa" points to y_t . "ex pl cativa retardada" points to r_{t-1} .

Donde:

C_t, i_t, Y_t : Son variables endógenas corrientes. (Las incógnitas, aquellas a la izquierda) // Con su variación con el subíndice-1 que en dicho caso serían endógenas retardadas.

r_t, g_t : Son variables exógenas corrientes. // Con su variación con el subíndice-1 que en dicho caso serían exógenas retardadas.

Todas aquellas con subíndices-1 son variables retardadas.

Además tenemos que saber que: **Los parámetros** (fijos y desconocidos) son los betas y las alphas del ejemplo anterior.

Las **relaciones** entre los distintos tipos de variables vienen especificadas mediante un sistema de ecuaciones.

FORMAS DE EXPRESAR UN MODELO ECONOMETRICO

Según el NÚMERO DE ECUACIONES:

- **Uniecuacionales:** Sistema descrito por una ecuación
 1. Uniecuacional simple $Y_t = \beta_0 + \beta_1 X_t + u_t$.
 2. Uniecuacional múltiple $Y_t = \beta_0 + \beta_1 X_{1t} + \beta_2 X_{2t} + \dots + \beta_k X_{kt} + u_t$.
- **Multiecuacionales:** Formado por un sistema de ecuaciones donde puede haber o no relación entre las variables.

$$Y_{1t} = \gamma_{12} Y_{2t} + \beta_{11} X_{1t} + u_{1t}$$

$$Y_{2t} = \gamma_{21} Y_{1t} + \beta_{21} X_{1t} + \beta_{22} X_{2t} + u_{2t}$$

$$Y_{1t} = \beta_0 + \beta_1 X_{1t} + u_t$$

$$Y_{2t} = \alpha_0 + \alpha_1 X_{2t} + v_t$$

Según la FORMA FUNCIONAL:

Modelos lineales $Y_t = \beta_0 + \beta_1 X_{1t} + \beta_2 X_{2t} + u_t$.

Modelos no lineales intrínsecamente linealizables

$$Y_t = \beta_0 + \beta_1 e^{X_{1t}} + \beta_2 X_{2t} X_{3t} + u_t$$

Modelos no lineales intrínsecamente no linealizables $Y_t = \beta_0 + X_t^{\beta_1} + u_t$.

Según las CARACTERÍSTICAS DE LA PERTURBACIÓN:

- **Modelos clásicos:** Cuando la perturbación es “ruido blanco” es decir:

1. $E[u_t] = 0 \forall t$
2. $\text{var}[u_t] = \sigma^2 \forall t$
3. $\text{cov}[u_t u_s] = 0 \forall t \neq s$.

Cumplen dichas hipótesis.

- **Modelos generalizados:** Cuando las perturbaciones incumplen alguna de las DOS últimas hipótesis descritas en los modelos clásicos. Es decir, habrá un problema de **heterocedasticidad** (varianza no cte.) o **autocorrelación** (covarianza distinta de 0 para t distinto de s). Siendo t y s dos momentos diferentes.

NATURALEZA DE LA INFORMACIÓN

Las bases de datos en econometría se dividen en tres tipos:

1. **Sección Cruzada / Cortes transversal:** Datos de múltiples individuos observados en el mismo instante de tiempo. Ejemplo: Calificaciones de alumnos en un examen específico.
2. **Series Temporales:** Datos de un solo individuo observados en distintos momentos del tiempo. Ejemplo: Cotizaciones diarias de las acciones de un banco.
3. **Datos de Panel:** Datos de varios individuos observados en diferentes momentos del tiempo. Ejemplo: Cotizaciones diarias de todas las empresas de un índice bursátil.

EJEMPLOS:

Ejemplo de corte transversal:

N. de obs.	Calif. media /dist(Y)	Ratio estud-maest(X ₁)	C/estud(X ₂)	% estud. inglés(X ₃)
1	660.8	17.89	6.385	0.8
2	661.2	21.52	5.099	4.6
3	643.6	18.70	5.502	30.8
...	...	...	...	...
419	692.2	20.20	4.776	3.0
420	765.8	19.04	5.993	5.8

Ejemplo de series temporales:

Año	Consumo (Y)	Renta (X)
1982	3981.5	6203.9
1983	3246.6	4903.7
1984	3407.6	5140.1
1985	3596.5	5323.5
1986	3708.7	5487.7
1987	3622.3	5448.5
1988	3672.7	5665.2
1989	4064.6	6062.0
1990	4132.2	6136.3
1991	4305.8	6279.4
1992	4219.5	6244.4
1993	4343.6	6389.6
1994	4486.0	6632.7
1995	4595.3	6742.1
1996	4714.1	6826.4

Ejemplo de datos de panel:

N. de obs.	Estado	Año	Vta de cig (Y)	P medio/paquete(X ₁)	Implos tot(X ₂)
1	Alabama	1985	116.5	1.022	0.333
2	Arkansas	1985	128.5	1.015	0.370
...	...	...	...	...	...
48	Wyoming	1985	129.4	0.935	0.240
49	Alabama	1986	117.2	1.080	0.334
...	...	...	...	...	...
96	Wyoming	1986	127.8	1.007	0.240
97	Alabama	1987	115.8	1.135	0.335
...	...	...	...	...	...
528	Wyoming	1995	127.8	1.007	0.240

Autocorrelación

Tema 2: Modelo Lineal (I)

Recordando del tema anterior:

Modelo Lineal Simple:

$$y_i \sim a + bx_i$$

Modelo Lineal Múltiple:

Permite estudiar el comportamiento de **una** variable endógena Y , a partir de k variables independientes ($X_2, X_3, \dots, X_k$)

$$Y_i = \beta_1 + \beta_2 X_{2i} + \beta_3 X_{3i} + \dots + \beta_k X_{ki} + u_i$$

Notación MLP

$$Y_i = \beta_1 + \beta_2 X_{2i} + \beta_3 X_{3i} + \dots + \beta_k X_{ki} + u_i$$

- ▶ el subíndice $i = 1, \dots, n$, indica la i -ésima de las n observaciones de la muestra.
- ▶ El subíndice i será empleado cuando trabajemos con observaciones de sección cruzada o de corte transversal, en el caso de trabajar con información de carácter temporal se empleará el subíndice t .
- ▶ Y_i es la i -ésima observación de la variable dependiente.
- ▶ $X_{2i}, \dots, X_{ki}$ son las i -ésimas observaciones de las k variables independientes.
- ▶ u_i es la perturbación de la i -ésima observación, encargada de recoger aquella parte de la variable Y_i que no es explicada por $X_{1i}, X_{2i}, \dots, X_{ki}$.
- ▶ $\beta_1, \dots, \beta_k$ son los parámetros del modelo. Además, β_1 es el denominado término constante.

Así que realmente estaremos trabajando con matrices:

$$\vec{y} = X\vec{\beta} + \vec{u}$$

$$\vec{y} = \begin{pmatrix} Y_1 \\ Y_2 \\ \vdots \\ Y_n \end{pmatrix} \quad X = \begin{pmatrix} 1 & X_{21} & \dots & X_{k1} \\ 1 & X_{22} & \dots & X_{k2} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & X_{2n} & \dots & X_{kn} \end{pmatrix} \quad \vec{\beta} = \begin{pmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_k \end{pmatrix} \quad \vec{u} = \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix}$$

SUPUESTOS

Para el análisis estadístico del modelo, comenzaremos considerando las siguientes

hipótesis básicas:

- 1. Supuesto de Linealidad:** Existe una relación lineal entre la variable dependiente y las variables explicativas, lo que permite usar álgebra de matrices y vectores para simplificar el análisis. *Es decir, podemos escribir la relación muestral tipo así:*

$$Y_i = \beta_1 + \beta_2 X_{2i} + \beta_3 X_{3i} + \dots + \beta_k X_{ki} + u_i$$

$$\text{con } i = 1, \dots, n \text{ o equivalentemente } Y_i = \vec{X}_i^t \vec{\beta} + u_i.$$

(*Aquí está la definición de Y_i ;) *)

2. Supuesto de Rango Completo por Columnas: La matriz de observaciones debe tener rango completo por columnas, lo que significa que las variables explicativas no deben ser combinaciones lineales entre sí, evitando problemas de *multicolinealidad*. Es decir, las columnas “k” de la matriz X son linealmente independientes. Esto nos permite además calcular la inversa de la matriz de productos cruzados: $(X^t X)^{-1}$

3. Supuesto de Exogeneidad: Las variables explicativas no deben tener ninguna relación con el término de perturbación. Es decir, el valor esperado de la perturbación debe ser cero, lo que garantiza que las variables explicativas no predicen el término de error. Es decir, se cumple esto (*Recuerda la definición de Y_i):

$$E[u_i | (X_{1i}, X_{2i}, \dots, X_{ki})] = 0 \quad E[Y_i | (X_{1i}, X_{2i}, \dots, X_{ki})] = \vec{X}_i^t \vec{\beta}.$$

4. Supuesto de Causalidad: La relación entre las variables explicativas y la variable dependiente es unidireccional; las variables explicativas causan cambios en la variable dependiente, pero no al revés. (Esto ya se enunció en el T1, las de la derecha explican la de la izquierda, y esta última depende de la perturbación)

5. Supuesto sobre la Perturbación: El término de perturbación debe ser: (Ya visto en T1)

- **Centrado:** Su valor esperado es cero. $E[u_i] = 0 \quad \forall i = 1, \dots, n.$
- **Homocedástico:** Su varianza es constante para todas las observaciones.

$$\text{var}[u_i] = E[u_i^2] = \sigma^2 \quad \text{Si no se cumple habrá heterocedasticidad.}$$

- **Incorrelado:** Las perturbaciones en diferentes observaciones no están correlacionadas. $[u_i u_j] = E[u_i u_j] = 0$

Si no se cumple habrá autocorrelación.

Consecuentemente, se tiene que la matriz de varianzas-covarianzas del término

perturbación es: $\text{var}[\vec{u}] = E[(\vec{u} - E[\vec{u}])(\vec{u} - E[\vec{u}])^t] = E[\vec{u} \vec{u}^t]$ siendo este el procedimiento que utilizaremos para el cálculo de la varianza en clase (saberse).

6. Supuesto de Normalidad de la Perturbación: Se asume que las perturbaciones siguen una distribución normal con media cero y varianza constante, lo que permite hacer inferencias sobre los parámetros del modelo.

$$u_i \sim N(0, \sigma^2) \quad \forall i = 1, \dots, n \quad \vec{u} \sim N(\vec{0}, \sigma^2 I_n)$$

Con la función de densidad asociada a la distribución (extra):

$$f(\vec{u} | \sigma^2) = \frac{1}{(\sqrt{2\pi}\sigma^2)^n} \exp\left[-\frac{1}{2\sigma^2} \vec{u}^t \vec{u}\right]$$

ESTIMACIÓN: MÍNIMOS CUADRADOS ORDINARIOS (MCO)

Definición: Es una técnica utilizada para estimar los parámetros en modelos de regresión lineal, buscando minimizar la diferencia entre los valores observados y los valores predichos por el modelo.

Para esto denotaremos dentro del **modelo lineal múltiple clásico**: $\vec{y} = X\vec{\beta} + \vec{u}$
Los siguiente valores:

Denotando por $\vec{\hat{y}}$ el valor **ajustado** de $\vec{y}$ por el modelo: $\vec{\hat{y}} = X\vec{\hat{\beta}}$

[Residuo] Diferencia existente entre el valor observado de la variable explicada y su estimación: $\vec{e} = \vec{y} - \vec{\hat{y}} = \vec{y} - X\vec{\hat{\beta}}$

Denotaremos a la función a minimizar por: $f(\vec{\beta}) = \vec{e}^t \vec{e}$.

Una vez realizado las siguientes definiciones, comenzaremos con los pasos a seguir:

1. **Definimos la expresión a minimizar**, definiendo la suma de los cuadrados de los residuos como:

$$\begin{aligned} f(\vec{\beta}) &= (\vec{y} - X\vec{\beta})^t (\vec{y} - X\vec{\beta}) \\ &= \vec{y}^t \vec{y} - \vec{y}^t X\vec{\beta} - \vec{\beta}^t X^t \vec{y} + \vec{\beta}^t X^t X \vec{\beta} \\ &= \vec{y}^t \vec{y} - 2\vec{\beta}^t X^t \vec{y} + \vec{\beta}^t X^t X \vec{\beta} \end{aligned}$$

2. **Derivamos con respecto al vector de beta estimada e igualamos a 0:**

$$\nabla f(\vec{\beta}) = -2X^t \vec{y} + 2X^t X \vec{\beta} = \vec{0}$$

3. **Despejamos el vector de parámetros y calculamos la inversa:**

$$-2X^t \vec{y} + 2X^t X \vec{\beta} = \vec{0} \Leftrightarrow X^t X \vec{\beta} = X^t \vec{y}$$

Considerando el supuesto de rango completo por columnas, sabemos que se verifica que $\rho(X^t X) = k$, condición necesaria para la existencia de la matriz $(X^t X)^{-1}$, luego:

$$\vec{\hat{\beta}} = (X^t X)^{-1} (X^t \vec{y})$$

4. Para que la solución se ha de cumplir:

$$Hess(f) = 2X^t X \succ 0 \quad (A \succ 0 \text{ sii } z^t A z > 0)$$

$$\vec{\hat{\beta}} = (X^t X)^{-1} (X^t \vec{y})$$

es el **Estimador Mínimo Cuadrático Ordinario (EMCO)** del vector de parámetros $\vec{\beta}$ del modelo de regresión lineal múltiple $\vec{y} = X\vec{\beta} + \vec{u}$.

Si el modelo tiene término independiente, este dentro de su matriz en la posición 0,0 tendrá el término n, en caso contrario no. (Si no se entiende consultar diapo. nº18 T2)

PROPIEDADES ALGEBRAICAS DE LOS MCO

- ▶ La suma de los residuos mínimo cuadráticos es cero: $\sum_{i=1}^n e_i = 0$.
- ▶ La suma de los valores observados coincide con la suma de los valores estimados: $\sum_{i=1}^n Y_i = \sum_{i=1}^n \hat{Y}_i$.

Además debemos saber que:

- **SCT**: Cuánto varía la observada de la media. (*Suma de los cuadrados totales*)
- **SCR**: Cuánto varía la observada estimada. (*Suma de los cuadrados de los residuos*)
- **SCE**: Cuánto varía la estimada respecto a la media. (*Suma de cuadrados explicada*)
- **SCT = SCR + SCE**

$$SCT = \sum_{i=1}^n (y_i - \bar{y})^2 = \vec{y}'\vec{y} - n \cdot \bar{y}^2$$

$$SCR = \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \vec{y}'\vec{y} - \vec{\beta}'X'\vec{y} \quad \text{▶ Se cumple que } \vec{y}'\vec{e} = 0.$$

$$SCE = \sum_{i=1}^n (\hat{y}_i - \bar{y})^2 = \vec{\beta}'X'\vec{y} - n \cdot \bar{y}^2$$

Considerando que **SI** hay término independiente en el modelo ($X_{1i} = 1$), se tiene:

- ▶ Las variables exógenas son ortogonales al vector de los residuos:

$$X'\vec{e} = \vec{0}.$$

La condición necesaria de EMCO era:

$$-2X'\vec{y} + 2X'X\vec{\beta} = \vec{0}$$

Equivalentemente:

$$X'(\vec{y} - X\vec{\beta}) = \vec{0}$$

Como $\vec{y} = X\vec{\beta}$, tenemos que: $X'(\vec{y} - X\vec{\beta}) = \vec{0}$, ó equivalentemente

$$X'\vec{e} = \vec{0}.$$

- A continuación en las diapositivas del tema, daba el ejerr Ecuaciones Normales (estarán en el archivo de ejercicios importancia)

BONDAD DEL AJUSTE

[Coeficiente de determinación ordinario / R^2] Porcentaje de variabilidad explicada por el modelo.

$$R^2 = \frac{\frac{1}{n} \cdot \sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2}{\frac{1}{n} \cdot \sum_{i=1}^n (Y_i - \bar{Y})^2} = \frac{\sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2}{\sum_{i=1}^n (Y_i - \bar{Y})^2} = \frac{SCE}{SCT}.$$

Si el modelo tiene término independiente: $R^2 = 1 - \frac{SCR}{SCT}$ y los valores del coeficiente estarán comprendidos entre 0 y 1.

Además de clase sabemos que: $SCT=SCR+SCE$ y por lo tanto:

- Si $SCR=0 \rightarrow SCT=SCE$ y $R^2=1$
- Si $SCE=0 \rightarrow SCT=SCR$ y $R^2=0$

Suele decirse que: Series temporales(R^2 aprox.=0,85) y Sección cruzada(R^2 aprox.=0,6-0,7)

[Coeficiente de determinación corregido / $\bar{R}^2$] Porcentaje de variación de la variable explicada considerando el número de variables incluidas en el modelo, es decir, considerando el valor de k.

$$\bar{R}^2 = 1 - \frac{SCR/(n-k)}{SCT/(n-1)} = 1 - (1 - R^2) \cdot \frac{n-1}{n-k}.$$

Adviértase que obtener un R^2 o $\bar{R}^2$ cercano a 1 no indica que los resultados sean fiables, ya que, por ejemplo, puede ser que no se cumpla alguna de las hipótesis básicas y los resultados no ser válidos.

Es necesario señalar que ambos coeficientes no son capaces de detectar si alguna de las variables incluidas en el modelo son o no estadísticamente significativas.

Por tanto, estos indicadores han de ser considerados como una herramienta más a tener en cuenta dentro del análisis.

CRITERIOS DE AKAIKE Y SCHWARZ (k parámetros, n observaciones)

Para comparar distintos modelos podríamos utilizar el coeficiente de determinación. Sin embargo, para poder llevar cabo dicha comparación será necesario que tengan la misma variable explicada.

Con el fin de buscar una solución al problema utilizaremos los criterios de selección de modelos de Akaike (AIC) y el bayesiano de Schwarz (BIC). Estos criterios se obtienen a partir de la suma de cuadrados de los residuos y de un factor que penaliza la inclusión de parámetros.

El criterio de información de Akaike:

$$AIC = \ln \left(\frac{SCR}{n} \right) + \frac{2k}{n}$$

El criterio de información de Schwarz:

$$BIC = \ln \left(\frac{SCR}{n} \right) + \frac{k}{n} \cdot \ln(n)$$

Utilizando estos criterios se escogería aquel modelo con un menor valor de AIC o BIC.

26 / 32

Acompañado de un ejemplo que hicimos en clase:

Ejemplo de uso

Ingresos: $(\beta_1) + (\beta_2) n^{\circ} \text{empleados} + \mu_i$ $n=15$ $k=2$

$$SCT = 590.247.358 \quad d' R^2? \quad SCT = SCR + SCE$$

$$SCR = 9.425.584'39 \quad d' \bar{R}^2?$$

$$\bar{R}^2 = 1 - \frac{SCR/(15-2)}{SCT/(15-1)} = 0'9828$$

$$R^2 = \frac{SCE}{SCT} = \frac{SCT - SCR}{SCT} = 0'984$$

$$AIC = \ln \left(\frac{SCR}{n} \right) + \frac{2k}{n} = \ln \left(\frac{...}{15} \right) + \frac{2 \cdot 2}{15} = 13'61$$

(Akaike) $\rightarrow + \text{peq} = \text{mejor}$

Ingresos: $(\beta_1) + (\beta_2) n^{\circ} \text{empleados} + (\beta_3) n^{\circ} \text{tiendas} + \mu_i$ $n=15$ $k=3$

$$R^2 = 0'9932$$

$$\bar{R}^2 = 0'9921$$

$$SCT = 590.247.358$$

$$AIC = \ln \left(\frac{SCR}{n} \right) + \frac{2k}{n} = 12'88 \leftarrow \text{este es mejor, } + \text{peque.}$$

El resto de cosas apuntadas son ejercicios que pondré en el archivo correspondiente.

Tema 3: Modelo Lineal (II)

Sabemos del tema anterior que:

Donde u sería la variable aleatoria con tal media y varianza. En caso de que el modelo no tuviera u , sería la beta la variable aleatoria.

$$\vec{y} = X\vec{\beta} + \vec{u} \quad \begin{aligned} E(u) &= 0 \\ \text{Var}(u) &= \sigma^2 \\ E(\vec{\beta}) & \\ \text{Var}(\vec{\beta}) & \end{aligned}$$

(Con la esperanza) -> No creo que caiga en el examen, es necesario saber que es.

DEMOSTRACIÓN DE QUE UN ESTIMADOR INSESGADO $\hat{\vec{\beta}}$ es un estimador insesgado de $\vec{\beta}$.

Debemos saber las siguientes afirmaciones:

El estimador EMCO de $\vec{\beta}$ del modelo $\vec{y} = X\vec{\beta} + \vec{u}$ es:

$$\hat{\vec{\beta}} = (X^T X)^{-1} (X^T \vec{y})$$

Y sustituyendo ...

$$\left. \begin{aligned} \hat{\vec{\beta}} &= (X^T X)^{-1} X^T \vec{y} \\ \vec{y} &= X\vec{\beta} + \vec{u} \end{aligned} \right\} \hat{\vec{\beta}} = (X^T X)^{-1} X^T (X\vec{\beta} + \vec{u}) = \underbrace{(X^T X)^{-1} X^T X}_{=I} \vec{\beta} + (X^T X)^{-1} X^T \vec{u}$$

$$E(\hat{\vec{\beta}}) = \vec{\beta} + E[(X^T X)^{-1} X^T \vec{u}] = \vec{\beta}$$

Vemos que si es un estimador insesgado.

(Con la varianza) -> No creo que caiga en el examen, es necesario saber que es.

DEMOSTRACIÓN DE LA VARIANZA DE LOS EMCO

Sabemos de TC II que:

$$\text{var}(x) = E(x - E(x))^2$$

al ser por matrices

Pues para esta demostración necesitaremos saber que*:

$$\text{var}(\vec{\beta}) = E[(\hat{\vec{\beta}} - E(\hat{\vec{\beta}}))(\hat{\vec{\beta}} - E(\hat{\vec{\beta}}))^T]$$

Y esas diferencias, gracias a la demostración anterior sabemos que:

$$\hat{\vec{\beta}} - \vec{\beta} = (X^T X)^{-1} X^T \vec{u}$$

Ahora podremos desarrollar lo de la varianza*:

$$= E[(X^T X)^{-1} X^T \vec{u} \cdot \vec{u}^T X (X^T X)^{-1}] =$$

$$= E[(X^T X)^{-1} X^T \underbrace{E(\vec{u} \cdot \vec{u}^T)}_{\text{solo de las vars. aleatorias}} X (X^T X)^{-1}] = \sigma^2 [(X^T X)^{-1} X^T X (X^T X)^{-1}] = \sigma^2 (X^T X)^{-1}$$

DEMOSTRACIÓN DE TEOREMA GAUSS-MARKOV (95% de que caiga en el examen)**¿En qué consiste?**

El Teorema de Gauss-Markov establece que, en un modelo de regresión lineal bajo ciertas condiciones, **los estimadores de mínimos cuadrados ordinarios** (OLS, por sus siglas en inglés) son **los estimadores lineales insesgados con varianza mínima**. Esto significa que, entre todos los estimadores lineales insesgados, los estimadores OLS **son los más eficientes** (es decir, tienen la menor varianza).

¿Para qué sirve?

Este teorema es fundamental en econometría y estadística, ya que garantiza que el método de mínimos cuadrados produce estimaciones óptimas bajo los supuestos adecuados. Se utiliza en la estimación de parámetros en modelos lineales para obtener predicciones precisas y fiables.

DEMOSTRACIÓN: Debemos saber ciertas afirmaciones en primer lugar (denotando W):

$$\boxed{\text{MCO}} - \text{LINEAL} \rightarrow \vec{\beta} = (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \vec{y} = \overset{\text{lineal}}{\mathbf{W}} \vec{y}$$

$$- \text{INSESADO} \rightarrow E(\vec{\hat{\beta}}) = \vec{\beta}$$

5 / 34

Otro estimador lineal e insesgado:

$$\vec{\hat{\beta}}^* = \mathbf{C} \cdot \vec{y} ; E(\vec{\hat{\beta}}^*) = \vec{\beta}^* = E(\mathbf{C} \cdot \vec{y})$$

$$- E(\mathbf{C} \cdot \vec{y}) = E[\mathbf{C} \cdot (\mathbf{X} \vec{\beta} + \vec{u})] = E[\mathbf{C} \mathbf{X} \vec{\beta}] + E[\mathbf{C} \vec{u}] =$$

$$= \underbrace{\mathbf{C} \mathbf{X} \vec{\beta}}_{\text{valores fijos}} + 0 = \mathbf{C} \mathbf{X} \vec{\beta} \quad \hookrightarrow \vec{\hat{\beta}}^* - \vec{\hat{\beta}} = \mathbf{C} \cdot \vec{u}$$

$$- \text{Para que sea insesgado: } \mathbf{C} \cdot \mathbf{X} = \mathbf{I} \rightarrow E(\vec{\hat{\beta}}^*) = \vec{\beta}^*$$

$$- \text{var}(\vec{\hat{\beta}}^*) = E[(\vec{\hat{\beta}}^* - E(\vec{\hat{\beta}}^*))(\vec{\hat{\beta}}^* - E(\vec{\hat{\beta}}^*))^t] = E(\mathbf{C} \cdot \vec{u} \cdot \vec{u}^t \cdot \mathbf{C}^t)$$

$$= \mathbf{C} E(\vec{u} \vec{u}^t) \mathbf{C}^t = \sigma^2 \cdot \mathbf{C} \cdot \mathbf{C}^t$$

$$\quad \quad \quad \sigma^2 \mathbf{I}$$

Tenemos que: $\hat{\beta} = (X^t X)^{-1} X^t y = W y$ } $D = C - W$
 $\hat{\beta}^* = 0 \cdot y$ } $D = C - W \Rightarrow C = W + D$

$$Dx = (C - W)x = Cx - Wx = Cx - (X^t X)^{-1} X^t x = 0$$

$\downarrow$ Param q sea insesgado

$$\begin{aligned} \text{Var}(\hat{\beta}^*) &= \sigma^2 C C^t = \sigma^2 (W + D)(W + D)^t = \\ &= \sigma^2 [W W^t + W D^t + D W^t + D D^t] \\ &= \sigma^2 [W W^t + 0 + 0 + D D^t] = \sigma^2 [W W^t + D D^t] \end{aligned}$$

$$W W^t = (X^t X)^{-1} X^t X (X^t X)^{-1} = (X^t X)^{-1}$$

$$\left\{ \begin{array}{l} \text{Var}(\hat{\beta}) = \sigma^2 (X^t X)^{-1} \\ \text{Var}(\hat{\beta}^*) = \sigma^2 (X^t X)^{-1} + \sigma^2 (D D^t) \end{array} \right\}$$

≥ 0

Finalizando con la conclusión de que la varianza del otro estimador beta será mayor a la varianza del beta original. (Es decir, es el más eficiente)

ESTIMACIÓN DE σ_u^2 (Lo que hay que saber)

$$\hat{\sigma}_u^2 = \frac{\vec{e}^t \vec{e}}{n - k} = \frac{SCR}{n - k}$$

$$E[\vec{e}^t \vec{e}] = \sigma^2 (n - k)$$

(Ignora la l dibujada en la igualdad)

Recuerda que SCR:

$$\vec{y}^t \vec{y} - \vec{\beta}^t X^t \vec{y}$$

$$\begin{aligned} \hat{\beta} &= (X^t X)^{-1} X^t y \\ E(\hat{\beta}) &= \beta \\ \text{var}(\hat{\beta}) &= \sigma^2 (X^t X)^{-1} \\ \Downarrow \\ \widehat{\text{var}}(\hat{\beta}) &= \hat{\sigma}^2 (X^t X)^{-1} \end{aligned}$$

(Aquí hicimos ejercicios, que estarán en el archivo correspondiente)

ESTIMACIÓN DE LA VARIANZA DE LOS EMCO (~~Saber usarlos, nos darán hoja formu.~~)

$$\text{var} \left(\begin{pmatrix} \vec{\beta} \end{pmatrix} \right) = \sigma^2 \cdot (X^t X)^{-1} \quad \hat{\sigma}^2 = \frac{\vec{e}^t \vec{e}}{n-k} = \frac{SCR}{n-k}, \quad \widehat{\text{var}} \left(\begin{pmatrix} \vec{\beta} \end{pmatrix} \right) = \frac{SCR}{n-k} \cdot (X^t X)^{-1}$$

SUPUESTO DE NORMALIDAD E INFERENCIA

(Más fórmulas que no hacen falta saberse pero sí usar en ejercicios)

► Suponíamos que $\vec{u} \sim \mathcal{N}(\vec{0}, \sigma^2 I_n)$.

► Como $\vec{\beta} = \vec{\beta} + (X^t X)^{-1} X^t \vec{u}$, y por T. Gauss-Markov:

$$\vec{\beta} \sim \mathcal{N} \left(E \left[\begin{pmatrix} \vec{\beta} \end{pmatrix} \right], \text{var} \left(\begin{pmatrix} \vec{\beta} \end{pmatrix} \right) \right) = \mathcal{N} \left(\vec{\beta}, \sigma^2 \cdot (X^t X)^{-1} \right)$$

o equivalentemente:

$$\left(\begin{pmatrix} \vec{\beta} \end{pmatrix} - \vec{\beta} \right) \sim \mathcal{N} \left(\vec{0}, \sigma^2 \cdot (X^t X)^{-1} \right)$$

Concluimos con que sigue una chi-cuadrado de n-k grados de libertad:

$$\text{Concluimos que } \frac{1}{\sigma^2} \vec{u}^t M_X \vec{u} \sim \chi_{n-k}^2 \Rightarrow \frac{(n-k)\hat{\sigma}^2}{\sigma^2} \sim \chi_{n-k}^2$$

INTERVALOS DE CONFIANZA:

“Para la media con varianza desconocida” de TC2:

► IC para β_i :

$$\hat{\beta}_i \pm t_{n-k, 1-\frac{\alpha}{2}} \cdot \sqrt{\widehat{\text{var}}[\hat{\beta}_i]}, \quad i = 1, \dots, k. \quad \text{var} \left(\begin{pmatrix} \vec{\beta} \end{pmatrix} \right) = \sigma^2 \cdot (X^t X)^{-1}$$

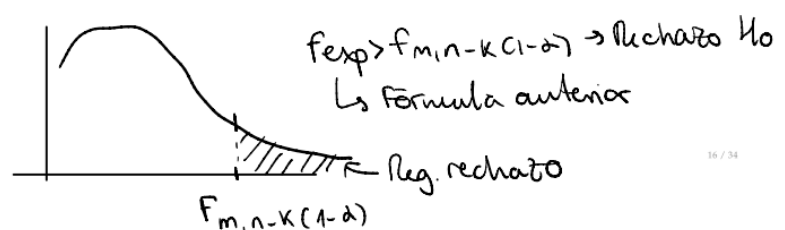
► IC para σ^2

$$\left[\frac{(n-k) \cdot \hat{\sigma}^2}{\chi_{n-k, 1-\frac{\alpha}{2}}^2}, \frac{(n-k) \cdot \hat{\sigma}^2}{\chi_{n-k, \frac{\alpha}{2}}^2} \right], \quad \hat{\sigma}_u^2 = \frac{\vec{e}^t \vec{e}}{n-k} = \frac{SCR}{n-k}$$

CONTRASTE DE HIPÓTESIS:

La hipótesis nula se rechazará si el valor del estadístico cae en la región de rechazo.

Es decir, si $F_{exp} > F_{m, n-k, 1-\alpha}$ entonces se rechaza la hipótesis $H_0 : R\vec{\beta} = \vec{r}$.



PARA LOS EJERCICIOS**CASOS PARTICULARES (i): SIGNIFICACIÓN INDIVIDUAL**

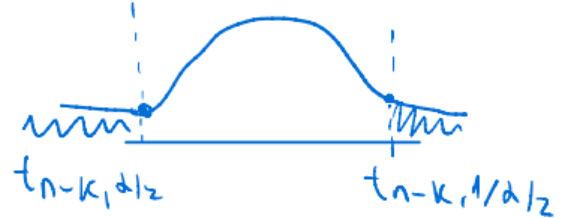
Si $m = 1$ y $r_i = 0 \forall i$.

► Hipótesis:

$$\begin{cases} H_0 : \beta_i = 0 \\ H_1 : \beta_i \neq 0 \end{cases}$$

► Estadístico de contraste:

$$t_{exp} = \frac{\hat{\beta}_i}{\sqrt{\widehat{\text{var}}[\hat{\beta}_i]}}$$



► Conclusión: Se rechaza H_0 si $t_{exp} > t_{n-k, 1-\frac{\alpha}{2}}$.
→ Val. abs.

CASOS PARTICULARES (ii): ÚNICA RESTRICCIÓN LINEAL

$m = 1$ y $r_i \neq 0 \forall i$.

► Hipótesis: $\begin{cases} H_0 : \beta_i = r_i \\ H_1 : \beta_i \neq r_i \end{cases}$

► Estadístico de contraste:

$$t_{exp} = \frac{\hat{\beta}_i - r_i}{\sqrt{\widehat{\text{var}}[\hat{\beta}_i]}}$$

► Conclusión: Se rechaza H_0 si $t_{exp} > t_{n-k, 1-\frac{\alpha}{2}}$.

CASOS PARTICULARES (iii): ÚNICA RESTRICCIÓN LINEAL

Si $m = k - 1$ y $r_i = 0 \ i = 2, 3, \dots, k$.

► Hipótesis: $\begin{cases} H_0 : \beta_2 = \beta_3 = \dots = \beta_k = 0 \\ H_1 : \exists \beta_i \neq 0 \text{ con } i = 2, 3, \dots, k \end{cases}$

► Estadístico de contraste:

$$F_{exp} = \left(R \vec{\beta} \right)^t \cdot \frac{\left[R (X^t X)^{-1} R^t \right]^{-1}}{(k-1) \cdot \hat{\sigma}^2} \cdot \left(R \vec{\beta} \right)$$

► Conclusión: Se rechaza H_0 si $F_{exp} > F_{k-1, n-k, 1-\alpha}$.

→ Todos menos el término indep. (β_1)

ANOVA

Considerando que se verifica que $SCT = SCR + SCE$

Test sign. global: $\begin{cases} H_0 : \beta_2 = \beta_3 = \dots = 0 \\ H_1 : \exists \beta_i \neq 0 \end{cases}$

$$F_{exp} = \frac{\frac{1}{k-1} SCE}{\frac{1}{n-k} SCR}$$

Regla decisión: $F_{exp} > F_{k-1, n-k} (1-\alpha)$
 rechaza H_0 .

TABLA ANOVA

Fuente de variación	Suma de Cuadrados	Grados de Libertad	Medias
Explicada	$SCE = \vec{\beta}^t X^t \vec{y} - n\bar{Y}^2$	$k - 1$	$\frac{SCE}{k-1}$
Residuos	$SCR = \vec{y}^t \vec{y} - \vec{\beta}^t X^t \vec{y}$	$n - k$	$\frac{SCR}{n-k}$
Total	$SCT = \vec{y}^t \vec{y} - n\bar{Y}^2$	$n - 1$	

Como

$$R^2 = \frac{SCE}{SCT} = 1 - \frac{SCR}{SCT}$$

$$F_{exp} = \frac{\frac{1}{k-1} SCE}{\frac{1}{n-k} SCT(1-R^2)} = \frac{\frac{1}{k-1} R^2}{\frac{1}{n-k} (1-R^2)} = \frac{n-k}{k-1} \frac{R^2}{(1-R^2)}$$

COTA R²

Se rechaza H_0 cuando $F_{exp} > F_{k-1, n-k, 1-\alpha}$, esto es si

$$\frac{n-k}{k-1} \frac{R^2}{(1-R^2)} > F_{k-1, n-k, 1-\alpha} \Rightarrow R^2 > \frac{(k-1) \cdot F_{k-1, n-k, 1-\alpha}}{(n-k) + (k-1) \cdot F_{k-1, n-k, 1-\alpha}}$$

PREDICCIÓN PUNTUAL

Tiene que coincidir el individual y el esperado en la pred. puntual, veamos en un ejemplo:

INDIVIDUAL $\rightarrow \hat{y}_0$ (predicción de y) $= x_0^t \vec{\beta}$

PREDICCIÓN $\left\{ \begin{array}{l} \text{VALOR ESPERADO} \rightarrow E(\hat{y}_0 | x_0) \text{ (Predicción del val. esp. de } y) \end{array} \right.$

Ej: $\hat{y}_t = 3 + 2x_{2t} + 4x_{3t}$

$\left. \begin{array}{l} x_{2t} = 1 \\ x_{3t} = 2 \end{array} \right\} \hat{y}_t = 3 + 2 \cdot 1 + 4 \cdot 2 = 13 = E(\hat{y}_0 | x_0) = \hat{y}_0$

$\hat{y}_0 = E(\hat{y}_0 | x_0) = \underset{\substack{\uparrow \\ \text{t.i.}}}{(1 \ 1 \ 2)} \begin{pmatrix} 3 \\ 2 \\ 4 \end{pmatrix} = 13$

PREDICCIÓN PUNTUAL para EL VALOR INDIVIDUAL (esperanza, no saber)

• Predicción del valor individual:

$$\left. \begin{array}{l} y_0 = x_0^t \vec{\beta} + \vec{u}_0 \\ \hat{y}_0 = x_0^t \vec{\hat{\beta}} \end{array} \right\} e_0 = x_0^t \vec{\beta} + \vec{u}_0 - x_0^t \vec{\hat{\beta}}$$

$$e_0 = \vec{u}_0 - x_0^t (\vec{\hat{\beta}} - \vec{\beta})$$

$$\cdot \left[E(e_0) = E(\vec{u}_0) - x_0^t [E(\vec{\hat{\beta}}) - E(\vec{\beta})] = 0 \right]$$

PREDICCIÓN PUNTUAL para EL VALOR INDIVIDUAL (varianza, esta tengo apuntado que si sabersela, pero no se)

$$\begin{aligned}
 \text{Var}(e_0) &= \text{Var}[\vec{u}_0 - x_0^t (\vec{\beta} - \vec{\hat{\beta}})] = \\
 &= \text{Var}(\vec{u}_0) + \text{Var}(x_0^t (\vec{\beta} - \vec{\hat{\beta}})) - 2\text{Cov}[\vec{u}_0, x_0^t (\vec{\beta} - \vec{\hat{\beta}})] = \\
 &= \sigma^2 + E[x_0^t (\vec{\beta} - \vec{\hat{\beta}}) (\vec{\beta} - \vec{\hat{\beta}})^t x_0] = \sigma^2 + x_0^t E[(\vec{\beta} - \vec{\hat{\beta}}) (\vec{\beta} - \vec{\hat{\beta}})^t] x_0 = \\
 &= \sigma^2 + x_0^t \sigma^2 (X^t X)^{-1} x_0 = \sigma^2 (1 + x_0^t (X^t X)^{-1} x_0)
 \end{aligned}$$

★ $\text{Var}(X) = E(X - E(X))^2$
 ↳ Con matrices en vez de x^2 , se hace por la traspuesta

$$e_0 \sim N(0, \sigma^2 (1 + x_0^t (X^t X)^{-1} x_0))$$

• Predicción del valor esperado:

$$\begin{aligned}
 E(y_0/x_0) &= x_0^t \vec{\beta} \\
 E(y_0/x_0) &= x_0^t \vec{\hat{\beta}}
 \end{aligned}
 \quad \left. \begin{array}{l} \\ \end{array} \right\} e_0^* = x_0^t (\vec{\hat{\beta}} - \vec{\beta})$$

↳ me viene bien dpo para llegar a la matriz de vars-covars

$$E(e_0^*) = E[-x_0^t (\vec{\hat{\beta}} - \vec{\beta})] = -x_0^t (E(\vec{\hat{\beta}}) - E(\vec{\beta})) = 0$$

$$\begin{aligned}
 \text{Var}(e_0^*) &= E(x_0^t (\vec{\hat{\beta}} - \vec{\beta}) (\vec{\hat{\beta}} - \vec{\beta})^t x_0) = \\
 &= x_0^t E[(\vec{\hat{\beta}} - \vec{\beta}) (\vec{\hat{\beta}} - \vec{\beta})^t] x_0 = \sigma^2 x_0^t (X^t X)^{-1} x_0
 \end{aligned}$$

Resumen:

Valor teórico	Valor estimado	Residuo	Distribución
y_0	$\hat{y}_0$	e_0	$N(0, \sigma^2 (1 + x_0^t (X^t X)^{-1} x_0))$
$E(y_0/x_0)$	$E(\hat{y}_0/x_0)$	e_0^*	$N(0, \sigma^2 x_0^t (X^t X)^{-1} x_0)$

$$\hat{y}_0 = E(\hat{y}_0/x_0) = x_0^t \vec{\hat{\beta}}$$

PREDICCIÓN POR INTERVALO

	Predicción puntual	Predicción por intervalos
v. individual	$x_0^t \vec{\hat{\beta}}$	$x_0^t \vec{\hat{\beta}} \pm t_{n-k(1-\alpha/2)} \hat{\sigma} \sqrt{1 + x_0^t (X^t X)^{-1} x_0}$
v. esperado	$x_0^t \vec{\hat{\beta}}$	$x_0^t \vec{\hat{\beta}} \pm t_{n-k}(1-\alpha/2) \hat{\sigma} \sqrt{x_0^t (X^t X)^{-1} x_0}$

Recordatorio de distribuciones (TC2) + Ej 11/12 importante (en archivo de ejercicios)

DISTRIBUCIONES χ^2 Y t -STUDENT

- La distribución χ^2 de Pearson con n grados de libertad, χ_n^2 , se construye como la suma de los cuadrados de n variables aleatorias independientes e idénticamente distribuidas según una $\mathcal{N}(0, 1)$:

$$\chi_n^2 = \sum_{i=1}^n X_i^2, \quad X_i \sim N(0, 1), \quad \forall i.$$

- La distribución t -Student con n grados de libertad, t_n , se construye como el cociente entre una variable aleatoria $\mathcal{N}(0, 1)$ y la raíz cuadrada de una χ^2 -cuadrado de n grados de libertad dividida entre sus grados de libertad, siendo ambas distribuciones independientes:

$$t_n = \frac{X}{\sqrt{\frac{Y}{n}}}, \quad X \sim N(0, 1), \quad Y \sim \chi_n^2.$$

Propiedades:

- χ^2 NO es simétrica, t -Student SÍ es simétrica.
- Para $n > 30$, la distribución t -Student se puede aproximar a una distribución Normal.

32 /

DISTRIBUCIÓN F -SNEDECOR

- La distribución F de Snedecor con n y m grados de libertad, que denotaremos por $F_{n,m}$, se construye a partir del cociente de dos variables aleatorias independientes y distribuidas según chi-cuadrados con n y m grados de libertad, respectivamente, divididas entre sus correspondientes grados de libertad. Por tanto, la distribución F de Snedecor responde a la siguiente estructura:

$$F_{n,m} = \frac{\frac{X}{n}}{\frac{Y}{m}}, \quad X \sim \chi_n^2, \quad Y \sim \chi_m^2.$$

Propiedades:

- La distribución $F_{n,m}$ NO es simétrica.
- $F_{m,n,1-\alpha} = \frac{1}{F_{n,m,\alpha}}$.
- $t_n^2 = F_{1,n}$.

33 /

Notación matricial (nada)

Tema 4: Multicolinealidad

CONCEPTO:

La multicolinealidad ocurre cuando hay relaciones lineales entre las variables independientes en un modelo (2 o más variables).

Puede ser perfecta (exacta) o aproximada.

Dentro de la aproximada, puede clasificarse como errática (inesperada) o sistemática (predecible), y esencial (afecta el modelo) o no esencial (no afecta significativamente).

MULTICOLINEALIDAD EXACTA

Concepto: Hace referencia a que existe una relación lineal exacta entre 2 o más var. indep.

Causas: Esta multicolinealidad se traduce en el incumplimiento de la hipótesis básica del modelo uniecuacional múltiple: la matriz X no es de rango completo por columnas, $rg(X) < k$.

Consecuencias: No permite que se invierta la matriz $X^t \cdot X$ por lo que ya no hay una solución única, sino infinitas.

MULTICOLINEALIDAD APROXIMADA

Concepto: Hace referencia a que existe una relación lineal aprox. entre 2 o más var. indep.

Causas:

1. Relación causal entre las variables explicativas del modelo
2. Escasa variabilidad en las observaciones de las variables independientes
3. Reducido tamaño de la muestra

TIPOS DE MULTICOLINEALIDAD APROXIMADA según sus causas:

1. Spanos y McGuirk:

- **Multi. Sistemática:** Por alta correlación lineal entre variables exógenas.
- **Multi. Errática:** Por problemas numéricos o mal manejo de datos.

2. Marquandt y Snee:

- **Multi. No esencial:** Relación lineal de las variables exógenas con la constante (se corrige centrando las variables).
- **Multi. Esencial:** Relación lineal entre las variables exógenas, excluyendo la constante. (Endógena=Explicada $\rightarrow$ Exógenas=Explicativas/Independientes)

Luego, se podrían distinguir los siguientes cuatro casos:

Multicolinealidad	Sistemática	Errática
No esencial	1	2
Esencial	3	4

Consecuencias: No se cumplirá la hipótesis básica de que la matriz X es completa por columnas ($\text{rg}(X)=k$) por lo que se podrá invertir $X^t \cdot X$ y obtener estimadores por MCO, pero el determinante será próximo a 0, por lo que la inversa de lo anterior será alto y por tanto:

- Varianzas altas en los estimadores.
- Los contrastes de significación individuales no rechazan la hipótesis nula, pero las conjuntas sí.
- Los coeficientes estimados son inestables ante cambios pequeños.
- R^2 y R^2 corregido elevado (coeficiente de determinación elevado)

Ejemplo: La edad y la experiencia suelen tener una alta relación, por lo tanto será difícil separar el efecto de cada una sobre la variable dependiente y que se produzca multicolinealidad debido a la relación entre ambas variables.

EFFECTOS NOCIVOS SOBRE EL ANÁLISIS ESTADÍSTICO DEL MODELO

(Ejercicios y ejemplos que estarán en el archivo de ejercicios)

FACTOR DE INFLACIÓN DE LA VARIANZA (FIV) ó VIF en inglés

Es una de las medidas más usadas para detectar el grado de multicolinealidad existente:

$$FIV(i) = \frac{1}{1 - R_i^2}, \quad i = 2, \dots, p,$$

↓
Para la esencial

Si esta medida es superior a 10, se supone que el grado de multicolinealidad presente es preocupante. Recordemos que el VIF no tiene en cuenta la relación de las variables exógenas del modelo, por tanto, no detecta la multicolinealidad no esencial.

NÚMERO DE CONDICIÓN (NC)

$$NC = \sqrt{\frac{\lambda_{\max}}{\lambda_{\min}}},$$

Si esta medida es superior a 20 se supone que el grado de multicolinealidad es moderado, y si es superior a 30 es preocupante. Este si tiene en cuenta la relación de las variables exógenas del modelo.

→ Para la no esencial

OTRAS MEDIDAS

1. Matriz de correlaciones simples (R):

- Analiza las relaciones entre pares de variables independientes (dos a dos).
- Ignora la constante y no puede detectar multicolinealidad más compleja.

2. Determinante de la matriz de correlaciones (det(R)):

- Considera estructuras más complejas que las relaciones dos a dos.

- También ignora la constante.
- Valores cercanos a cero indican una multicolinealidad (del tipo esencial) grave.

3. Coeficiente de variación de las variables explicativas (CV(Xi)):

- Identifica baja variabilidad en las variables independientes.
- Esto puede causar multicolinealidad no esencial.

SOLUCIONES Para abordar el problema de multicolinealidad:

1. **Mejorar el diseño muestral:** Maximizar la información obtenida de las variables observadas.
2. **Eliminar variables problemáticas:** Quitar aquellas que sean sospechosas de causar multicolinealidad.
3. **Aumentar el tamaño de la muestra:** Especialmente útil cuando se tienen pocas observaciones.
4. **Usar información extramuestral:** Incorporar información a priori y estimar el modelo mediante métodos como mínimos cuadrados restringidos.

Además, algunos autores sugieren tratar el problema con métodos numéricos, como:

- Estimación cresta.
- Estimación alzada. (desde la perspectiva geométrica)
- Residualización.

Resumenillo

	Solución
Errática (esencial- no esencial)	• Ampliar la muestra • Mejorar la muestra (↑ variabilidad)
Sistemática (esencial)	• A priori ($\beta_i = b_i$) • métodos alternativos de estimación • Eliminar variable - cresta - alzada - residualización - etc...
Sistemática (no esencial)	• Centrar la variable • Eliminar la variable

Tema 5: Heterocedasticidad

CONCEPTO: Se dice que un modelo de regresión presenta heteroscedasticidad cuando la varianza del término de perturbación no permanece constante a lo largo del tiempo.
(Varianza no cte.)

CAUSAS:

1. **Datos de sección cruzada:**
 - La escala de la variable dependiente y el poder explicativo del modelo pueden variar entre observaciones.
2. **Uso de datos agrupados:**
 - Al agrupar observaciones en categorías y trabajar con sus medias en lugar de los datos originales, se introduce heterocedasticidad, aunque los datos originales sean homoscedásticos.
3. **Omisión de una variable relevante:**
 - Si una variable importante no está incluida en el modelo, la varianza del error puede depender de esta variable omitida, causando heterocedasticidad.

CONSECUENCIAS:

Puesto que en el método de estimación por MCO no influye la matriz de varianzas-covarianzas de la perturbación aleatoria es claro que el estimador por MCO será igualmente:

$$\hat{\beta}_{MCO} = (X^t X)^{-1} X^t y.$$

$\hat{\beta}_{MCO}$ es lineal e insesgado. Sin embargo, ya no se tiene asegurado que la varianza sea mínima:

$$\begin{aligned} \text{Var}(\hat{\beta}_{MCO}) &= (X^t X)^{-1} X^t \mathbb{E}[u \cdot u^t] X (X^t X)^{-1} \\ &= \sigma^2 (X^t X)^{-1} X^t \Omega X (X^t X)^{-1}, \end{aligned}$$

distinta a la del modelo con perturbaciones esféricas: $\sigma^2 (X^t X)^{-1}$.

Por tanto, la consecuencia de la presencia de heteroscedasticidad en un modelo lineal es que los estimadores obtenidos, aunque serán lineales e insesgados, no serán óptimos.

7 / 26

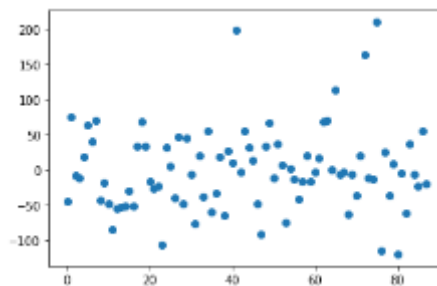
PROCEDIMIENTOS DE DETECCIÓN de HETEROCEDASTICIDAD, 2 enfoques:

1. **Métodos gráficos:**
 - Analizar gráficos de residuos y de dispersión para identificar visualmente posibles variables que causen heterocedasticidad.
2. **Métodos analíticos:**
 - Realizar tests específicos, como:

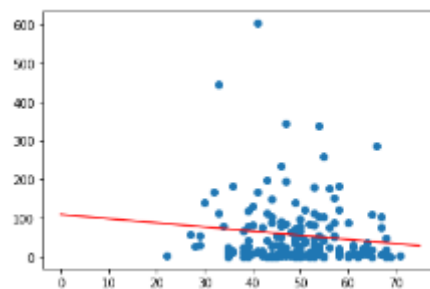
- Glejser y Goldfeld-Quandt: Para muestras pequeñas, cuando se sospecha que una variable específica causa el problema.
- Breusch-Pagan y White: Para muestras grandes, cuando no se conoce la causa exacta de la heterocedasticidad.
- En todos estos casos, la hipótesis nula es que el modelo es homocedástico (varianza constante).

MÉTODOS GRÁFICOS

GRÁFICOS DE LOS RESIDUOS: Si en dichos gráficos observamos grupos de observaciones con distinta varianza, podemos pensar en la presencia de heteroscedasticidad.



GRÁFICOS DE DISPERSIÓN: Si la variabilidad de los residuos aumenta o disminuye conforme aumenta el valor de la variable independiente, entonces podemos pensar que la varianza de la perturbación aleatoria depende de dicha variable y, por tanto, habría presencia de heteroscedasticidad.



CONTRASTE DE GOLDFELD QUANDT

(Ejemplo en el archivo de ejercicios y presentaciones)

Para muestras pequeñas, y si sospechamos que σ_i^2 está relacionada positivamente con la variable X_i :

1. Ordenar la información muestral, según los valores crecientes de este regresor, de menor a mayor.
2. Omitir los m valores muestrales centrales, (por ejemplo $m = \frac{n}{3}$).
3. Ajustar, mediante MCO, cada uno de los subgrupos obtenidos en el paso 2, una vez eliminados los valores centrales. Cada uno de estos subgrupos estarán formados por $\frac{n-m}{2}$ observaciones. La estimación del modelo del primer subgrupo nos permitirá obtener la SCR_1 y la del subgrupo 2, la SCR_2 .
4. Bajo ausencia de heteroscedasticidad ($\sigma_1^2 = \sigma_2^2 = \sigma^2$):

$$F_{exp} = \frac{SCR_2}{SCR_1} > F_{\frac{n-m}{2}-k, \frac{n-m}{2}-k, 1-\alpha}$$

se rechaza la hipótesis de homocedasticidad (hay heteroscedasticidad)

`sms.het_goldfeldquandt(mco.resid, mco.model.exog)`

CONTRASTE DE BREUSCH-PAGAN Y WHITE (Solo lo hicimos por comando)

Para número elevado de observaciones:

1. Estimamos el modelo $y = X\beta + u$ y el vector correspondiente a los residuos de mínimos cuadrados ordinarios, $e = y - \hat{y}$.

2. Breusch-Pagan:

$$e^2 = \delta_0 + \delta_1 x_1 + \dots + \delta_p X_p + v$$

3. White:

$$e^2 = \delta_0 + \delta_1 x_1 + \dots + \delta_p X_p + \delta_{p+1} x_1^2 + \delta_{p+2} x_1 x_2 + \dots + \delta_{p+p^2} x_p^2 + v$$

4. Se contrasta $H_0: \delta_1 = \delta_2 = \dots = 0$ (HOMOCEDASTICIDAD)

5. Si se rechaza H_0 : Existe heterocedasticidad!

```
sms.het_breuschpagan(mco.resid, mco.model.exog)
sms.het_white(mco.resid, mco.model.exog)
```

CONTRASTE DE GLEJER (ejemplo en diapos, pero no es importante)

Si sospechamos de que la heteroscedasticidad está ligada a X_i :

1. Estimamos el modelo $y = X\beta + u$ y $e = y - \hat{y}$.
2. Se ajusta por MCO el modelo en el que la variable endógena es $|e_t|$ y la variable exógena es la variable X_i , la cual pensamos que es la que provoca la heteroscedasticidad, que será denotada por z_t , es decir

$$|e_t| = \delta_0 + \delta_1 z_t^h + v_t$$

donde v_t es ruido blanco. La regresión se realiza para distintos valores de $h = \pm 1, \pm 2, \pm \frac{1}{2}, \dots$

3. Para cada h , contrastar $H_0: \delta_1 = 0$ (t-test). Si se rechaza H_0 , entonces existe heteroscedasticidad en el modelo.

ESTIMACIÓN DE MODELOS CON HETEROCEDASTICIDAD

Si conocemos *omega* Ω = Aplicación directa de MCG

Si no conocemos MCGF

MÍNIMOS CUADRADOS GENERALIZADOS (MCG)

Como $\Omega \succ 0$ y simétrica, por Descomposición de Cholesky, existirá una matriz $P (= Q^{-1})$, no singular, tal que,

$$\Omega^{-1} = P^t P$$

De manera que se podría transformar el modelo original para que la matriz Ω se convierta en la matriz identidad y se pudiera aplicar MCO sobre el modelo transformado aplicando estas expresiones:

$$\begin{aligned} \hat{\beta}_{MCG} &= [(X^*)^t X^*]^{-1} [(X^*)^t y^*] = [X^t P^t P X]^{-1} [X^t P^t P y] = \\ &= [X^t \Omega^{-1} X]^{-1} [X^t \Omega^{-1} y] \end{aligned}$$

$$\text{var} \left[\hat{\beta}_{MCG} \right] = \sigma^2 [(X^*)^t X^*]^{-1} = \sigma^2 [X^t P^t P X]^{-1} = \sigma^2 [X^t \Omega^{-1} X]^{-1}$$

ESTIMADOR MÍNIMOS CUADRADOS PONDERADOS

En el caso particular de la existencia de heterocedastidad, ¿Cómo se descompone $\Omega^{-1} = P^t P$?

$$\Omega = \begin{pmatrix} w_1 & 0 & \dots & 0 \\ 0 & w_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & w_n \end{pmatrix} \Rightarrow \Omega^{-1} = \begin{pmatrix} \frac{1}{w_1} & 0 & \dots & 0 \\ 0 & \frac{1}{w_2} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \frac{1}{w_n} \end{pmatrix}$$

Considerando:

$$P = \begin{pmatrix} \frac{1}{\sqrt{w_1}} & 0 & \dots & 0 \\ 0 & \frac{1}{\sqrt{w_2}} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \frac{1}{\sqrt{w_n}} \end{pmatrix}$$

MÍNIMOS CUADRADOS PONDERADOS

$$y^* = Py \quad X^* = PX \quad u^* = Pu$$

$$y^* = \begin{pmatrix} \frac{y_1}{\sqrt{w_1}} \\ \vdots \\ \frac{y_n}{\sqrt{w_n}} \end{pmatrix}, X^* = \begin{pmatrix} \frac{1}{\sqrt{w_1}} & \frac{x_{21}}{\sqrt{w_1}} & \dots & \frac{x_{k1}}{\sqrt{w_1}} \\ \frac{1}{\sqrt{w_2}} & \frac{x_{22}}{\sqrt{w_2}} & \dots & \frac{x_{k2}}{\sqrt{w_2}} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{1}{\sqrt{w_n}} & \frac{x_{2n}}{\sqrt{w_n}} & \dots & \frac{x_{kn}}{\sqrt{w_n}} \end{pmatrix}, u^* = \begin{pmatrix} \frac{u_1}{\sqrt{w_1}} \\ \vdots \\ \frac{u_n}{\sqrt{w_n}} \end{pmatrix}.$$

- $\mathbb{E}[u_i^*] = \mathbb{E}\left[\frac{u_i}{\sqrt{w_i}}\right] = \frac{1}{\sqrt{w_i}} \mathbb{E}[u_i] = 0.$
- $Var(u_i^*) = Var\left(\frac{u_i}{\sqrt{w_i}}\right) = \frac{1}{w_i} Var(u_i) = \frac{1}{w_i} \sigma^2 w_i = \sigma^2:$
HOMOCEDASTICIDAD!!!
- $\mathbb{E}[u_i^* u_j^*] = \mathbb{E}\left[\frac{u_i}{\sqrt{w_i}} \frac{u_j}{\sqrt{w_j}}\right] = \frac{1}{\sqrt{w_i w_j}} \mathbb{E}[u_i u_j] = 0:$ INCORRELACIÓN!!

Podemos aplicar MCO para estimar $y^* = X^* \beta + u^*$.

(En presentaciones hay un ejemplo del test de Glesjer al final, con algo que hizo en la pizarra pero tampoco tiene mucha importancia, solo por revisar)

Tema 6: Autocorrelación

¿De dónde viene la autocorrelación, cual es su naturaleza?

$$y = X\beta + u$$

con $E[u] = 0$ y $\text{var}[u] = E[uu^t] = \sigma^2 I_n$, lo que implica:

- ▶ $E[u_t] = 0, \forall t \in 1, \dots, n$
- ▶ $E[u_t^2] = \sigma^2, \forall t \in 1, \dots, n$ (Homocedasticidad)
- ▶ $E[u_i, u_j] = \text{Cov}[u_i, u_j] = 0, \forall i \neq j \in 1, \dots, n$ (Incorrelación).

Si $\text{Cov}[u_t, u_{t-k}] \neq 0 \forall k > 0 \rightarrow$ **Autocorrelación.**

CAUSAS

1. **Error de especificación:** Omitir variables exógenas relevantes para explicar la variable endógena. $\rightarrow$ explicador
2. **Ciclos y tendencias:** Si los ciclos de la variable endógena no son explicados adecuadamente por las variables exógenas, surge autocorrelación en los errores (en la perturbación).
3. Existencia de relaciones no lineales
4. Existencia de relaciones dinámicas

CONSECUENCIAS

Los estimadores MCO son insesgados pero **no óptimos**.
El estimador por Mínimos Cuadrados Generalizados (MCG)

$$\hat{\beta}_{MCG} = [X^t \Omega^{-1} X]^{-1} [X^t \Omega^{-1} y]$$

cuya matriz de varianzas-covarianzas viene dada por

$$\text{var}[\hat{\beta}_{MCG}] = \sigma^2 [X^t \Omega^{-1} X]^{-1}$$

Sí es óptimo.

Estimador mco
↓
No óptimo
Estimador MCG
↓
óptimo

DETECCIÓN DE LA AUTOCORRELACIÓN

- Gráficamente

- **Gráfico temporal de los residuos obtenidos por MCO:** Si los residuos están incorrelados deben distribuirse de forma aleatoria alrededor de 0.
- **Gráfico de dispersión de los mismo frente a algún retardo suyo:** Los residuos estarán autocorrelacionados cuando el gráfico presenta una tendencia creciente o decreciente.

- Analíticamente:

- **Contraste de Durbin-Watson.**
- **Estadístico H de Durbin.**
- **Contraste de Ljung-Box**

MÉTODOS GRÁFICOS

Para detectar autocorrelación en un modelo se pueden utilizar gráficos:

1. Gráfico temporal de los residuos:

- Residuos incorrelados: Se distribuyen aleatoriamente alrededor de cero.
- Residuos correlados:
 - Correlación Positiva: Rachas de residuos por encima y por debajo de la media.
 - Correlación Negativa: Alternancia en el signo de los residuos.

2. Gráfico de dispersión:

- Positiva: Tendencia creciente en el gráfico.
- Negativa: Tendencia decreciente.

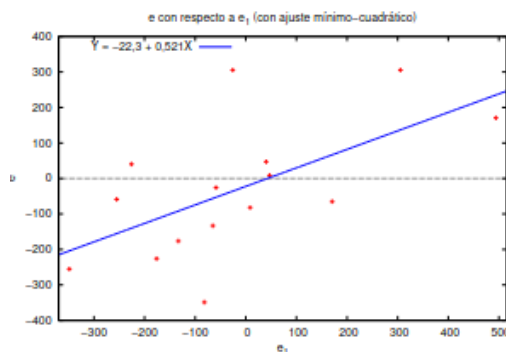


Figure. Gráfico de dispersión

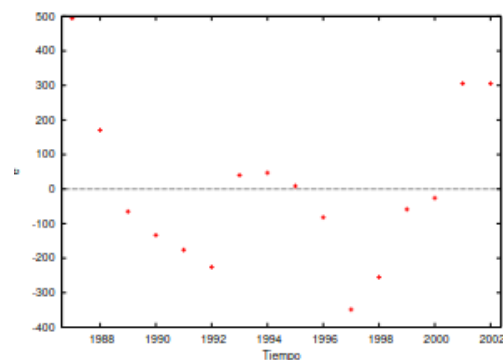


Figure. Gráfico temporal de los residuos

CONTRASTE DE DURBIN-WATSON

El contraste de Durbin-Watson resulta adecuado si suponemos que la autocorrelación que presenta el término de perturbación pudiera ser descrito mediante un proceso autorregresivo de orden uno:

$$u_t = \rho u_{t-1} + \varepsilon_t$$

donde ε_t es ruido blanco.

Estudiamos si el coeficiente ρ es significativo o no mediante los contrastes:

$$\begin{cases} H_0: \rho = 0 & \text{(Incorrelación)} \\ H_1: \rho < 0 & \text{(Correlación negativa)} \end{cases} \quad \begin{cases} H_0: \rho = 0 & \text{(Incorrelación)} \\ H_1: \rho > 0 & \text{(Correlación positiva)} \end{cases}$$

Estadístico de Contraste de Durbin-Watson:

$$DW = \frac{\sum_{t=2}^n (e_t - e_{t-1})^2}{\sum_{t=1}^n e_t^2} = \frac{\sum_{t=2}^n (e_t^2 - 2e_t e_{t-1} + e_{t-1}^2)}{\sum_{t=1}^n e_t^2}$$

siendo e_t los residuos MCO: $e_t = y_t - \hat{y}_t = y_t - x_t \hat{\beta}$.

Si la muestra seleccionada es suficientemente grande, se puede considerar que $\sum_{t=2}^n e_t^2 \approx \sum_{t=2}^n e_{t-1}^2$, ya que ambas sumatorias difieren únicamente en la primera de las observaciones muestrales:

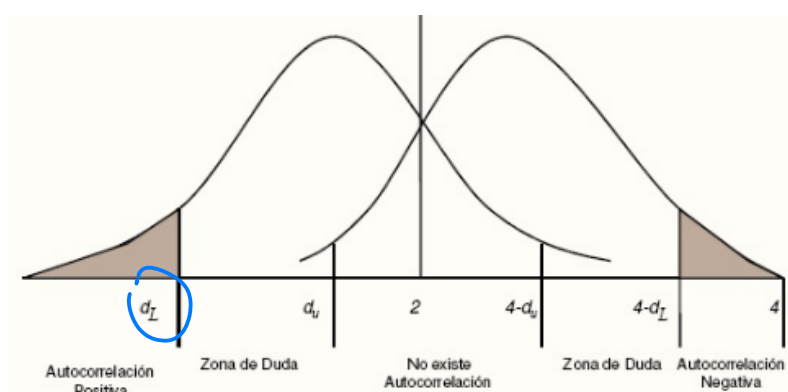
$$DW = \frac{2 \sum_{t=2}^n e_t^2 - 2 \sum_{t=2}^n e_t e_{t-1}}{\sum_{t=1}^n e_t^2} = 2 \left(\frac{\sum_{t=2}^n e_t^2}{\sum_{t=1}^n e_t^2} - \frac{\sum_{t=2}^n e_t e_{t-1}}{\sum_{t=1}^n e_t^2} \right) \approx 2(1 - \hat{\rho})$$

donde $\hat{\rho}$ es la estimación del parámetro ρ .

Como $|\rho| < 1$:

- $DW \approx 0$ cuando $\hat{\rho} \approx 1$ (correlación positiva)
- $DW \approx 2$ cuando $\hat{\rho} \approx 0$ (residuos incorrelacionados)
- $DW \approx 4$ cuando $\hat{\rho} \approx -1$ (correlación es negativa)

from statsmodels.stats.stattools import durbin_watson
durbin_watson(mco.resid)



H DE DURBIN

El contraste de Durbin-Watson es válido cuando la autocorrelación de los errores es autorregresiva de orden 1 y cuando la regresión no incluye entre las variables explicativas algún retardo de la variable dependiente.

En este segundo caso se recurre al contraste h de Durbin. Esto es, en modelos del tipo

$$Y_t = \alpha Y_{t-1} + \beta_1 + \beta_2 X_{2t} + \dots + \beta_k X_{kt} + u_t,$$

donde $u_t = \rho u_{t-1} + v_t$, para contrastar

$$H_0 : \rho = 0 \text{ (incorrelación)}$$

$$H_1 : \rho \neq 0 \text{ (correlación)}$$

utilizaremos el estadístico $h = \rho \cdot \sqrt{\frac{n}{1 - n \cdot \widehat{\text{Var}}(\hat{\rho})}} \sim N(0, 1)$, de forma que se rechazará la hipótesis nula si $|h| > Z_{1-\frac{\alpha}{2}}$.

H de Durbin solo me dice si hay o no autocorrelación, no me da el signo. Aunque podemos intuir por la estimación de “ ρ ” -> p (su signo)

CONTRASTE DE LJUNG-BOX (En presentaciones ejemplo de aplicación directa con fórmula)

Si los residuos son independientes, sus primeras m autocorrelaciones son cero, para cualquier valor de m . El test de Ljung-Box contrasta la hipótesis nula de que las primeras m autocorrelaciones, ρ_m , son cero. Esto es:

$$\left. \begin{array}{l} H_0 : \rho_1 = \rho_2 = \dots = \rho_m = 0 \\ H_1 : \exists i \in \{1, 2, \dots, m\} \text{ tal que } \rho_i \neq 0 \end{array} \right\}$$

Se rechaza la hipótesis nula (de incorrelación) si $Q_{LB} = n(n+2) \sum_{s=1}^m \frac{r(s)^2}{n-s} > \chi_{m-1}^2(1-\alpha)$ donde

$$r(s) = \frac{\sum_{t=s+1}^n e_t e_{t-s}}{\sum_{t=1}^n e_t^2},$$

es el coeficiente de autocorrelación muestral de orden k .

Si las observaciones son independientes (incorrelación), los coeficientes $r(s)$ serán próximos a cero, por lo que no se rechazaría la hipótesis nula.

```
from statsmodels.stats.diagnostic import acorr_ljungbox
acorr_ljungbox(mco.resid, lags=3)
```

ESTIMACIÓN BAJO AUTOCORRELACIÓN (MCG)

$$y = X\beta + u$$

con:

- ▶ $\mathbb{E}[u] = 0$,
- ▶ $\mathbb{E}[u_i u_j] \neq 0$ ($i \neq j$ con $i, j \in \{1, \dots, n\}$) y
- ▶ $\text{var}[u] = \sigma^2 \Omega$.

(no se verifican los supuestos de MCO)

Por el método de Mínimos Cuadrados Generalizado: necesitamos calcular P tal que $\Omega^{-1} = P^t P$.

Suponemos que u_t viene determinada por un proceso autorregresivo de orden 1, $AR(1)$: $u_t = \rho u_{t-1} + \varepsilon_t$ (ε_t un ruido blanco) y ρ el coeficiente de autocorrelación con $-1 < \rho < 1$.

$$u_t \cdot u_{t-1} = (\rho u_{t-1} + \varepsilon_t) \cdot u_{t-1} = \rho u_{t-1}^2 + \varepsilon_t \cdot u_{t-1} \Rightarrow \mathbb{E}[u_t u_{t-1}] = \rho \sigma^2$$

De la misma forma:

$$\mathbb{E}[u_t u_{t-2}] = \mathbb{E}[(\rho u_{t-1} + \varepsilon_t) u_{t-2}] = \rho \mathbb{E}[u_{t-1} u_{t-2}] = \rho(\rho \sigma^2) = \rho^2 \sigma^2$$

$$\mathbb{E}[u_t u_{t-3}] = \mathbb{E}[(\rho u_{t-1} + \varepsilon_t) u_{t-3}] = \rho \mathbb{E}[u_{t-1} u_{t-3}] = \rho(\rho^2 \sigma^2) = \rho^3 \sigma^2$$

$$\Omega = \begin{pmatrix} \mathbb{E}[u_t u_{t-k}] = \rho^k \sigma^2 \\ 1 & \rho & \rho^2 & \rho^3 & \dots & \rho^{n-1} \\ \rho & 1 & \rho & \rho^2 & \dots & \rho^{n-2} \\ \rho^2 & \rho & 1 & \rho & \dots & \rho^{n-3} \\ & & & \ddots & & \\ \rho^{n-1} & \rho^{n-2} & \dots & \rho^2 & \rho & 1 \end{pmatrix}$$

$$\Rightarrow \Omega^{-1} = \frac{1}{1 - \rho^2} \begin{pmatrix} 1 & -\rho & 0 & 0 & \dots & 0 \\ -\rho & 1 + \rho^2 & -\rho & 0 & \dots & 0 \\ 0 & -\rho & 1 + \rho^2 & -\rho & \dots & 0 \\ 0 & 0 & \ddots & 1 + \rho^2 & -\rho & \\ 0 & 0 & \dots & -\rho & 1 \end{pmatrix}$$

Suponemos;

$$y_t = \beta_1 + \beta_2 X_{2t} + \dots + \beta_k X_{kt} + u_t$$

y

$$y_{t-1} = \beta_1 + \beta_2 X_{2t-1} + \dots + \beta_k X_{kt-1} + u_{t-1}$$

Luego:

$$y_t - \rho y_{t-1} = \beta_1(1 - \rho) + \beta_2(X_{2t} - \rho X_{2t-1}) + \dots + \beta_k(X_{kt} - \rho X_{kt-1}) + u_t - \rho u_{t-1}$$

Llamando $y_t^* = y_t - \rho y_{t-1}$ y $X_{it}^* = X_{it} - \rho X_{it-1}$:

$$y_t^* = \beta_1 + \beta_2 X_{2t}^* + \dots + \beta_k X_{kt}^* + \varepsilon_t$$

donde ahora ε cumple las hipótesis para aplicar MCO.

Nota: Se perdería la primera observación, para ello Prais-Winsten proponen

$$y_1^* = \sqrt{1 - \rho^2} y_1, X_{i1}^* = \sqrt{1 - \rho^2} X_{i1} \dots$$

MODIFICACIÓN PRAIS-WINSTEN

1. Estimando por MCO y se obtienen los residuos e_t .
2. Se estima $u_t = \rho u_{t-1} + \varepsilon_t$ el valor de ρ empleando los e_t :

$$\hat{\rho} = \frac{\sum_{t=2}^n e_t e_{t-1}}{\sum_{t=2}^n e_t^2}$$

3. Se estima por MCO $y^* = X^* \beta + u^*$: e_t^* .
4. Estimamos $\hat{\rho}$ de la regresión $e_t^* = \hat{\rho} e_{t-1}^* + \varepsilon_t^*$: $\hat{\rho} = \frac{\sum_{t=2}^n e_t^* e_{t-1}^*}{\sum_{t=2}^n (e_t^*)^2}$.
5. Repetimos... La estimación de ρ es el valor que estabiliza la secuencia $\hat{\rho}$, $\hat{\hat{\rho}}, \dots$ (hasta que la diferencia sea $< 10^{-3}$).

Proceso iterativo de Cochrane-Orcutt: Obviamos la primera observación en las transformaciones.

```
sm.GLSAR(y, sm.add_constant(x), rho=rho).iterative_fit(maxiter
```

Esto es mucho más fácil en la práctica que en la teoría (que no hay que saberse), de hecho es simplemente aplicar la fórmula y ya está. (Ejemplo en las diapositivas últimas del tema 6)