

Lecture 2: Properties of the OLS estimator

Economics 527: Econometric Methods

Vadim Marmer, UBC

The model and the assumptions

Where the first lecture ended

- The first lecture produced the estimator and stopped there:

$$\hat{\beta} = (X^\top X)^{-1} X^\top Y = \left(\sum_{i=1}^n X_i X_i^\top \right)^{-1} \sum_{i=1}^n X_i Y_i,$$

the unique minimiser of $\sum_{i=1}^n (Y_i - X_i^\top b)^2$.

- Nothing so far says whether it is a good estimator.
- $\hat{\beta}$ is a function of the sample, so it is random, and in general it will not equal β .
- Three questions, taken in this order: whether $\hat{\beta}$ is right on average, how far it moves from sample to sample, and whether a better estimator of the same kind exists.

The model in matrix form

- Stack the n observations, as in the first lecture:

$$\begin{matrix} Y \\ n \times 1 \end{matrix} = \begin{pmatrix} Y_1 \\ \vdots \\ Y_n \end{pmatrix}, \quad \begin{matrix} X \\ n \times k \end{matrix} = \begin{pmatrix} X_1^\top \\ \vdots \\ X_n^\top \end{pmatrix}, \quad \begin{matrix} U \\ n \times 1 \end{matrix} = \begin{pmatrix} U_1 \\ \vdots \\ U_n \end{pmatrix}.$$

- The n scalar equations $Y_i = X_i^\top \beta + U_i$ become one:

$$\begin{matrix} Y \\ n \times 1 \end{matrix} = \begin{matrix} X \\ n \times k \end{matrix} \begin{matrix} \beta \\ k \times 1 \end{matrix} + \begin{matrix} U \\ n \times 1 \end{matrix}.$$

- β is fixed and unknown. X , Y and U are random, and U is not observed.

The four assumptions

- **Assumption 1.** $Y = X\beta + U$ for some $\beta \in \mathbb{R}^k$.
- **Assumption 2.** $E[U \mid X] = 0$.
- **Assumption 3.** $\text{Var}(U \mid X) = \sigma^2 I_n$ for some $\sigma^2 > 0$.
- **Assumption 4.** $\text{rank}(X) = k$, so $n \geq k$ and no regressor is an exact linear combination of the others.
- The numbering is the typed notes'. Assumption 4 is the sample rank condition, the sample counterpart of $\text{rank}(E[X_i X_i^\top]) = k$ from the first lecture.
- Assumption 3 is new, and nothing before the variance of $\hat{\beta}$ uses it.

What unbiased means

- $\hat{\beta}$ is a function of the sample, hence a random vector:

$$\hat{\beta}(\omega) = \left(\sum_{i=1}^n X_i(\omega) X_i(\omega)^\top \right)^{-1} \sum_{i=1}^n X_i(\omega) Y_i(\omega).$$

- With a continuously distributed error, $P(\hat{\beta} = \beta) = 0$. Hitting β exactly is not the standard.
- **Definition.** An estimator $\hat{\theta}$ of θ is **unbiased** if $E[\hat{\theta}] = \theta$, and **conditionally unbiased given X** if

$$E[\hat{\theta} | X] = \theta,$$

whatever the true value of θ : the estimator is a rule that does not know θ , and the equality must hold for every value θ could take.

- The conditional statement is the stronger one, and it is the one this lecture proves.

Unbiasedness

The sampling error

- Substitute Assumption 1 into the estimator; Assumption 4 is what makes $(X^\top X)^{-1}$ exist:

$$\begin{aligned} \hat{\beta} &= (X^\top X)^{-1} X^\top Y \\ &= (X^\top X)^{-1} X^\top (X\beta + U) \\ &= \underbrace{(X^\top X)^{-1} X^\top X}_{=I_k} \beta + (X^\top X)^{-1} X^\top U. \end{aligned}$$

- The first term collapses, and what is left is the **sampling error**:

$$\boxed{\hat{\beta} = \underbrace{\beta}_{\text{signal}} + \underbrace{(X^\top X)^{-1} X^\top U}_{\text{noise}}}$$

- Everything in this section and the next is a statement about the noise term.

Taking expectations, first attempt

- Take expectations of the sampling error:

$$E[\hat{\beta}] = \beta + E[(X^\top X)^{-1} X^\top U].$$

- Suppose only that $E[X_i U_i] = 0$ for every i . Then

$$E[X^\top U] = E\left[\sum_{i=1}^n X_i U_i\right] = \sum_{i=1}^n E[X_i U_i] = 0.$$

- That is not enough, because in general

$$E[(X^\top X)^{-1} X^\top U] \text{ need not equal } E[(X^\top X)^{-1}] E[X^\top U].$$

- The two factors are built from the same X , and once $E[U | X]$ is not zero, the expectation of their product need not factor.

The expectation of a ratio

- The one-regressor case makes the failure visible:

$$\mathbb{E} \left[\frac{\sum_i X_i U_i}{\sum_i X_i^2} \right] \text{ need not equal } \frac{\mathbb{E} [\sum_i X_i U_i]}{\mathbb{E} [\sum_i X_i^2]} = 0.$$

- In general the expectation of a ratio is not the ratio of expectations. For any two scalar random variables Z and S ,

$$\mathbb{E} \left[\frac{Z}{S} \right] = \iint \frac{z}{s} f_{Z,S}(z, s) dz ds \text{ need not equal } \frac{\int z f_Z(z) dz}{\int s f_S(s) ds} = \frac{\mathbb{E}[Z]}{\mathbb{E}[S]}.$$

- The way out is to condition on X , which turns $(X^\top X)^{-1} X^\top$ into a known matrix.

Conditional unbiasedness: the claim

- **Claim.** Under Assumptions 1, 2 and 4,

$$\mathbb{E} [\hat{\beta} \mid X] = \beta.$$

- Assumption 3 is not used. Unbiasedness does not depend on the variance of the errors.

Proof of conditional unbiasedness

- Condition the sampling error on X :

$$\begin{aligned} \mathbb{E} [\hat{\beta} \mid X] &= \mathbb{E} [\beta + (X^\top X)^{-1} X^\top U \mid X] \\ &= \beta + \mathbb{E} [(X^\top X)^{-1} X^\top U \mid X] \\ &= \beta + (X^\top X)^{-1} X^\top \underbrace{\mathbb{E} [U \mid X]}_{=0 \text{ by Assumption 2}} \\ &= \beta. \end{aligned}$$

- The third line takes out what is known, as in the first lecture. Given X , the matrix $(X^\top X)^{-1} X^\top$ is known and comes out of the expectation.
- This is exactly the step the unconditional version could not take.

From conditional to unconditional

- The law of iterated expectations, from the first lecture, finishes the proof whenever $\mathbb{E} [\hat{\beta}]$ exists:

$$\mathbb{E} [\hat{\beta}] = \mathbb{E} [\mathbb{E} [\hat{\beta} \mid X]] = \mathbb{E} [\beta] = \beta.$$

- Conditional unbiasedness implies unbiasedness whenever the unconditional mean exists. The converse does not hold.
- Existence is not automatic: with $n = k = 1$, a regressor uniform on $(0, 1)$ and a mean-zero error independent of it, the noise term U/X has no mean, although its conditional mean is zero.

The scalar case

- Take $k = 1$, so X is $n \times 1$ and

$$X^\top X = \sum_{i=1}^n X_i^2, \quad X^\top Y = \sum_{i=1}^n X_i Y_i.$$

- The estimator and its sampling error:

$$\hat{\beta} = \frac{\sum_i X_i Y_i}{\sum_i X_i^2} = \beta + \frac{\sum_i X_i U_i}{\sum_i X_i^2}.$$

- Conditioning on $X_1, \dots, X_n$ makes the denominator and every X_i in the numerator known:

$$E[\hat{\beta} | X] = \beta + \frac{\sum_i X_i \overbrace{E[U_i | X]}^{=0}}{\sum_i X_i^2} = \beta.$$

A weaker condition

- Replace Assumption 2 by

$$E[X_i U_i] = 0 \quad \text{for all } i,$$

which, together with $E[U_i] = 0$, says that each regressor is uncorrelated with the error.

- The first lecture showed that β is still identified under this condition.
- It is not enough for unbiasedness:

$$E[\hat{\beta}] = \beta + \underbrace{E[(X^\top X)^{-1} X^\top U]}_{\text{need not be 0}}.$$

- So $\hat{\beta}$ need not be unbiased. The noise term is a nonlinear function of X and U , and once $E[U | X]$ is not zero, nothing forces its mean to vanish.

Zero correlation does not pin down the conditional mean

- Take one regressor, so X_i is a scalar with density f_X , and write $g(x) = E[U_i | X_i = x]$. By iterated expectations,

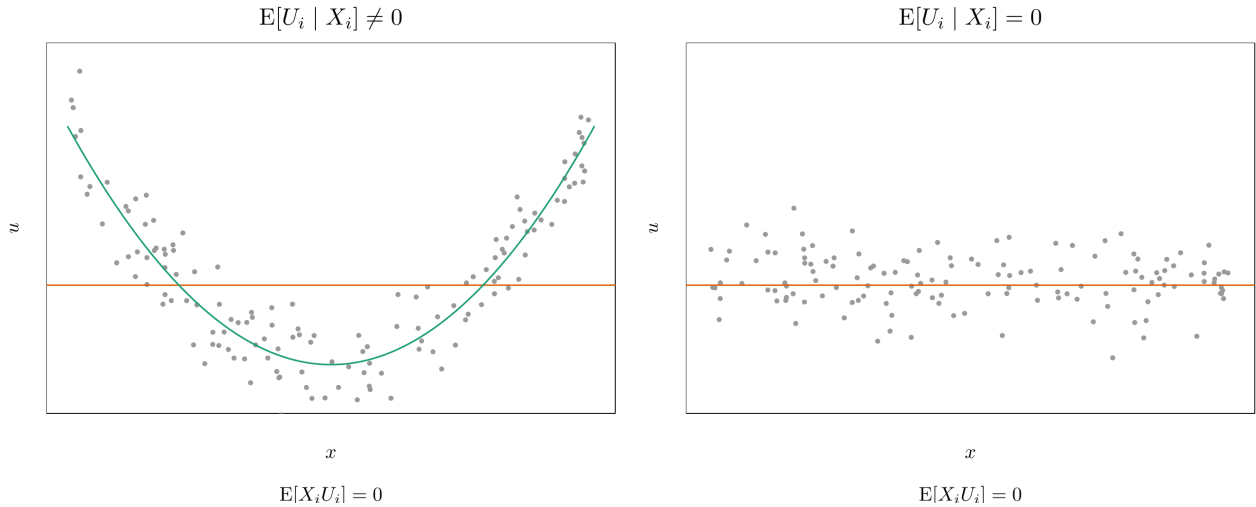
$$E[X_i U_i] = E[E[X_i U_i | X_i]] = E[X_i g(X_i)] = \int x g(x) f_X(x) dx.$$

- Assumption 2 says $g(x) = 0$ at every x , which makes the integral zero. The weaker condition says only that the integral is zero.
- An integral vanishes when the positive and negative parts cancel, so

$$\int x g(x) f_X(x) dx = 0 \quad \text{does not imply} \quad g(x) = 0 \quad \text{for all } x.$$

The parabola

- Let the conditional mean of the error be a parabola in the regressor, centred at zero and shifted down so that $E[U_i] = E[g(X_i)] = 0$, and let the regressor be symmetric about zero. Then g is even, so $x g(x)$ is odd and $E[X_i U_i] = 0$, while $E[U_i | X_i] \neq 0$ at every value of X_i except the two zeros of g .
- The parabola shows that the conditional mean is not zero. It does not by itself bias $\hat{\beta}$: with g even, the conditional mean of the noise term, $\sum_i X_i g(X_i) / \sum_i X_i^2$, changes sign when every X_i does. The regressors are symmetric about zero, so the noise term has mean zero. Bias needs one of those two symmetries broken, as on the next slide but one.
- Read the horizontal axis as education and the vertical axis as ability: ability is low at middling education and high at both ends.



What survives under the weaker condition

- Under $E[X_i U_i] = 0$ and $E[U_i] = 0$ in place of Assumption 2: β is identified, and $\hat{\beta}$ is still the method-of-moments and least-squares estimator.
- What is lost: $X_i^\top \beta$ need not be the conditional mean any more, and $\hat{\beta}$ is no longer guaranteed to be unbiased at any sample size.

The variance of a random vector

The definition

- For a random vector V that is $n \times 1$ with finite second moments, and with $E[V] = (E[V_1], \dots, E[V_n])^\top$,

$$\text{Var}(V) = E \left[\begin{matrix} (V - E[V]) \\ n \times 1 \end{matrix} \begin{matrix} (V - E[V])^\top \\ 1 \times n \end{matrix} \right] \quad (n \times n).$$

- The product inside is an outer product, not an inner product, so the result is a matrix rather than a number.

The matrix entry by entry

- Entry (i, j) is $E[(V_i - E[V_i])(V_j - E[V_j])]$, so

$$\text{Var}(V) = \begin{pmatrix} \text{Var}(V_1) & \text{Cov}(V_1, V_2) & \cdots & \text{Cov}(V_1, V_n) \\ \text{Cov}(V_2, V_1) & \text{Var}(V_2) & & \vdots \\ \vdots & & \ddots & \\ \text{Cov}(V_n, V_1) & \cdots & & \text{Var}(V_n) \end{pmatrix}.$$

- The variances sit on the diagonal and the covariances off it.

Symmetry

- $\text{Var}(V)^\top = \text{Var}(V)$, because $\text{Cov}(V_i, V_j) = \text{Cov}(V_j, V_i)$.
- The same statement from the definition, using that the transpose passes through the expectation entry by entry, and then $(AB)^\top = B^\top A^\top$:

$$\text{Var}(V)^\top = E \left[((V - E[V])(V - E[V])^\top)^\top \right] = E[(V - E[V])(V - E[V])^\top] = \text{Var}(V).$$

Positive semidefiniteness

- **Claim.** $\text{Var}(V)$ is positive semidefinite.
- **Proof.** Take any non-random $a \in \mathbb{R}^n$ and set $W = a^\top(V - \mathbb{E}[V])$, a scalar. Then

$$\begin{aligned} a^\top \text{Var}(V) a &= a^\top \mathbb{E}[(V - \mathbb{E}[V])(V - \mathbb{E}[V])^\top] a \\ &= \mathbb{E} \left[\underbrace{a^\top(V - \mathbb{E}[V])}_W \underbrace{(V - \mathbb{E}[V])^\top a}_W \right] \\ &= \mathbb{E}[W^2] \geq 0. \end{aligned}$$

- The constant a moves inside the expectation because expectation is linear.
- This is the argument the first lecture used for $\mathbb{E}[X_i X_i^\top]$, with $W_i = X_i^\top c$.

When it is positive definite

- For $a \neq 0$, with $W = a^\top(V - \mathbb{E}[V])$ as above, track the equality case:

$$a^\top \text{Var}(V) a = 0 \iff \mathbb{E}[W^2] = 0 \iff a^\top(V - \mathbb{E}[V]) = 0,$$

that is,

$$a_1(V_1 - \mathbb{E}[V_1]) + \dots + a_n(V_n - \mathbb{E}[V_n]) = 0.$$

- So $\text{Var}(V)$ is positive definite exactly when no non-zero linear combination of $V_1, \dots, V_n$ is constant, that is, when no exact linear relation holds among the centred variables $V_i - \mathbb{E}[V_i]$.

Variance under a linear map

- **Claim.** For a non-random $m \times n$ matrix Γ ,

$$\boxed{\text{Var}(\Gamma V) = \Gamma \text{Var}(V) \Gamma^\top} \quad (m \times m).$$

- **Proof.** $\mathbb{E}[\Gamma V] = \Gamma \mathbb{E}[V]$, so

$$\begin{aligned} \text{Var}(\Gamma V) &= \mathbb{E}[(\Gamma V - \Gamma \mathbb{E}[V])(\Gamma V - \Gamma \mathbb{E}[V])^\top] \\ &= \mathbb{E}[\Gamma(V - \mathbb{E}[V])(V - \mathbb{E}[V])^\top \Gamma^\top] \\ &= \Gamma \mathbb{E}[(V - \mathbb{E}[V])(V - \mathbb{E}[V])^\top] \Gamma^\top = \Gamma \text{Var}(V) \Gamma^\top, \end{aligned}$$

using $(AB)^\top = B^\top A^\top$ on the second factor, and linearity to move the non-random Γ and $\Gamma^\top$ outside the expectation.

- Check against the scalar case: with $m = n = 1$ this is $\text{Var}(\gamma V) = \gamma^2 \text{Var}(V)$.

Adding a constant

- For a non-random α that is $m \times 1$,

$$\text{Var}(\alpha + \Gamma V) = \text{Var}(\Gamma V) = \Gamma \text{Var}(V) \Gamma^\top.$$

- The constant shifts the mean, so the deviations from the mean are unchanged and it drops out.
- This is why β will drop out of the variance of the sampling error.

The conditional version

- The same definition with conditional expectations throughout:

$$\text{Var}(V | X) = E \left[(V - E[V | X])(V - E[V | X])^\top | X \right].$$

- Everything above holds conditionally, with Γ allowed to be a function of X . Given X , such a Γ is known and comes out of the conditional expectation, by taking out what is known, as in the proof of unbiasedness.
- When the conditional mean is zero, the definition simplifies. Under Assumption 2,

$$\text{Var}(U | X) = E[UU^\top | X].$$

The variance of the OLS estimator

The variance matrix of $\hat{\beta}$

- $\hat{\beta}$ is $k \times 1$, so $\text{Var}(\hat{\beta} | X)$ is $k \times k$: the conditional variance of each estimated coefficient on the diagonal, the conditional covariances between them off it.
- Conditioning on X again, for the same reason it worked for unbiasedness: it makes $(X^\top X)^{-1}X^\top$ a known matrix.

The sandwich formula

- Start from the sampling error and name the matrix that multiplies U :

$$\hat{\beta} = \beta + \underbrace{(X^\top X)^{-1}X^\top}_{\Gamma_X} U.$$

- β is a constant and drops out, and Γ_X is known given X , so the rule for a linear map applies:

$$\begin{aligned} \text{Var}(\hat{\beta} | X) &= \text{Var}(\beta + \Gamma_X U | X) \\ &= \text{Var}(\Gamma_X U | X) \\ &= \Gamma_X \text{Var}(U | X) \Gamma_X^\top. \end{aligned}$$

The sandwich formula written out

- Substituting $\Gamma_X = (X^\top X)^{-1}X^\top$, with $\Gamma_X^\top = X(X^\top X)^{-1}$ because $X^\top X$ and its inverse are symmetric:

$$\text{Var}(\hat{\beta} | X) = (X^\top X)^{-1}X^\top \text{Var}(U | X) X(X^\top X)^{-1}$$

- The box uses Assumptions 1 and 4, and needs the errors to have finite conditional second moments, so that $\text{Var}(U | X)$ exists. Assumption 2 enters only to write the middle factor as $E[UU^\top | X]$.
- Apart from requiring $\text{Var}(U | X)$ to exist, nothing yet has been assumed about the errors' conditional variances or covariances.

Assumption 3: homoskedasticity and no autocorrelation

- **Assumption 3.** $\text{Var}(U | X) = \sigma^2 I_n$ for some $\sigma^2 > 0$, that is,

$$\text{Var}(U | X) = \begin{pmatrix} \sigma^2 & & 0 \\ & \ddots & \\ 0 & & \sigma^2 \end{pmatrix}.$$

- Read the diagonal: every error has the same conditional variance,

$$\text{Var}(U_1 | X) = \dots = \text{Var}(U_n | X) = \sigma^2.$$

- This is **homoskedasticity**. It rules out errors that are more variable at some values of the regressors than at others.
- Read the off-diagonal entries: for $i \neq j$, with the first equality by Assumption 2,

$$\text{Cov}(U_i, U_j | X) = E[U_i U_j | X] = 0.$$

- The errors are conditionally uncorrelated across observations. This is **no autocorrelation**, also called no serial correlation.
- Assumption 3 is these two statements together, and nothing more. It does not say the errors are independent, and it does not say they are normal.

The variance of OLS under Assumption 3

- Put $\text{Var}(U | X) = \sigma^2 I_n$ into the sandwich:

$$\begin{aligned} \text{Var}(\hat{\beta} | X) &= (X^\top X)^{-1} X^\top \sigma^2 I_n X (X^\top X)^{-1} \\ &= \sigma^2 (X^\top X)^{-1} \underbrace{X^\top X (X^\top X)^{-1}}_{= I_k} \\ &= \sigma^2 (X^\top X)^{-1}. \end{aligned}$$

- Under Assumptions 1 to 4:

$$E[\hat{\beta} | X] = \beta, \quad \text{Var}(\hat{\beta} | X) = \sigma^2 (X^\top X)^{-1}$$

- This pair is the central result of the lecture.

What σ^2 is

- From Assumptions 2 and 3 together,

$$\sigma^2 = \text{Var}(U_i | X) = E[U_i^2 | X],$$

and the law of iterated expectations gives

$$\sigma^2 = E[U_i^2] = \text{Var}(U_i),$$

the last step because $E[U_i] = E[E[U_i | X]] = 0$ by Assumption 2.

- σ^2 is one unknown number, the same for every observation. Estimating it is a question for the next lecture.

One regressor

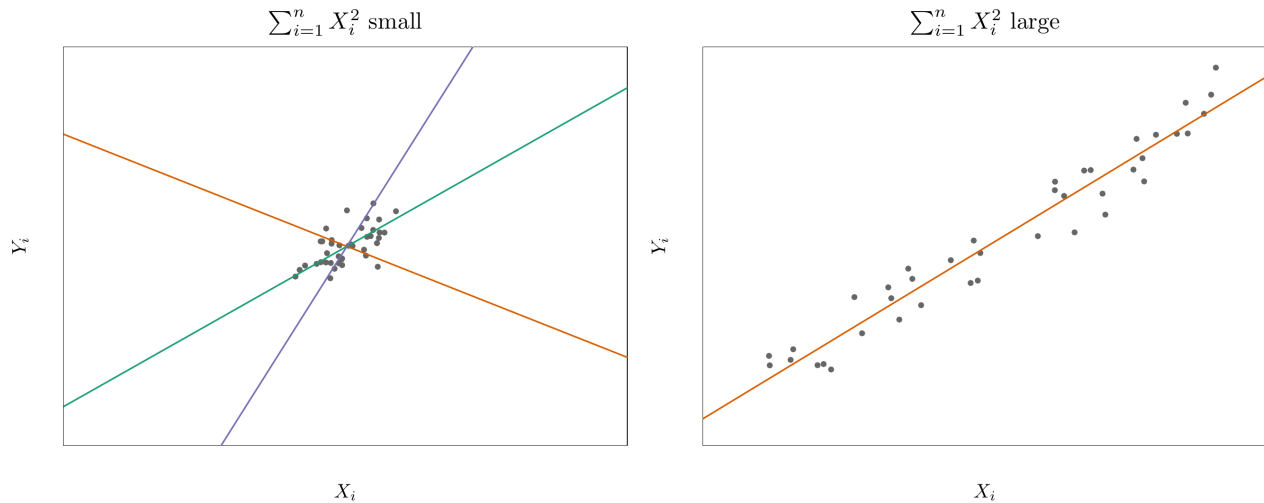
- Take $k = 1$ and no intercept, so $Y_i = \beta X_i + U_i$ and X is $n \times 1$. Then $X^\top X = \sum_i X_i^2$ and

$$\text{Var}(\hat{\beta} | X) = \frac{\sigma^2}{\sum_{i=1}^n X_i^2}.$$

- Two things move the variance. Noisier errors raise it. A larger $\sum_i X_i^2$ lowers it: with no intercept this is the sum of squares of the regressor about zero. With an intercept the same formula gives the variance of the slope, with $\sum_i (X_i - \bar{X})^2$, the spread of the regressor about its mean, in place of $\sum_i X_i^2$.

Why spread in the regressor helps

- With the regressor bunched together, many quite different slopes fit the data about equally well, so the estimate moves a lot from sample to sample.
- With the regressor spread out, the slope is pinned down.



Without homoskedasticity

- If the errors are heteroskedastic but still conditionally uncorrelated across observations, so that $\text{Var}(U | X)$ is diagonal with entries that may depend on X ,

$$\text{Var}(U | X) = \Omega = \begin{pmatrix} \sigma_1^2 & & 0 \\ & \ddots & \\ 0 & & \sigma_n^2 \end{pmatrix},$$

the sandwich does not collapse. The middle factor becomes

$$X^\top \Omega X = \begin{pmatrix} X_1 & \cdots & X_n \end{pmatrix} \Omega \begin{pmatrix} X_1^\top \\ \vdots \\ X_n^\top \end{pmatrix} = \sum_{i=1}^n \sigma_i^2 X_i X_i^\top,$$

so

$$\text{Var}(\hat{\beta} | X) = (X^\top X)^{-1} \left(\sum_{i=1}^n \sigma_i^2 X_i X_i^\top \right) (X^\top X)^{-1}.$$

- Unbiasedness is untouched. Only the variance formula changes.

Where each assumption was used

Assumptions	What they deliver
1, 2, 4	$E[\hat{\beta} X] = \beta$, and $E[\hat{\beta}] = \beta$ whenever that mean exists
1, 4, finite conditional second moments	$\text{Var}(\hat{\beta} X) = \Gamma_X \text{Var}(U X) \Gamma_X^\top$, with $\Gamma_X = (X^\top X)^{-1} X^\top$
1, 3, 4	$\text{Var}(\hat{\beta} X) = \sigma^2 (X^\top X)^{-1}$

The Gauss-Markov theorem

Linear estimators

- **Definition.** An estimator of β is **linear** if it can be written

$$\tilde{\beta} = AY$$

for a matrix A that is $k \times n$ and depends on X alone, not on Y .

- OLS is linear, with $A = (X^\top X)^{-1} X^\top$.
- Given X , each entry of a linear estimator is a fixed linear combination of the observations on the dependent variable.

Two more linear estimators

- **The two-point slope.** With $k = 1$, $Y_i = \beta X_i + U_i$ and $X_1 \neq X_n$,

$$\tilde{\beta} = \frac{Y_n - Y_1}{X_n - X_1} = -\frac{1}{X_n - X_1} Y_1 + \frac{1}{X_n - X_1} Y_n.$$

It uses two observations and throws the rest away.

- **Weighted least squares.** For a known $n \times n$ matrix W that is symmetric and positive definite, and that may depend on X but not on Y ,

$$\tilde{\beta} = (X^\top W X)^{-1} X^\top W Y.$$

OLS is the case $W = I_n$. With W diagonal, this estimator weights observation i by the i th diagonal entry w_i , as in the first lecture's slide on other linear estimators. The inverse $(X^\top W X)^{-1}$ exists: for $c \neq 0$, $Xc \neq 0$ by Assumption 4, so $c^\top X^\top W X c = (Xc)^\top W (Xc) > 0$.

Both are unbiased

- The two-point slope, substituting $Y_i = \beta X_i + U_i$ and using Assumption 2:

$$\tilde{\beta} = \beta + \frac{U_n - U_1}{X_n - X_1}, \quad \mathbb{E}[\tilde{\beta} | X] = \beta.$$

- Weighted least squares, substituting $Y = X\beta + U$ and using Assumption 2:

$$\tilde{\beta} = \beta + (X^\top W X)^{-1} X^\top W U, \quad \mathbb{E}[\tilde{\beta} | X] = \beta.$$

- Being unbiased does not single OLS out. When $n > k$, there are many conditionally unbiased linear estimators.

The variance of the weighted estimator

- The route that gave the sandwich formula applies, with $A = (X^\top W X)^{-1} X^\top W$ in place of Γ_X . Since $W^\top = W$, the matrix $X^\top W X$ and its inverse are symmetric too, so $A^\top = W X (X^\top W X)^{-1}$. Under Assumptions 1, 3 and 4,

$$\text{Var}(\tilde{\beta} | X) = \sigma^2 (X^\top W X)^{-1} X^\top W W X (X^\top W X)^{-1},$$

which reduces to $\sigma^2 (X^\top X)^{-1}$ when $W = I_n$.

- Which of these matrices is smallest, and what smallest means for matrices, are the next questions.

The comparison to settle

- For every linear conditionally unbiased $\tilde{\beta}$, compare the two conditional variance matrices:

$$\text{Var}_{k \times k}(\tilde{\beta} | X) - \text{Var}_{k \times k}(\hat{\beta} | X) \stackrel{?}{\geq} 0.$$

- Here ≥ 0 means positive semidefinite.

The Gauss-Markov theorem

- **Theorem.** Under Assumptions 1 to 4, let $\tilde{\beta}$ satisfy
 1. **linearity:** $\tilde{\beta} = AY$, with A a $k \times n$ matrix depending on X alone;
 2. **conditional unbiasedness:** $E[\tilde{\beta} | X] = \beta$ for every value of β .
- Then

$$\text{Var}(\tilde{\beta} | X) - \text{Var}(\hat{\beta} | X) \text{ is positive semidefinite.}$$

- OLS is the **best linear unbiased estimator** (BLUE): best in the sense above, linear in Y , conditionally unbiased, and an estimator of β .
- The theorem is about conditional unbiasedness, because it compares conditional variances.

Reading the matrix ordering

- If a $k \times k$ matrix M is positive semidefinite, then $c^\top M c \geq 0$ for **every** $c \in \mathbb{R}^k$.
- Take c to be the j th coordinate vector e_j , the vector with a 1 in position j and a 0 everywhere else:

$$e_j = (0, \dots, 0, 1, 0, \dots, 0)^\top, \quad \text{the 1 sitting in position } j.$$

- Writing the quadratic form out entry by entry, and using $(e_j)_l = 0$ for every $l \neq j$ and $(e_j)_j = 1$,

$$e_j^\top M e_j = \sum_{l=1}^k \sum_{m=1}^k (e_j)_l M_{lm} (e_j)_m = M_{jj}.$$

- So every diagonal entry of a positive semidefinite matrix is non-negative.
- Apply this to $M = \text{Var}(\tilde{\beta} | X) - \text{Var}(\hat{\beta} | X)$, whose j th diagonal entry is the difference of the two variances of the j th coefficient:

$$\text{Var}(\tilde{\beta}_j | X) \geq \text{Var}(\hat{\beta}_j | X), \quad j = 1, \dots, k.$$

- A general c says the same about every linear combination $c^\top \beta$ of the coefficients, since $\text{Var}(c^\top V | X) = c^\top \text{Var}(V | X) c$:

$$\text{Var}(c^\top \tilde{\beta} | X) \geq \text{Var}(c^\top \hat{\beta} | X) \quad \text{for every } c \in \mathbb{R}^k.$$

Proof, step 1: linearity and unbiasedness force $AX = I_k$

- Take conditional expectations of $\tilde{\beta} = AY$:

$$\begin{aligned} \beta &= E[\tilde{\beta} | X] = E[A(X\beta + U) | X] \\ &= AX\beta + A \underbrace{E[U | X]}_{=0 \text{ by Assumption 2}} \\ &= AX\beta. \end{aligned}$$

- This must hold for every value of β , which is what unbiasedness means, so

$$\boxed{AX = I_k}$$

- Linearity plus conditional unbiasedness is exactly this one condition.

Step 1 checked on the examples

- OLS: $A = (X^\top X)^{-1} X^\top$, so $AX = (X^\top X)^{-1} X^\top X = I_k$.
- Weighted least squares: $A = (X^\top W X)^{-1} X^\top W$, so $AX = (X^\top W X)^{-1} X^\top W X = I_k$.
- The two-point slope, with $k = 1$: A has $-1/(X_n - X_1)$ in position 1 and $1/(X_n - X_1)$ in position n , so

$$AX = \frac{X_n - X_1}{X_n - X_1} = 1.$$

Proof, step 2: the two sampling errors

- Using $AX = I_k$ for the general estimator, and the sampling error for OLS,

$$\hat{\beta} = \beta + (X^\top X)^{-1} X^\top U, \quad \tilde{\beta} = AX\beta + AU = \beta + AU.$$

- Both estimators differ from β only through a known matrix times U . Everything that follows is an expectation of $UU^\top$.

The transpose rule for covariance matrices

- For a $k \times 1$ random vector V and an $\ell \times 1$ random vector Z ,

$$\text{Cov}(V, Z) = E[(V - E[V])(Z - E[Z])^\top] \quad (k \times \ell),$$

and transposing gives

$$\text{Cov}(V, Z)^\top = E[(Z - E[Z])(V - E[V])^\top] = \text{Cov}(Z, V).$$

- The same definition and the same rule hold conditionally, with $E[\cdot]$ replaced by $E[\cdot | X]$ throughout. Step 4 needs this, because both orders appear there.

Proof, step 3: the covariance equals the variance of OLS

- Both estimators are conditionally unbiased, so their deviations from their conditional means are the two noise terms of step 2:

$$\begin{aligned} \text{Cov}(\tilde{\beta}, \hat{\beta} | X) &= E[(\tilde{\beta} - \beta)(\hat{\beta} - \beta)^\top | X] \\ &= E[AU U^\top X (X^\top X)^{-1} | X] \\ &= A \underbrace{E[UU^\top | X]}_{= \sigma^2 I_n \text{ by Assumptions 2 and 3}} X (X^\top X)^{-1} \\ &= \sigma^2 \underbrace{AX}_{= I_k} (X^\top X)^{-1} \\ &= \sigma^2 (X^\top X)^{-1}. \end{aligned}$$

- The right-hand side is $\text{Var}(\hat{\beta} | X)$, so

$$\boxed{\text{Cov}(\tilde{\beta}, \hat{\beta} | X) = \text{Var}(\hat{\beta} | X)}$$

- Every linear conditionally unbiased estimator has the same covariance with OLS, namely the variance of OLS itself. This is the one surprising step in the proof.

Proof, step 4: the difference of the variances

- As for a scalar, expand the variance of the difference using the definition; the cross terms are the two covariance matrices:

$$\begin{aligned}\text{Var}(\tilde{\beta} - \hat{\beta} \mid X) &= \text{Var}(\tilde{\beta} \mid X) + \text{Var}(\hat{\beta} \mid X) \\ &\quad - \text{Cov}(\tilde{\beta}, \hat{\beta} \mid X) - \text{Cov}(\hat{\beta}, \tilde{\beta} \mid X).\end{aligned}$$

- Step 3 gives the first covariance term, the second is its transpose by the rule above, and $\text{Var}(\hat{\beta} \mid X)$ is symmetric. All three terms in blue equal $\text{Var}(\hat{\beta} \mid X)$, and they combine into one with a minus sign:

$$\text{Var}(\tilde{\beta} - \hat{\beta} \mid X) = \text{Var}(\tilde{\beta} \mid X) - \text{Var}(\hat{\beta} \mid X).$$

- The left-hand side is a conditional variance matrix, hence positive semidefinite. Therefore

$$\text{Var}(\tilde{\beta} \mid X) - \text{Var}(\hat{\beta} \mid X) \geq 0,$$

which is the theorem.

Uniqueness

- Suppose some linear conditionally unbiased $\tilde{\beta}$ attains the same variance matrix. By the identity of step 4,

$$\text{Var}(\tilde{\beta} \mid X) = \text{Var}(\hat{\beta} \mid X) \implies \text{Var}(\tilde{\beta} - \hat{\beta} \mid X) = 0.$$

- A random vector with zero conditional variance equals its conditional mean (the equality case of positive semidefiniteness, one coordinate at a time), so

$$\tilde{\beta} - \hat{\beta} = d(X)$$

for some function of X alone.

- Both are conditionally unbiased, so taking conditional expectations,

$$\underbrace{\text{E}[\tilde{\beta} \mid X]}_{=\beta} = \underbrace{\text{E}[\hat{\beta} \mid X]}_{=\beta} + d(X) \implies d(X) = 0.$$

- Therefore $\tilde{\beta} = \hat{\beta}$. A linear conditionally unbiased competitor that is also best attains the same variance matrix: both $M = \text{Var}(\tilde{\beta} \mid X) - \text{Var}(\hat{\beta} \mid X)$ and $-M$ are then positive semidefinite, so $c^\top M c = 0$ for every c , and a symmetric matrix with that property is zero. So OLS is not just a best linear unbiased estimator; it is the only one.

Where Assumption 3 was used

- Assumption 3 entered in step 3, where $\text{E}[UU^\top \mid X]$ was replaced by $\sigma^2 I_n$, and again when $\sigma^2(X^\top X)^{-1}$ was recognised as $\text{Var}(\hat{\beta} \mid X)$. Without it the cancellation that produced $\sigma^2 AX$ fails, and the covariance need not equal the variance of OLS.
- Assumptions 1, 2 and 4 give unbiasedness, and Assumptions 1 and 4 with finite conditional second moments give the sandwich formula. Assumption 3 is what makes OLS efficient.
- If Assumption 3 fails, OLS is still unbiased but need not be best: Assumption 3 is sufficient for efficiency, not necessary.

Two extensions, and the summary

Adding normality

- **Assumption 5.** $U | X \sim N(0, \sigma^2 I_n)$. Assumptions 1 to 5 define the classical normal regression model.
- Given X , $\hat{\beta} = \beta + (X^\top X)^{-1} X^\top U$ is a constant plus a linear function of U , and Assumption 5 fixes the conditional distribution of U . For non-random α and Γ , such a function of a normal vector is normal:

$$V \sim N(\mu, \Sigma) \implies \alpha + \Gamma V \sim N(\alpha + \Gamma \mu, \Gamma \Sigma \Gamma^\top).$$

- Under Assumptions 1, 4 and 5, since Assumption 5 implies Assumptions 2 and 3, the conditional mean and variance are the ones already computed:

$$\boxed{\hat{\beta} | X \sim N(\beta, \sigma^2 (X^\top X)^{-1})}$$

- Nothing before this slide needed normality. It is what confidence intervals and tests will be built on.

If Assumption 3 fails: generalised least squares

- Suppose now that $\text{Var}(U | X) = \Omega$ for a known positive definite Ω , not necessarily diagonal, with $\Omega \neq \sigma^2 I_n$ for every $\sigma^2 > 0$. Let $\Omega^{-1/2}$ be the symmetric positive definite matrix with $\Omega^{-1/2} \Omega^{-1/2} = \Omega^{-1}$, built from the eigenvalue decomposition $\Omega = C \Lambda C^\top$, and write $\Omega^{1/2}$ for the inverse of $\Omega^{-1/2}$.
- Multiply the model through by $\Omega^{-1/2}$:

$$\underbrace{\Omega^{-1/2} Y}_{Y^*} = \underbrace{\Omega^{-1/2} X}_{X^*} \beta + \underbrace{\Omega^{-1/2} U}_{U^*}, \quad \mathbb{E}[U^* | X] = 0, \quad \text{Var}(U^* | X) = \Omega^{-1/2} \Omega \Omega^{-1/2} = I_n,$$

by taking out what is known and by the rule for a linear map.

- The transformed model satisfies all four assumptions (X^* has rank k because $\Omega^{-1/2}$ is nonsingular), so Gauss-Markov applies to it. Since $A \mapsto A \Omega^{1/2}$ matches the linear conditionally unbiased estimators of the two models one to one, OLS on the transformed model is also best among estimators linear in Y and conditionally unbiased, and it is

$$\hat{\beta}^* = (X^{*\top} X^*)^{-1} X^{*\top} Y^* = (X^\top \Omega^{-1} X)^{-1} X^\top \Omega^{-1} Y,$$

which is weighted least squares with $W = \Omega^{-1}$.

- In practice Ω is not known. Later lectures return to this.

Summary: unbiasedness and variance

- Substituting the model into the estimator gives the sampling error $\hat{\beta} = \beta + (X^\top X)^{-1} X^\top U$. Everything follows from it.
- **Unbiasedness.** Under Assumptions 1, 2 and 4, conditioning on X makes $(X^\top X)^{-1} X^\top$ known, so $\mathbb{E}[\hat{\beta} | X] = \beta$, and the law of iterated expectations gives $\mathbb{E}[\hat{\beta}] = \beta$ whenever that mean exists. The weaker condition $\mathbb{E}[X_i U_i] = 0$ is not enough, because once $\mathbb{E}[U | X]$ is not zero, the expectation of the noise term need not vanish, and the parabola shows that zero correlation does not make the conditional mean zero.
- **Variance.** Under Assumptions 1 and 4, given finite conditional second moments,

$$\text{Var}(\hat{\beta} | X) = (X^\top X)^{-1} X^\top \text{Var}(U | X) X (X^\top X)^{-1},$$

and adding Assumption 3, that $\text{Var}(U | X) = \sigma^2 I_n$, gives

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

With one regressor and no intercept this is $\sigma^2 / \sum_i X_i^2$: noisier errors raise it; a larger $\sum_i X_i^2$ lowers it.

Summary: efficiency

- **Gauss-Markov.** Under Assumptions 1 to 4, OLS has the smallest conditional variance among all estimators that are linear in Y and conditionally unbiased, whatever the value of β . For every such $\tilde{\beta}$, the difference of the two conditional variance matrices is positive semidefinite:

$$\text{Var}(\tilde{\beta} | X) - \text{Var}(\hat{\beta} | X) \geq 0,$$

so in particular $\text{Var}(\tilde{\beta}_j | X) \geq \text{Var}(\hat{\beta}_j | X)$ for every coefficient.

- **The proof.** Linearity and conditional unbiasedness force $AX = I_k$. Assumption 3 then makes $\text{Cov}(\tilde{\beta}, \hat{\beta} | X) = \text{Var}(\hat{\beta} | X)$. Expanding $\text{Var}(\tilde{\beta} - \hat{\beta} | X)$, which is positive semidefinite, leaves exactly the difference of the two variances.
- **Uniqueness.** Equality of the variances forces $\tilde{\beta} = \hat{\beta}$.
- Assumption 3 is what buys efficiency. Assumptions 1, 2 and 4 give unbiasedness on their own.