

FINANCIAL ECONOMETRICS 2026
SHANGHAI JIAO TONG UNIVERSITY
SUMMER SCHOOL 2026
QUIZ 1
With Solutions

Reading time: 10 minutes. **Writing time:** 60 minutes.

Total marks: 100. **Calculators are NOT permitted.** Every calculation is designed to be done by hand.

Instructions: Answer **all** questions. In Section A, circle the letter of the single best answer. In Sections B and C, show your working and write in clear sentences.

Structure:

- **Section A** — Multiple choice: 20 questions, 2.5 marks each (**50 marks**).
- **Section B** — One analytical problem in five parts (**30 marks**).
- **Section C** — One short interpretation (**20 marks**).

Useful facts. For a 95% confidence interval use 1.96 (for 90% use 1.64, for 99% use 2.58). $SE(\bar{Y}) = s_Y/\sqrt{n}$; $t = \frac{\text{estimate} - \text{null}}{SE}$; $\hat{\beta}_1 = \frac{s_{XY}}{s_X^2} = r \frac{s_Y}{s_X}$; $\hat{\beta}_0 = \bar{Y} - \hat{\beta}_1 \bar{X}$;

$r = \frac{s_{XY}}{s_X s_Y}$; $R^2 = \frac{ESS}{TSS}$ (with one regressor, $R^2 = r^2$).

Section A: Multiple Choice (20 questions $\times$ 2.5 marks = 50 marks)

1. An analyst has the monthly return of the S&P 500 index for every month from January 1990 to December 2025. This data set is an example of:

- (A) cross-sectional data.
- (B) time-series data.
- (C) panel (longitudinal) data.
- (D) experimental data.

Solution. (B). One entity (the index) is followed over many time periods, which is time-series data. Cross-sectional data would be many entities at a single date (say 500 stocks on one day); panel data would follow many entities over several periods (say 48 funds each over 10 years).

2. You lend a friend money. With probability 0.95 you are repaid \$110; with probability 0.05 the friend defaults and you receive \$0. The expected repayment is:

- (A) \$110.
- (B) \$105.
- (C) \$104.50.
- (D) \$100.

Solution. (C). The expected value is the probability-weighted average: $\$110 \times 0.95 + \$0 \times 0.05 = \$104.50$. It need not equal any outcome that can actually occur; it is the long-run average over many such loans.

3. In finance, the standard deviation of an asset's returns is used as the standard measure of:

- (A) its risk, also called its volatility.
- (B) its expected return.
- (C) its correlation with the market.
- (D) its average price.

Solution. (A). The standard deviation of returns measures how widely outcomes spread around the mean, which is exactly what we mean by risk; in finance it has the special name volatility. The mean measures the expected return, a different thing.

4. Monthly returns on a portfolio are approximately normal with mean 0% and standard deviation 5%. About 95% of monthly returns fall between:

- (A) -5% and 5% .
- (B) -2.5% and 2.5% .
- (C) -15% and 15% .
- (D) -10% and 10% .

Solution. (D). For a normal distribution, 95% of the probability lies within 1.96 standard deviations of the mean: $0 \pm 1.96 \times 5\% \approx \pm 9.8\%$, or roughly -10% to $+10\%$. Option (A) is the ± 1 standard-deviation band (about 68%), and (C) is roughly ± 3 standard deviations.

5. Which statement about the *correlation* between two random variables is correct?

- (A) A correlation of zero proves that the two variables are independent.
- (B) It always lies between -1 and $+1$ and measures the strength of their *linear* relationship.
- (C) It has the same units as X multiplied by Y .
- (D) It can exceed $+1$ when the relationship is very strong.

Solution. (B). Correlation is covariance divided by the two standard deviations, so it is unit-free and bounded in $[-1, +1]$, and it captures only *linear* association. Zero correlation does *not* prove independence (a strong curved pattern can have $r \approx 0$); having units is the drawback of covariance, not correlation.

6. An investor puts half of her money in each of two stocks that have the same volatility. Combining them lowers the portfolio's risk below that of either stock alone, provided the two returns are:

- (A) less than perfectly positively correlated (correlation below $+1$).
- (B) perfectly positively correlated (correlation equal to $+1$).
- (C) equal in expected return.
- (D) both riskier than the market.

Solution. (A). The variance of a sum contains a covariance term, so risk does not simply add. As long as the two returns do not move in perfect lockstep ($\rho < 1$), some of their fluctuations offset, and the portfolio's volatility falls below the individual volatility. If $\rho = +1$ there is no diversification benefit at all.

7. Individual daily stock returns have fat tails and are clearly not normal. Yet we can still use the value 1.96 to build a confidence interval for an *average* return computed from many observations. This is justified by:

- (A) the assumption that individual returns are normal.
- (B) the law of large numbers, which makes the sample mean exactly equal to the population mean.
- (C) the fact that 1.96 is valid for every distribution and every sample size.
- (D) the central limit theorem: the sample average is approximately normal in large samples, whatever the shape of the underlying data.

Solution. (D). The central limit theorem says the standardised sample average approaches a standard normal as n grows, *regardless* of the parent distribution, which is why 1.96 works for the average even though single returns are not normal. The law of large numbers is about the mean settling near the truth (consistency), not about the shape, and it never makes the estimate *exactly* equal to the truth.

8. The sample mean $\bar{Y} = \frac{1}{n} \sum_i Y_i$, seen as a rule applied to a random sample, is best described as:

- (A) a fixed number, known before the data are drawn.
- (B) equal to the population mean μ_Y .
- (C) a random variable, because a different sample would give a different value.
- (D) the same thing as the standard error.

Solution. (C). Because the sample is random, the estimator $\bar{Y}$ is itself a random variable: before you draw the data it could come out many ways. Its distribution across all possible samples is the sampling distribution. It is centred on μ_Y but is not equal to it in any one sample.

9. A fund database includes only funds that are still operating today; funds that closed (usually after poor performance) were removed. Using it to estimate the average return of *all funds ever launched* gives an estimate that is:

- (A) biased upward, and collecting more surviving funds will not remove the bias.
- (B) unbiased, because the sample is large.
- (C) biased downward.
- (D) unbiased, provided returns are normally distributed.

Solution. (A). This is survivorship bias: the sample is not a random draw from the intended population, because the worst performers have been dropped, so the average of the survivors is too high. Bias is a systematically wrong target; adding more surviving funds reduces noise but does not fix a sample that is aimed at the wrong thing.

10. The standard error of the sample mean, $SE(\bar{Y}) = s_Y/\sqrt{n}$, measures:

- (A) the spread of the individual observations around the sample mean.
- (B) roughly how far the estimate $\bar{Y}$ typically lands from the true mean μ_Y .
- (C) the probability that the null hypothesis is true.
- (D) the fraction of the variation in Y explained by a model.

Solution. (B). The standard error is the estimated standard deviation of the estimator $\bar{Y}$: how far your one estimate typically falls from the truth. The spread of the individual observations is s_Y (option A), which is a different quantity; the standard error is s_Y shrunk by $\sqrt{n}$.

11. A two-sided test of $H_0 : \mu = 0$ at the 5% level gives a t -statistic of 1.2. The correct conclusion is:

- (A) reject H_0 : the mean is not zero.
- (B) accept H_0 : the mean is exactly zero.
- (C) fail to reject H_0 : the data are consistent with $\mu = 0$, though this does not prove $\mu = 0$.
- (D) the test cannot be carried out with a positive t .

Solution. (C). Since $|1.2| < 1.96$, the estimate is within the ordinary range of sampling noise, so we fail to reject H_0 at 5%. We never “accept” the null: failing to reject is an acquittal for lack of evidence, not proof that the mean is exactly zero.

12. A strategy’s average monthly return over $n = 64$ months is 0.8%, with a sample standard deviation $s = 3.2\%$. Using $SE = s/\sqrt{n}$, the t -statistic for $H_0 : \mu = 0$ is:

- (A) 0.25.
- (B) 6.4.
- (C) 16.
- (D) 2.0.

Solution. (D). $SE = 3.2/\sqrt{64} = 3.2/8 = 0.4$, so $t = (0.8 - 0)/0.4 = 2.0$. Since $2.0 > 1.96$, the edge is (just) statistically significant at 5%. Option (A) forgets the $\sqrt{n}$ (using $0.8/3.2$); the others misuse the standard error.

13. A regression coefficient has a p -value of 0.03. This means:

- (A) there is a 3% probability that the null hypothesis is true.
- (B) if the null hypothesis were true, a result at least this extreme would occur about 3 times in 100 by chance alone.
- (C) the coefficient explains 3% of the variation in the outcome.
- (D) there is a 97% probability that the estimate equals the true value.

Solution. (B). The p -value is the probability of a result at least this extreme *given* that H_0 is true. It is *not* the probability that H_0 is true (option A), which is the single most common misreading; a small p -value is evidence against H_0 , so at $0.03 < 0.05$ we would reject at the 5% level.

14. A 95% confidence interval for a fund's mean monthly return is [0.1%, 0.9%]. From this alone, a two-sided 5% test of $H_0 : \mu = 0$ would:

- (A) reject H_0 , because 0 lies outside the interval.
- (B) fail to reject H_0 , because the interval is narrow.
- (C) fail to reject H_0 , because the interval contains only positive values.
- (D) be impossible to decide from the interval alone.

Solution. (A). A 95% confidence interval is exactly the set of null values a two-sided 5% test would *not* reject. Since 0 lies outside [0.1%, 0.9%], the test rejects $H_0 : \mu = 0$. This equivalence between intervals and tests is the duality from Lecture 2.

15. Ordinary least squares (OLS) chooses the intercept and slope that make which quantity as small as possible?

- (A) the sum of the residuals.

- (B) the largest single residual.
- (C) the sum of the absolute values of the residuals.
- (D) the sum of the *squared* residuals.

Solution. (D). OLS minimises $\sum_i (Y_i - b_0 - b_1 X_i)^2$. We square rather than simply add the residuals because positive and negative misses would otherwise cancel (many bad lines give a plain sum of exactly zero), and squaring makes large misses count much more.

16. In a regression of a stock's return Y on the market's return X , the sample covariance of X and Y is 12 and the sample variance of X is 8. The estimated slope (the stock's beta) is:

- (A) 0.67.
- (B) 3.
- (C) 1.5.
- (D) 96.

Solution. (C). $\hat{\beta}_1 = s_{XY}/s_X^2 = 12/8 = 1.5$. Option (A) inverts the ratio (8/12); (D) multiplies instead of dividing. A slope (beta) above 1 means the stock tends to move more than one-for-one with the market.

17. A fitted regression line is $\hat{Y} = 2 + 0.5X$. For an observation with $X = 10$ and an actual value $Y = 9$, the residual (actual minus predicted) is:

- (A) +9.
- (B) +2.
- (C) -2.
- (D) +7.

Solution. (B). The fitted value is $\hat{Y} = 2 + 0.5 \times 10 = 7$, so the residual is $Y - \hat{Y} = 9 - 7 = +2$: the point sits two units above the line. Option (D), 7, is the fitted value itself, not the residual.

18. In a regression of test scores on class size across many school districts, $R^2 = 0.05$. The best interpretation is:

- (A) class size accounts for about 5% of the variation in scores across districts; other factors account for the rest.

- (B) the regression has failed, because a good model needs an R^2 close to 1.
- (C) 5% of the slope estimate is due to chance.
- (D) the slope must be statistically insignificant.

Solution. (A). R^2 is the share of the variation in Y that the line explains, here about 5%. A low R^2 is normal, not a failure: it simply says many other things also affect scores. It does not by itself tell us whether the slope is precisely estimated or statistically significant.

19. OLS gives a slope of -2.28 for a regression of test scores on class size. For this slope to be read as the *causal* effect of class size, rather than a mere association, the key requirement is that:

- (A) R^2 is close to 1.
- (B) the sample is large enough that the t -statistic exceeds 1.96.
- (C) the two variables are strongly correlated.
- (D) other factors in the error term (such as district income) do not move systematically with class size, that is, $\mathbb{E}(u \mid X) = 0$.

Solution. (D). The causal reading needs the first least-squares assumption, $\mathbb{E}(u \mid X) = 0$: the omitted factors in u must not move with X . Here they plausibly do (richer districts have both smaller classes and higher scores), so the slope may credit class size with gains that are really due to income. A large t or a high R^2 makes the association more precise, not more causal.

20. A stock's return is regressed on the market's return and the estimated slope (its beta) is 1.4. On average, when the market rises by 1%, this stock:

- (A) rises by exactly 1%.
- (B) falls by 1.4%.
- (C) rises by about 1.4%, so it swings more than the market.
- (D) is unaffected by the market.

Solution. (C). The slope is the average change in the stock's return per one-unit change in the market's return, so a beta of 1.4 means the stock moves about 1.4% for each 1% market move: it is more volatile than the market. A beta below 1 would mean it swings less; a beta near 0 would mean little market sensitivity.

Section B: Analytical Problem (30 marks)

You have $n = 60$ months of data on a stock. For each month you observe the stock's excess return Y (in %) and the market's excess return X (in %). The summary statistics are:

$$\bar{X} = 0.5, \quad \bar{Y} = 1.0, \quad s_X^2 = 16, \quad s_{XY} = 24, \quad s_Y^2 = 100.$$

You will fit the “market model” regression $Y = \beta_0 + \beta_1 X + u$ by OLS. (Here s_X^2 and s_Y^2 are the sample variances and s_{XY} the sample covariance.)

Useful roots: $\sqrt{16} = 4$ and $\sqrt{100} = 10$. No calculator is needed anywhere in this question.

- (a) [6 marks] Compute the OLS slope $\hat{\beta}_1$ (the stock's *beta*) and the intercept $\hat{\beta}_0$ (its *alpha*). In one sentence, interpret the slope.
- (b) [6 marks] Compute the correlation r between X and Y , and the R^2 of the regression. What fraction of the variation in the stock's return does the market explain, and what does the rest represent?
- (c) [6 marks] Use the fitted line to predict the stock's excess return in a month when the market's excess return is $X = 2.5\%$. If the stock actually returned 5.0% that month, compute the residual.
- (d) [6 marks] The standard error of the slope is $SE(\hat{\beta}_1) = 0.25$. Test $H_0 : \beta_1 = 1$ (the stock is exactly as risky as the market) against a two-sided alternative at the 5% level, and give the 95% confidence interval for beta.
- (e) [6 marks] The estimated alpha is $\hat{\beta}_0 = 0.25\%$ per month. Explain what a *positive* alpha would mean for an investor, and why you must test whether it is statistically different from zero before concluding that the manager has genuine “skill”.

Solution. (a) $\hat{\beta}_1 = s_{XY}/s_X^2 = 24/16 = \mathbf{1.5}$ and $\hat{\beta}_0 = \bar{Y} - \hat{\beta}_1\bar{X} = 1.0 - 1.5(0.5) = 1.0 - 0.75 = \mathbf{0.25}$. Interpretation: when the market's excess return is 1 percentage point higher, this stock's excess return is on average 1.5 points higher, so with a beta of 1.5 it swings more than the market.

(b) $s_X = \sqrt{16} = 4$ and $s_Y = \sqrt{100} = 10$, so $r = s_{XY}/(s_X s_Y) = 24/(4 \times 10) = \mathbf{0.6}$, and $R^2 = r^2 = \mathbf{0.36}$. The market explains about 36% of the variation in the stock's return (its systematic, or market, risk); the remaining 64% is stock-specific (idiosyncratic) variation that diversification can reduce.

(c) $\hat{Y} = 0.25 + 1.5(2.5) = 0.25 + 3.75 = \mathbf{4.0\%}$. The residual is actual minus predicted, $5.0 - 4.0 = \mathbf{+1.0\%}$: that month the stock did 1 percentage point better than its market exposure alone predicts.

(d) $t = \frac{\hat{\beta}_1 - 1}{SE(\hat{\beta}_1)} = \frac{1.5 - 1}{0.25} = \mathbf{2.0}$. Since $2.0 > 1.96$, we **reject** $H_0 : \beta_1 = 1$ at the 5% level: the stock is significantly more volatile than the market. The 95% confidence interval for beta is $\hat{\beta}_1 \pm 1.96 SE = 1.5 \pm 1.96(0.25) = 1.5 \pm 0.49 = [\mathbf{1.01}, \mathbf{1.99}]$, which excludes 1, the same conclusion.

(e) Alpha is the average return the stock earns *beyond* what its market exposure (beta) would justify. A positive alpha therefore looks like genuine outperformance, or manager “skill”. But the estimated 0.25% is only one noisy estimate, so before claiming skill you must test $H_0 : \alpha = 0$: a small positive alpha could easily be luck (sampling variation), and only if it is both statistically significant and economically large after costs is “skill” a credible conclusion. This is the skill-versus-luck problem from Lecture 2.

Section C: Interpretation (20 marks)

Read the following and answer in **no more than five sentences**. Use the concepts from the course; you do not need any calculation.

*A financial analyst collects one year of data on 200 companies. For each firm she records its stock return over the year (Y) and the number of times the firm was mentioned in the financial press that year (X). She regresses Y on X , finds a positive slope that is highly statistically significant ($t = 4.5$, $p < 0.001$), and concludes that generating more press coverage **causes** a company's stock to rise. She advises firms to buy more media attention in order to lift their share price.*

Question [20 marks]. Explain why this significant positive slope does *not* establish that press coverage causes higher returns. Name the key least-squares assumption a causal interpretation requires, and give at least one concrete reason it is likely to fail here. Answer in at most five sentences.

Solution. Model answer (correlation is not causation). A regression slope, however statistically significant, measures *association*, not causation; a large t -statistic only tells us the relationship is unlikely to be exactly zero by chance, not that X causes Y . A causal reading requires the first least-squares assumption, $\mathbb{E}(u \mid X) = 0$: the regressor must be uncorrelated with the other factors in the error term. That almost certainly fails here. *Reverse causality*: firms whose stock is already rising attract more press coverage, so Y drives X rather than X driving Y . *Omitted variable (a confounder)*: larger, faster-growing or more profitable firms both receive more coverage *and* earn higher returns, so an unmeasured factor (firm size or performance) drives both X and Y ; to establish causation you would need a randomised experiment or a way to hold such confounders constant.

Marking guide (20 marks).

- **8 marks** — states clearly that a significant slope shows **association / correlation, not causation** (statistical significance is not a causal effect).
- **6 marks** — names the required assumption $\mathbb{E}(u \mid X) = 0$ (exogeneity: the regressor uncorrelated with the error / no omitted confounders).
- **4 marks** — gives at least one concrete failure: a **confounder / omitted variable** (firm size or performance) **or reverse causality** (returns attract coverage).
- **2 marks** — notes what *would* support causation (a randomised experiment, controls, or holding confounders fixed) or writes clearly within the five-sentence limit.

Answers that stop at “the relationship is significant, so coverage works” earn little: significance addresses chance, not confounding.