

Exam II: Mathematical Foundations of Risk Measurement - 2015 Edition

PRMIA 8007

Version Demo

Total Demo Questions: 15

Total Premium Questions: 159

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

QUESTION NO: 1

A 2-year bond has a yield of 5% and an annual coupon of 5%. What is the Modified Duration of the bond?

- A. 2
- B. 1.95
- C. 1.86
- D. 1.75

ANSWER: C

Explanation:

1.86 is correct. Since the bond's annual coupon rate equals its annual yield, the bond is priced at par. Using a face value of 100, its cash flows are 5 at the end of year one and 105 at the end of year two. Discounting at 5% gives present values of 4.7619 and 95.2381, respectively, which sum to 100.

The Macaulay duration is the present-value-weighted average time at which the cash flows are received: $\text{Macaulay duration} = [(1 \text{ d7 } 4.7619) + (2 \text{ d7 } 95.2381)] / 100 = 1.9524$ years. Modified duration adjusts Macaulay duration for the yield compounding interval. Because coupons and yield are annual, modified duration equals $1.9524 / (1 + 0.05) = 1.8594$ years. Rounded to two decimal places, this is 1.86.

Modified duration is the standard local sensitivity measure relating a small change in yield to an approximate percentage change in bond price. See the duration definitions and calculation framework in the [PRM certification information](#) and this [modified-duration reference](#).

QUESTION NO: 2

Let $N(\cdot)$ denote the cumulative distribution function and suppose that X and Y are standard normally distributed and uncorrelated. Using the fact that $N(1.96) = 0.975$, the probability that $X < 0$ and $Y < 1.96$ is approximately

- A. 0.25%
- B. 48.8%
- C. 0.495%

ANSWER: B

Explanation:

The correct probability is 0.488, or 48.8%. The intended inequalities are that X is less than 0 and Y is less than 1.96. Under the standard joint-normal interpretation used in this type of risk-measurement question, zero correlation implies independence. The joint probability can therefore be obtained by multiplying the two marginal probabilities. By symmetry of the standard normal distribution, $P(X < 0) = 0.5$. The supplied value $N(1.96) = P(Y < 1.96) = 0.975$. Thus, $P(X < 0, Y < 1.96) = 0.5 \times 0.975 = 0.4875$, which rounds to 0.488. Expressed as a percentage, this is 48.75%, or approximately 48.8%, rather than 0.488%. This calculation applies the standard-normal CDF and the product rule for independent events, both central tools in quantitative risk calculations. For background on the standard normal distribution and its CDF, see the [NIST standard normal distribution reference](#). PRMIA's professional-risk-management curriculum is available through [PRMIA](#).

QUESTION NO: 3

At what point x does the function $f(x) = x^3 - 4x^2 + 1$ have a local minimum?

- A. -0.666666667
- B. 0
- C. 2.66667
- D. 2

ANSWER: C

Explanation:

The correct local-minimum point is **2.66667**, which is the decimal approximation of **8/3**. To locate a local extremum, first differentiate the function: $f'(x) = 3x^2 - 8x$. Setting this first derivative equal to zero gives $x(3x - 8) = 0$, so the critical points occur at $x = 0$ and $x = 8/3$. The second-derivative test determines the nature of the critical point. Here, $f''(x) = 6x - 8$. At $x = 8/3$, the second derivative equals 8, which is positive. A positive second derivative means that the graph is concave upward at that critical point, so the function has a local minimum there. Therefore, the requested x-coordinate is 8/3, or approximately 2.66667. This use of stationary points together with curvature is the standard calculus procedure for identifying local minima in smooth functions. See [OpenStax, Maxima and Minima](#) and [Paul's Online Math Notes on critical points](#).

QUESTION NO: 4

Consider the linear regression model for the returns of stock A and the returns of stock B.

Stock A is 50% more volatile than stock B. Which of the following statements is TRUE?

- A. The stocks must be positively correlated ()
- B. Beta must be positive ()
- C. Beta must be greater in absolute value than the correlation of the stocks ()
- D. Alpha must be positive ()

ANSWER: C

Explanation:

"Beta must be greater in absolute value than the correlation of the stocks" is correct when stock A's return is the dependent variable and stock B's return is the regressor. The population regression beta is $\beta = \text{Cov}(R_A, R_B) / \text{Var}(R_B)$. Writing covariance in terms of correlation and standard deviations gives $\beta = \rho_{AB}(\sigma_A/\sigma_B)$. Since stock A is 50% more volatile, $\sigma_A = 1.5\sigma_B$, so $\beta = 1.5\rho_{AB}$. Consequently, the magnitude of beta is 1.5 times the magnitude of the correlation: $|\beta| = 1.5|\rho_{AB}|$. Thus beta has a greater absolute magnitude than the correlation for the nonzero correlation relationship ordinarily intended in this regression comparison. This result follows directly from the regression-slope identity and illustrates that beta combines both co-movement, measured by correlation, and relative volatility. See the NIST discussion of the relationship between the simple-regression slope, correlation, and standard deviations at [NIST/SEMATECH e-Handbook](#) and PRMIA's certification information at [PRMIA PRM Certification](#).

QUESTION NO: 5

Which of the following statements is true?

- A. Discrete and continuous compounding produce the same results if the discount rate is positive.
- B. Continuous compounding is the better method because it results in higher present values compared to discrete compounding.
- C. Continuous compounding can be thought as making the compounding period infinitesimally small.

D. The constant plays an important role in the mathematical description of continuous compounding.

ANSWER: C

Explanation:

“Continuous compounding can be thought as making the compounding period infinitesimally small.” is true because continuous compounding is defined as the limiting case of increasingly frequent discrete compounding. For an annual nominal rate r , an investment compounded m times per year grows over time t as $(1 + r/m)^{mt}$. As the number of compounding intervals m increases without bound, each interval becomes arbitrarily short and the expression converges to e^{rt} . This limit is the mathematical basis for the continuously compounded accumulation factor.

Accordingly, a continuously compounded zero-coupon investment with present value PV has future value $PV e^{rt}$, while a future cash flow FV is discounted using $PV = FV e^{-rt}$. The associated continuously compounded rate is therefore naturally expressed using logarithms: $r = \ln(FV/PV)/t$. This convention is particularly useful in risk measurement and derivative pricing because continuously compounded returns add over successive time intervals, which makes time aggregation and many stochastic models more tractable. The limiting relationship between periodic and continuous compounding is described in [OpenStax College Algebra](#) and the continuous-compounding formula is summarized by [Investopedia](#).

QUESTION NO: 6

A 95% confidence interval for a parameter estimate can be interpreted as follows:

- A. If this confidence-interval procedure were repeated on many independent samples, 95% of the resulting intervals would contain the true parameter value.
- B. The probability that the real value of the parameter is outside this interval is 95%.
- C. The probability that the estimated value of the parameter is within this interval is 95%.
- D. The probability that the estimated value of the parameter is outside this interval is 95%.

ANSWER: A

Explanation:

A 95% confidence interval is a frequentist statement about the long-run performance of the interval-construction procedure. If independent samples were repeatedly drawn under the same assumptions, and a confidence interval were calculated from each sample using the same method, approximately 95% of those intervals would contain the fixed, true parameter value. The parameter itself is not treated as random in this interpretation; it is the sample, and therefore the calculated interval endpoints, that vary from sample to sample. Consequently, once a particular interval has been calculated, it either contains the true value or it does not. The 95% figure describes the method's pre-sample coverage probability, not a probability assigned to the already-fixed parameter after observing the data. This interpretation is central to statistical estimation and risk measurement because it distinguishes uncertainty arising from sampling variation from a probability distribution placed on an unknown parameter. See the [NIST/SEMATECH e-Handbook discussion of confidence intervals](#) and the [Carnegie Mellon University notes on confidence intervals](#).

QUESTION NO: 7

Which of the following can induce a 'multicollinearity' problem in a regression model?

- A. A large negative correlation between the dependent variable and one of the explanatory variables
- B. A high positive correlation between the dependent variable and one of the explanatory variables
- C. A high positive correlation between two explanatory variables
- D. The omission of a relevant explanatory variable

ANSWER: C

Explanation:

A high positive correlation between two explanatory variables can induce multicollinearity because multicollinearity concerns linear dependence among the regression model's independent (explanatory) variables. When two predictors contain very similar information, the model has difficulty separating their individual contributions to the dependent variable. In an ordinary least-squares regression, this weakens the precision with which the associated coefficients are estimated: their standard errors can become large, coefficient estimates may be sensitive to small changes in the sample, and individual statistical inference becomes less reliable. The issue is therefore about the relationships within the predictor matrix, rather than the strength or direction of an explanatory variable's relationship with the outcome. A correlation need not be exactly one to create a practical problem; sufficiently strong correlation, especially in combination with relationships involving other predictors, can materially inflate coefficient-estimate variance. This is commonly assessed using diagnostics such as the variance inflation factor (VIF), which measures how much the variance of a coefficient estimate is increased by linear relationships with the other explanatory variables. See Penn State's [multicollinearity overview](#) and the NIST [discussion of multicollinearity](#).

QUESTION NO: 8

Which statement regarding the matrix below is true?

- A. It is not positive definite
- B. It is positive semi-definite
- C. It is positive definite
- D. It is negative definite
- E. Cannot be determined because the matrix is not provided.

ANSWER: E

Explanation:

Cannot be determined because the matrix itself is not provided in the question data. Definiteness is a property that must be evaluated from a matrix's entries, not inferred from the wording of the question. For a real symmetric matrix, the standard tests include checking whether all eigenvalues are strictly positive (positive definite), all are nonnegative (positive semidefinite), or using Sylvester's criterion, under which positive definiteness requires every leading principal minor to be positive. Without the displayed matrix, none of these calculations can be performed: its eigenvalues, determinant, principal minors, symmetry, and quadratic-form behavior are all unknown. The only technically valid conclusion from the supplied material is therefore that the matrix classification cannot be established. This distinction is important in risk measurement because covariance and correlation matrices are expected to be positive semidefinite; verification requires the actual numerical matrix. See the discussion of positive-definite matrices and Sylvester's criterion in [Definite matrix](#) and the numerical linear-algebra reference in [NumPy eigenvalue documentation](#).

QUESTION NO: 9

Which statement regarding the matrix below is true?

- A. It is not positive definite
- B. It is positive semi-definite
- C. It is positive definite
- D. It is negative definite
- E. Cannot be determined because the matrix is not provided.

ANSWER: E

Explanation:

Cannot be determined because the referenced matrix is not present in the question. Definiteness is a property that must be evaluated from the matrix entries, usually under the assumption that the matrix is real and symmetric. For a real symmetric matrix, positive definiteness means that the quadratic form $x^T A x$ is strictly positive for every nonzero vector x ; equivalently, all eigenvalues are positive. Positive semidefiniteness permits zero eigenvalues, while negative definiteness requires the quadratic form to be strictly negative for every nonzero vector. Without the actual matrix, there is no basis for calculating eigenvalues, checking leading principal minors via Sylvester's criterion, or evaluating the relevant quadratic form. Therefore, no definitive classification as positive definite, positive semidefinite, negative definite, or otherwise can validly be made from the supplied content. This is particularly important in risk measurement because covariance matrices must be positive semidefinite to represent valid variances and covariances. For background on positive-definite matrices and their characterizations, see [Definite matrix](#) and the [Wolfram MathWorld overview of positive-definite matrices](#).

QUESTION NO: 10

A typical leptokurtotic distribution can be described as a distribution that is relative to a normal distribution

- A. more peaked at the center and with heavy (fat) tails
- B. peaked and thin at the center and with thin tails
- C. flat and thick at the center and with heavy (fat) tails
- D. flat and thick at the center and with thin tails

ANSWER: A

Explanation:

A typical leptokurtotic distribution is more sharply concentrated around its center and has heavier, or "fatter," tails than a normal distribution with comparable location and scale. Thus, "more peaked at the center and with heavy (fat) tails" is correct. Kurtosis describes aspects of distributional shape associated particularly with tail weight and the propensity for observations far from the mean. A leptokurtotic distribution has positive excess kurtosis relative to the normal benchmark, whose excess kurtosis is zero. In risk measurement, this feature is important because heavy tails imply that large gains and losses occur more frequently than a normal-model assumption would indicate. Consequently, models that assume normality can materially understate the probability and severity of extreme market moves when returns are leptokurtic. The central peak and heavy tails should be understood as a relative shape comparison with the normal distribution, rather than as a complete specification of a particular probability distribution. For background on kurtosis and excess kurtosis, see the [NIST Engineering Statistics Handbook](#). The risk-management significance of non-normal return distributions is also consistent with the [Basel Committee's market-risk framework](#), which emphasizes measurement of tail risk.

QUESTION NO: 11

I have a portfolio of two stocks. The weights are 60% and 40% respectively, the volatilities are both 20%, while the correlation of returns is 100%. The volatility of my portfolio is

- A. 4%
- B. 14.4%
- C. 20%
- D. 24%

ANSWER: C

Explanation:

QUESTION NO: 12

ANSWER: C

QUESTION NO: 13

ANSWER: D

Explanation:

The false statement is “the definite integral of the PDF from minus infinity to plus infinity is zero.” For a continuous random variable with probability density function f , probabilities are obtained by integrating the density over sets of possible values. Since the random variable must take some value on the real line with total probability one, its density is normalized so that $\int_{-\infty}^{\infty} f(x) dx = 1$, not zero. More generally, the cumulative distribution function is defined by $F(x) = P(X \leq x)$, and, where a density exists, $F(x)$ is the accumulated integral of that density up to x . Thus, integrating the PDF over the entire support accounts for all possible outcomes and necessarily yields total probability one. This normalization condition is fundamental to using a PDF in risk measurement, including calculating probabilities of losses over intervals and deriving distributional quantities. The distinction matters because a non-normalized function cannot serve as a probability density without first being scaled appropriately. See the probability-density-function definition in [NIST/SEMATECH e-Handbook](#) and the cumulative-distribution-function definition in [NIST/SEMATECH e-Handbook](#).

QUESTION NO: 14

In a portfolio there are 7 bonds: 2 AAA Corporate bonds, 2 AAA Agency bonds, 1 AA Corporate and 2 AA Agency bonds. By an unexplained characteristic the probability of any specific AAA bond outperforming the others is twice the probability of any specific AA bond outperforming the others. What is the probability that an AA bond or a Corporate bond outperforms all of the others?

- A. 5/7
- B. 8/11
- C. 6/11
- D. None of these

ANSWER: D**Explanation:**

“None of these” is correct because the required probability is 7/11, which is not listed among the numerical choices. Let the probability that any particular AA bond is the best-performing bond be x . The stated relationship makes the probability for each particular AAA bond $2x$. There are four AAA bonds and three AA bonds, and exactly one bond must outperform all others. Therefore, total probability is $4(2x) + 3(x) = 11x = 1$, giving $x = 1/11$.

The event in question is that the winning bond is AA, Corporate, or both. It includes the two AAA Corporate bonds, the one AA Corporate bond, and the two AA Agency bonds. Equivalently, it excludes only the two AAA Agency bonds. Summing the included probabilities gives $2(2/11) + 1/11 + 2(1/11) = 7/11$. This calculation applies the mutually exclusive, exhaustive outcome framework used in basic probability: probabilities of distinct possible winning bonds sum to one, and probabilities of mutually exclusive outcomes are added. See [OpenStax: Combined Events and the Sum Rule](#) and PRMIA's [PRM certification information](#).

QUESTION NO: 15

Which of the following statements are true about Maximum Likelihood Estimation?

- (i) MLE can be applied even if the error terms are not i.i.d. normal.
- (ii) MLE involves integrating a likelihood function or a log-likelihood function.
- (iii) MLE yields parameter estimates that are consistent.

- A. (i) and (ii)
- B. (i) only
- C. (i) and (iii)
- D. (i), (ii), and (iii)

Explanation:

(i) and (iii) is correct. Maximum likelihood estimation is a general estimation principle: the analyst specifies a probability model for the observed data conditional on unknown parameters and selects parameter values that maximize the resulting likelihood, or equivalently its logarithm. It is therefore not inherently restricted to regression models with independent, identically distributed normal error terms. For example, likelihood-based models can be constructed for non-normal data, heteroskedastic observations, time-series dependence, and other explicitly modelled dependence structures. The key requirement is that an appropriate joint probability model, or a suitable likelihood formulation, is available for the data.

Under standard regularity and identification conditions, maximum-likelihood estimators are consistent. In practical terms, as the sample size grows, the estimator converges in probability to the true parameter value. This large-sample property is central to the use of MLE in risk modelling, where parameter estimates for distributions, volatility processes, and dependence models are commonly assessed through asymptotic theory. The likelihood is maximized rather than integrated; integration is associated with operations such as marginalization or Bayesian posterior calculations, not the defining MLE optimization step. See the [NIST overview of maximum likelihood estimation](#) and Penn State's discussion of [properties of maximum likelihood estimators](#).