

PRM Certification - Exam II: Mathematical Foundations of Risk Measurement

PRMIA 8002

Version Demo

Total Demo Questions: 14

Total Premium Questions: 147

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Foundations of Risk Measurement	35
Topic 2, Probability Theory	22
Topic 3, Mathematical Statistics	36
Topic 4, Numerical Methods	54
Total	147

QUESTION NO: 1

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. Assign each individual AA bond an outperformance weight of 1. Since each specific AAA bond is twice as likely to outperform as each specific AA bond, assign each AAA bond a weight of 2. There are four AAA bonds, contributing total weight 8, and three AA bonds, contributing total weight 3. Thus, the total probability weight is 11.

The event in question is that the winning bond is either AA-rated, Corporate, or both. All three AA bonds qualify, with combined weight 3. In addition, the two AAA Corporate bonds qualify because they are Corporate; their combined weight is 4. The AA Corporate bond is already included among the AA bonds, so it must not be counted a second time. Therefore, the favorable weight is 3 + 4 = 7, giving a probability of 7/11. This uses the union principle for overlapping events: the AA Corporate bond belongs to both classifications but represents only one possible winning bond. For background on probability rules and unions of events, see [NIST's probability reference](#) and [OpenStax's discussion of event probabilities](#).

QUESTION NO: 2

Simple linear regression involves one dependent variable, one independent variable and one error variable. In contrast, multiple linear regression uses...

- A. One dependent variable, many independent variables, one error variable
- B. Many dependent variables, one independent variable, one error variable
- C. One dependent variable, one independent variable, many error variables
- D. Many dependent variables, many independent variables, many error variables

ANSWER: A

Explanation:

Multiple linear regression models a single response (dependent) variable as a linear function of two or more explanatory (independent) variables, together with an error term. Its population form is commonly written as $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \varepsilon$, where Y is the one dependent variable, the X terms are the multiple independent variables, the β terms are regression coefficients, and ε represents the disturbance or unexplained portion of the outcome. This framework is particularly useful in risk measurement because a risk-relevant outcome can depend simultaneously on several drivers—for example, returns may be related to market, sector, interest-rate, and volatility factors. The model estimates each explanatory variable's association with the dependent variable while holding the other included variables constant. Although a data set contains a residual observation for each fitted record, the regression specification contains one error term representing the random component of the dependent-variable equation. Therefore, “One dependent variable, many independent variables, one error variable” accurately describes multiple linear regression. For a formal treatment, see the [NIST Engineering Statistics Handbook discussion of multiple linear regression](#) and the [Penn State STAT 501 introduction to multiple linear regression](#).

QUESTION NO: 3

Which of the following is consistent with the definition of a Type I error?

- A. The probability of a Type I error is 100% minus the significance level
- B. A Type I error would have occurred if the performance of a stock was positively correlated with the performance of a hedge fund, but in a linear regression, the hypothesis of positive correlation was rejected
- C. A Type I error would have occurred if the performance of a stock was positively correlated with the performance of a hedge fund, but in a linear regression, the hypothesis of no correlation was rejected
- D. A Type I occurs whenever data series are serially correlated

ANSWER: B

Explanation:

“A Type I error would have occurred if the performance of a stock was positively correlated with the performance of a hedge fund, but in a linear regression, the hypothesis of positive correlation was rejected” is consistent with a Type I error because it describes rejecting a hypothesis that is in fact true. In hypothesis-testing terminology, a Type I error is a false positive: the test rejects the null hypothesis even though the null hypothesis holds in the underlying population. Here, the stated true relationship is positive correlation, while the regression-based test rejects the hypothesis asserting that positive relationship. Thus, the statistical conclusion incorrectly rejects a true proposition.

The probability of making a Type I error is conventionally controlled by the chosen significance level, denoted by α . This is why careful specification of the null hypothesis, the alternative hypothesis, and the significance threshold is essential when interpreting regression results. The U.S. National Institute of Standards and Technology defines a Type I error as rejecting the null hypothesis when it is true, and describes α as the probability of this error under the null. See the [NIST/SEMATECH e-Handbook discussion of hypothesis testing](#) and [NIST's definitions of Type I and Type II errors](#).

QUESTION NO: 4

You invest \$2m in a bank savings account with a constant interest rate of 5% p.a. What is the value of the investment in 2 years time if interest is compounded quarterly?

- A. \$2,208,972
- B. \$2,210,342
- C. \$2,205,000
- D. None of them

ANSWER: A

Explanation:

The correct value is **\$2,208,972**. With quarterly compounding, the stated annual interest rate must first be converted into the rate earned in each quarter. A 5% annual nominal rate compounded four times per year gives a quarterly rate of $5\% \div 4 = 1.25\%$, or 0.0125. Over two years there are $2 \times 4 = 8$ quarterly compounding periods. The future-value calculation is therefore: $\$2,000,000 \times (1 + 0.05/4)^8$. This equals $\$2,000,000 \times 1.104486\dots$, giving approximately \$2,208,972 when rounded to the nearest dollar. The key feature of compound interest is that each quarter's interest is added to the balance, so subsequent quarterly interest is earned on both the original \$2 million and interest previously credited. This is the standard future-value formula, $FV = PV(1 + r/m)^{mt}$, where m is the number of compounding periods per year. See the discussion of compound-interest calculations in [OpenStax College Algebra](#) and the future-value formula in [Investor.gov's compound interest calculator](#).

QUESTION NO: 5

You are given the following values of a quadratic function $f(x)$: $f(0)=0$, $f(1)=-2$, $f(2)=-5$. On the basis of these data, the derivative $f'(0)$ is

- A. in the interval $] -2.5, -2[$
- B. equal to -2
- C. in the interval $] -2, +\infty[$
- D. in the interval $] -\infty, -2.5]$

ANSWER: C

Explanation:

The correct answer is *in the interval $] -2, +\infty[$* . Write the quadratic as $f(x) = ax^2 + bx + c$. Since $f(0) = 0$, $c = 0$. The values at one and two then give $a + b = -2$ and $4a + 2b = -5$. Subtracting twice the first equation from the second yields $2a = -1$, so $a = -0.5$; consequently, $b = -1.5$. The derivative of a quadratic is $f'(x) = 2ax + b$, and therefore $f'(0) = b = -1.5$. Because -1.5 is strictly greater than -2 , it belongs to the stated open interval extending from -2 to positive infinity. This result follows from uniquely fitting a degree-two polynomial through the three supplied function values and then differentiating that fitted polynomial. The standard power rule used here states that the derivative of ax^n is nax^{n-1} ; see [OpenStax, Differentiation Rules](#). For the polynomial form and its derivative, see [Quadratic function](#).

QUESTION NO: 6

In a 2-step binomial tree, at each step the underlying price can move up by a factor of $u = 1.1$ or down by a factor of $d = 1/u$. The continuously compounded risk free interest rate over each time step is 1% and there are no dividends paid on the underlying. Use the Cox, Ross, Rubinstein parameterization to find the risk neutral probability and hence find the value of a European put option with strike 102, given that the underlying price is currently 100.

- A. 5.19
- B. 5.66
- C. 6.31
- D. 4.18

ANSWER: A

Explanation:

The correct value is **5.19**. In the Cox–Ross–Rubinstein binomial model, the one-step risk-neutral up probability is determined by matching the expected stock return under the risk-neutral measure to the continuously compounded risk-free growth factor: $p = (e^{0.01} - d)/(u - d)$. With $u = 1.1$ and $d = 1/1.1 = 0.90909$, this gives $p \approx 0.52883$ and a down probability of approximately 0.47117.

After two steps, the possible underlying prices are 121, 100, and approximately 82.6446. A put with strike 102 therefore has terminal payoffs of 0, 2, and approximately 19.3554, respectively. The associated risk-neutral probabilities are p^2 , $2p(1 - p)$, and $(1 - p)^2$. Thus, the expected terminal put payoff is approximately 5.2936. Discounting this amount for two periods at the continuously compounded risk-free rate gives $e^{-0.02} \times 5.2936 \approx 5.1888$, which rounds to 5.19. This is the standard no-arbitrage valuation approach in a recombining CRR tree; see the overview of the [Cox–Ross–Rubinstein formulas](#) and the [binomial option-pricing model](#).

QUESTION NO: 7

What is the total derivative of the function $f(x,y) = \ln(x+y)$, where $\ln()$ denotes the natural logarithmic function?

- A. $1 / (x+y)$

- B. $(\Delta x + \Delta y) / (x+y)$
- C. $-\Delta x/(x+y) - \Delta y/(x+y)$
- D. $\ln(x+y) \Delta x + \ln(x+y) \Delta y$

ANSWER: B

Explanation:

The total derivative (more precisely, the total differential) is $df = f_x dx + f_y dy$. For $f(x,y) = \ln(x + y)$, apply the chain rule to the outer logarithm and the inner expression $x + y$. The derivative of $\ln(u)$ with respect to u is $1/u$, while the differential of the inner expression is $dx + dy$. Therefore, $df = (1/(x + y))(dx + dy) = (dx + dy)/(x + y)$. Written with Δx and Δy , this gives “ $(\Delta x + \Delta y) / (x+y)$,” which represents the first-order total change in the function for small changes in both variables. This result is valid on the domain $x + y > 0$, required for the natural logarithm. See the discussion of multivariable differentials in [Paul's Online Math Notes](#) and the derivative rule for logarithms in [OpenStax Calculus](#).

QUESTION NO: 8

Newton-Raphson iteration is used to find a solution of $x^5 - x^3 + x = 1$. If $x_n = 2$, what is x_{n+1} ?

- A. 2.362
- B. 1.623
- C. 1.638
- D. 0.377

ANSWER: C

Explanation:

The correct answer is **1.638**. Write the equation in root-finding form as $f(x) = x^5 - x^3 + x - 1$, so that a solution to the original equation is a value for which $f(x) = 0$. Newton-Raphson updates an approximation using $x_{n+1} = x_n - f(x_n)/f'(x_n)$. The derivative is $f'(x) = 5x^4 - 3x^2 + 1$. At the supplied iterate, $x_n = 2$, the function value is $f(2) = 2^5 - 2^3 + 2 - 1 = 25$, while the derivative is $f'(2) = 5(2^4) - 3(2^2) + 1 = 80 - 12 + 1 = 69$. Therefore, the next Newton-Raphson iterate is $x_{n+1} = 2 - 25/69 = 1.637681\dots$. Rounding this value to three decimal places gives 1.638. This calculation applies the standard Newton update, which uses the tangent-line approximation to move from the current estimate toward a zero of the function. Further background on the Newton-Raphson method and its update formula is available from [NIST's Engineering Statistics Handbook](#) and [Newton's method reference](#).

QUESTION NO: 9

Assume that 40% of all financial organizations investigated by authorities turn out to be fraudulent.

What is the probability of randomly investigating 2 different organizations and finding that neither is fraudulent; and what is the probability of finding exactly one being fraudulent?

- A. 2/5 and 1/2
- B. 2/5 and 3/5
- C. 1/3 and 8/17
- D. 9/25 and 12/25

ANSWER: D

Explanation:

The correct result is **9/25 and 12/25**. Let the probability that an investigated organization is fraudulent be 0.40, so the probability that it is not fraudulent is 0.60. Treating the outcomes for the two randomly selected organizations as independent draws from this population, the probability that neither organization is fraudulent is $0.60 \times 0.60 = 0.36$, which equals 9/25.

Exactly one organization is fraudulent when either the first is fraudulent and the second is not, or the first is not fraudulent and the second is fraudulent. Each ordering has probability $0.40 \times 0.60 = 0.24$. Since these two mutually exclusive orderings are added, the required probability is $0.24 + 0.24 = 0.48$, or 12/25. This follows the standard binomial model for two trials, where the probability of exactly one occurrence is given by the binomial coefficient 2 multiplied by 0.40 multiplied by 0.60. See the [NIST Engineering Statistics Handbook on the binomial distribution](#) and [OpenStax's binomial-distribution overview](#).

QUESTION NO: 10

A linear regression gives the following output:

Figures in square brackets are estimated standard errors of the coefficient estimates.

Which of the following is an approximate 95% confidence interval for the true value of the coefficient of ?

- A. [0, 1.5]
- B. [1, 2]
- C. [0, 3]
- D. None of the above

ANSWER: C**Explanation:**

[0, 3] is the approximate 95% confidence interval obtained by applying the usual large-sample regression confidence-interval rule: estimated coefficient plus or minus approximately two estimated standard errors. In other words, for a coefficient estimate $\hat{\beta}$ with reported standard error $SE(\hat{\beta})$, the approximate interval is $\hat{\beta} \pm 1.96 SE(\hat{\beta})$, commonly rounded to $\hat{\beta} \pm 2SE(\hat{\beta})$. Using the estimate and bracketed standard error from the intended regression output gives lower and upper endpoints of approximately zero and three, respectively. This interval quantifies sampling uncertainty around the unknown population regression coefficient under the standard assumptions supporting the estimator and its standard error. It is a confidence procedure: across repeated samples generated under those assumptions, approximately 95% of intervals constructed in this way would contain the true coefficient. For further background on regression confidence intervals and use of estimated standard errors, see the [NIST/SEMATECH e-Handbook discussion of confidence intervals](#) and the [Penn State regression inference notes](#).

QUESTION NO: 11

For the function $f(x) = 3x - x^3$ which of the following is true?

- A. $x = 0$ is a minimum
- B. $x = -3$ is a maximum
- C. $x = 2$ is a maximum
- D. None of these

ANSWER: D**Explanation:**

"None of these

" is correct, interpreting the expression as the standard polynomial $f(x) = 3x - x^3$. To identify local extrema, differentiate the function: $f'(x) = 3 - 3x^2$. Setting the first derivative equal to zero gives $3 - 3x^2 = 0$, so $x^2 = 1$ and the stationary points are $x = -1$ and $x = 1$. The second derivative is $f''(x) = -6x$. At $x = -1$, the second derivative is positive, establishing a local minimum; at $x = 1$, it is negative, establishing a local maximum. Therefore, zero is neither a local minimum nor maximum, and neither negative three nor two is a stationary point or extremum for this function. Since the actual extrema occur at values not stated in the listed claims, the statement that none of the listed assertions applies is the correct single choice. This uses the first- and second-derivative tests for local extrema, as described in [OpenStax Calculus: Maxima and Minima](#) and [OpenStax Calculus: Derivatives and Graph Shape](#).

QUESTION NO: 12

Variance reduction is:

- A. A technique that is applied in regression models to improve the accuracy of the coefficient estimates
- B. A numerical method for finding portfolio weights to minimize the variance of a portfolio that has a given expected return
- C. A numerical method for finding the variance of the underlying that is implicit in a market price of an option
- D. A method for reducing the number of simulations required in a Monte Carlo simulation

ANSWER: D

Explanation:

Variance reduction is a collection of Monte Carlo techniques designed to obtain a more precise estimate for a given computational budget, or equivalently to achieve a target precision with fewer simulation runs. Standard Monte Carlo estimates converge relatively slowly because sampling error declines at a rate proportional to the inverse square root of the number of simulated paths. In risk measurement and option valuation, where portfolios may be complex and simulations computationally expensive, reducing the estimator's variance is therefore highly valuable.

Common variance-reduction methods include antithetic variates, control variates, importance sampling, stratified sampling, and conditional Monte Carlo. These methods use known relationships, alternative sampling schemes, or deliberately selected scenarios to reduce dispersion in the estimated result while preserving the quantity being estimated. The benefit is not that the true variance of the underlying asset is changed; rather, the uncertainty of the simulation-based estimator is reduced. Thus, *A method for reducing the number of simulations required in a Monte Carlo simulation* accurately captures the purpose of variance reduction. See the [NIST Engineering Statistics Handbook discussion of Monte Carlo methods](#) and the [Springer reference on Monte Carlo methods in financial engineering](#).

QUESTION NO: 13

A 2-step binomial tree is used to value an American put option with strike 105, given that the underlying price is currently 100. At each step the underlying price can move up by 10 or down by 10 and the risk-neutral probability of an up move is 0.6. There are no dividends paid on the underlying and the continuously compounded risk free interest rate over each time step is 1%. What is the value of the option in this model?

- A. 7.12
- B. 6.59
- C. 7.44
- D. 7.29

ANSWER: A

Explanation:

The value is **7.12**. At maturity, the additive two-step stock-price tree has terminal stock prices of 120, 100, and 80. The corresponding put payoffs, calculated as $\max(105 - S, 0)$, are 0, 5, and 25. Because the put is American, valuation proceeds

by backward induction and at every nonterminal node compares immediate exercise with the discounted risk-neutral continuation value.

At the first-step up node, where the stock price is 110, the continuation value is $e^{-0.01} \times [0.6 \times 0 + 0.4 \times 5] = 1.98$, which exceeds the zero exercise value. At the first-step down node, where the stock price is 90, immediate exercise produces 15. The continuation value is $e^{-0.01} \times [0.6 \times 5 + 0.4 \times 25] = 12.87$, so optimal exercise occurs at that node and its value is 15. Finally, today's value is $e^{-0.01} \times [0.6 \times 1.98 + 0.4 \times 15] = 7.12$, after rounding. This use of risk-neutral discounting and node-by-node early-exercise comparison is the standard binomial approach for American options; see [PRMIA's PRM Certification overview](#) and [NYU Stern's derivatives resources](#).

QUESTION NO: 14

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

- A. The probability that the real value of the parameter is within this interval is 95%.
- 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%.
- E. If the sampling procedure were repeated many times, 95% of the constructed intervals would contain the true parameter.

ANSWER: E

Explanation:

In the standard frequentist interpretation, a 95% confidence interval is the result of a procedure with 95% coverage. If the same population were sampled repeatedly and an interval were constructed from each sample using the same method, approximately 95% of those intervals would contain the fixed, true parameter. The parameter itself is not treated as random in this framework; the data and, consequently, the interval endpoints vary from sample to sample. Once a particular interval has been calculated, it either contains the fixed parameter or it does not, so assigning a 95% probability to the parameter being in that already-observed interval is not the technically correct frequentist statement. This distinction matters in risk measurement because confidence intervals quantify uncertainty arising from sampling and estimation procedures, rather than a probability distribution directly assigned to an unknown fixed model parameter. A probability statement about the parameter conditional on observed data requires a Bayesian credible interval and an explicitly specified prior distribution. See the [NIST/SEMATECH e-Handbook discussion of confidence intervals](#) and the [Penn State explanation of confidence-interval interpretation](#).