

DUMPS ARENA

Applied Probability and Statistics (FZO1 C955)

WGU Applied-Probability-and-Statistics

Version Demo

Total Demo Questions: 17

Total Premium Questions: 176

Buy Premium PDF

<https://dumpsarena.co>

sales@dumpsarena.co

sales@dumpsarena.co
dumpsarena.co

QUESTION NO: 1

Poisson mean = 4. Probability of exactly 2 events?

- A. $(e^{-4} \times 4^2)/2! = 0.1465$
- B. 0.2
- C. 0.5
- D. 0.25

ANSWER: A

Explanation:

$(e^{-4} \times 4^2)/2! = 0.1465$ is correct because a Poisson random variable models the number of events occurring in a fixed interval when events occur independently at a constant average rate. Its probability mass function is $P(X = k) = e^{-\lambda} \lambda^k / k!$, where λ is the expected number of events and k is the exact count of interest. In this case, the mean is $\lambda = 4$ and the requested count is $k = 2$. Substituting these values gives $P(X = 2) = e^{-4} \times 4^2 / 2!$. Since 4^2 is 16 and $2!$ is 2, the expression simplifies to $8e^{-4}$. Using $e^{-4} \approx 0.0183156$ gives $8 \times 0.0183156 \approx 0.146525$, which rounds to 0.1465. The result represents approximately a 14.65% chance of observing exactly two events in an interval whose long-run mean count is four. This direct substitution into the Poisson probability formula is the standard method for finding an exact count probability. See the [NIST Poisson distribution reference](#) and the [OpenStax Poisson distribution overview](#).

QUESTION NO: 2

Correlation $r = 0.8$ â†’ strong positive relationship â†’

- A. True
- B. False
- C. True only for categorical variables
- D. True only when r is negative

ANSWER: A

Explanation:

True is correct because a correlation coefficient of $r = 0.8$ describes a strong positive linear association between two quantitative variables. Correlation coefficients range from -1 to 1 . The positive sign indicates that the variables tend to move in the same direction: observations with larger values of one variable generally occur with larger values of the other variable. The magnitude, 0.8 , is relatively close to 1 , indicating that the points in a scatterplot would tend to follow a fairly strong upward linear pattern. Although classifications such as “strong” are context dependent, 0.8 is commonly interpreted as a strong association in introductory statistics. This result does not mean every pair of observations follows the pattern exactly; a perfect positive linear relationship would have $r = 1$. It also describes association rather than establishing that one variable causes changes in the other. For further discussion of interpreting the direction and strength of correlation, see [OpenStax Introductory Statistics](#) and [NIST's correlation coefficient reference](#).

QUESTION NO: 3

Regression $R^2 = 0.64$. Interpretation?

- A. 64% of variation in Y explained by X
- B. Correlation = 0.64
- C. Slope = 0.64

D. 64% chance of prediction

ANSWER: A

Explanation:

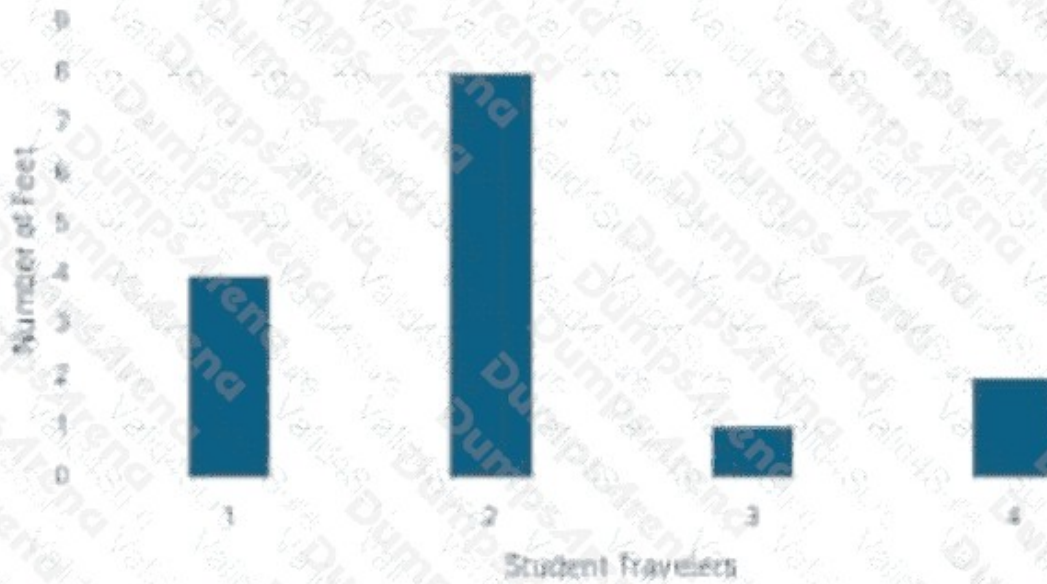
64% of variation in Y explained by X is correct. The coefficient of determination, R^2 , describes the proportion of the observed variability in the response variable that is accounted for by the fitted regression model. Thus, an R^2 value of 0.64 means that the model incorporating X explains 64% of the variation in Y around its mean. This is commonly called explained variation. The remaining 36% represents residual, or unexplained, variation: differences between the observed Y values and the values predicted by the regression line. That residual variation can arise from relevant factors not included in the model, random variability, measurement imprecision, or a relationship that is not fully represented by the selected linear model. R^2 is a descriptive measure of model fit and should be interpreted in the context of the data and subject area; it does not by itself establish causation or guarantee accurate predictions for every individual observation. In a simple linear regression, R^2 equals the square of the correlation coefficient, so the correlation's magnitude would be 0.80, with its sign determined by the direction of the slope. See the [NIST Engineering Statistics Handbook discussion of R-squared](#) and [Penn State STAT 200 regression interpretation guidance](#).

QUESTION NO: 4

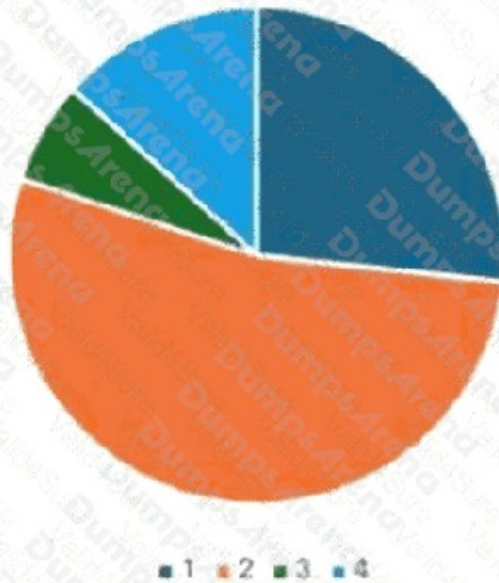
Which graph displays data in a manner that is accurate?

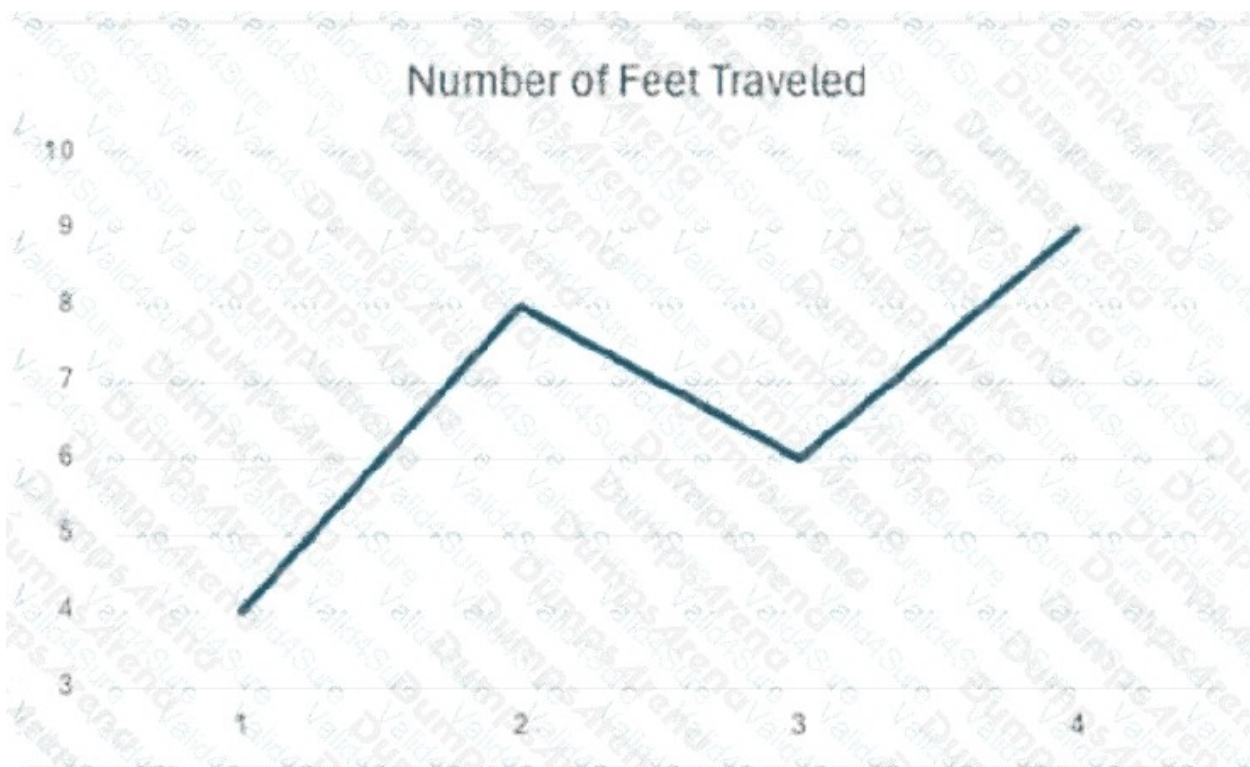


Number of Feet Traveled



Distance Traveled





- A. Bar graph titled “Number of Miles Traveled,” with the vertical axis labeled in feet.
- B. Bar graph titled “Number of Feet Traveled,” with the vertical axis labeled “Number of Feet” and the horizontal axis labeled “Student Travelers.”
- C. Pie chart titled “Distance Traveled.”
- D. Line graph titled “Number of Feet Traveled.”

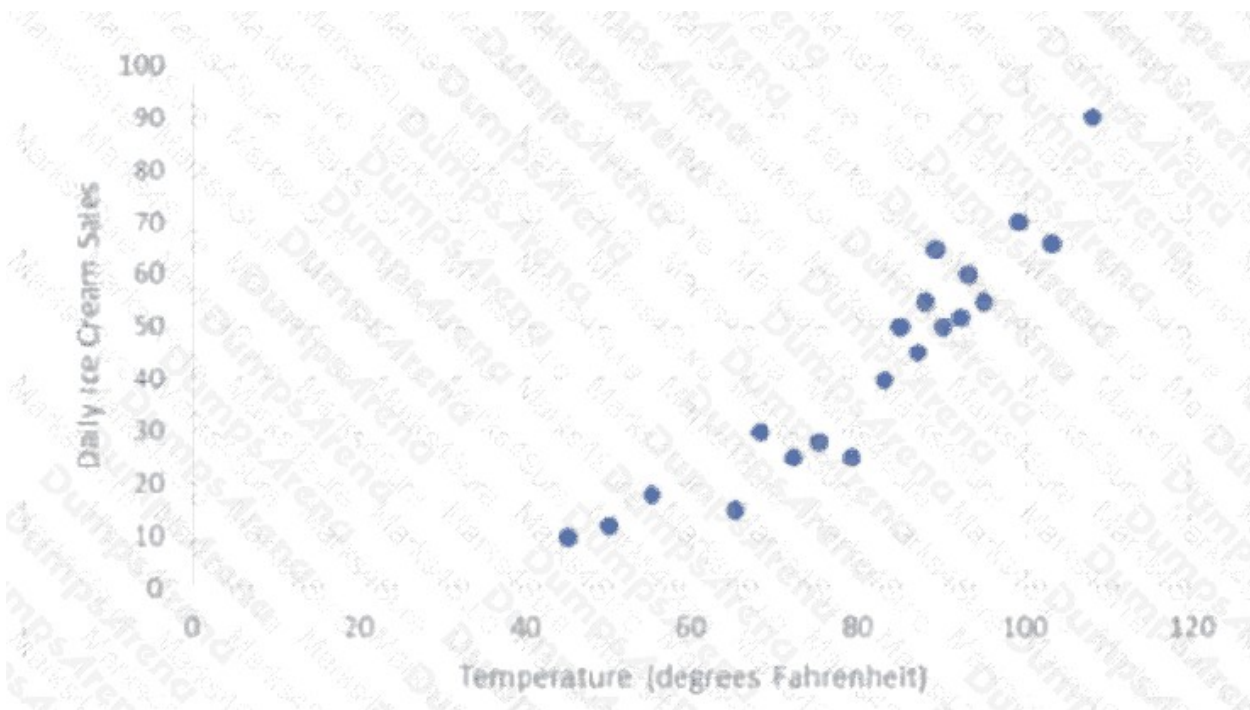
ANSWER: B

Explanation:

Bar graph titled “Number of Feet Traveled,” with the vertical axis labeled “Number of Feet” and the horizontal axis labeled “Student Travelers.” is accurate because it matches the graph type, labels, and units to the data being displayed. Each student traveler is a distinct category, while the corresponding distance is a quantitative measurement. A bar graph clearly represents this relationship: each separate bar corresponds to one student, and the height of that bar communicates the number of feet traveled. The vertical-axis label explicitly identifies the measured quantity and its unit, while the horizontal-axis label identifies the categories being compared. The title also uses the same unit as the measurement axis, so a reader can interpret the display without ambiguity or conversion. Accurate statistical graphics should identify variables clearly, use consistent units, and select a display suitable for the structure of the data. Bar charts are commonly used to compare values across categories, including named individuals or groups. See the OpenStax discussion of graphical displays at [OpenStax: Bar Graphs](#) and the NIST guidance on presenting data at [NIST: Bar Charts](#).

QUESTION NO: 5

A child’s parents want to analyze the relationship between the daily high temperature and ice cream sales at their child’s ice cream stand. They create the following scatterplot.



What is the estimated value of r , the correlation coefficient, between these variables?

- A. -0.93
- B. -0.39
- C. 0.39
- D. 0.93

ANSWER: D

Explanation:

0.93 is the appropriate estimate because the scatterplot displays a strong positive linear association between daily high temperature and ice cream sales. The plotted observations generally move upward from left to right: higher temperatures are associated with higher sales. In addition to being positive, the relationship appears strong because the points cluster relatively closely around an imagined upward-sloping straight line rather than being widely dispersed. The correlation coefficient r summarizes both direction and strength for a linear relationship between two quantitative variables. Its value must fall from -1 to 1 ; values close to 1 represent a strong positive linear pattern. Therefore, an estimate near positive 1 best matches this scatterplot, and 0.93 captures the evident tight upward trend. Correlation describes the strength and direction of the observed linear association; it does not, by itself, establish that temperature causes the changes in sales. For further guidance on interpreting scatterplots and correlation, see the [NIST Engineering Statistics Handbook discussion of correlation](#) and [OpenStax material on linear relationships and correlation](#).

QUESTION NO: 6

Probability of drawing red or black card = ?

- A. 1
- B. $1/2$
- C. 0
- D. 0.25

ANSWER: A

Explanation:

The correct answer is **1**, assuming a single card is drawn from a standard 52-card deck. Every card in a standard deck has one of two colors: red or black. The red cards are the hearts and diamonds, totaling 26 cards, and the black cards are the clubs and spades, also totaling 26 cards. Because these color categories include every possible card in the deck, the event “drawing a red or black card” is certain. Its probability is therefore the number of favorable outcomes divided by the total number of outcomes: $(26 + 26) / 52 = 52 / 52 = 1$. In probability, a value of 1 represents a certain event: the event will occur whenever the stated conditions hold. The addition is valid because a card cannot be both red and black, so the two events are mutually exclusive. This also illustrates the complement rule: there are no cards outside the red-or-black categories, making the probability of the complement zero and the probability of the stated event one. See the [OpenStax discussion of combined events](#) and [Encyclopaedia Britannica's probability overview](#).

QUESTION NO: 7

Probability of rolling an even number on a six-sided die?

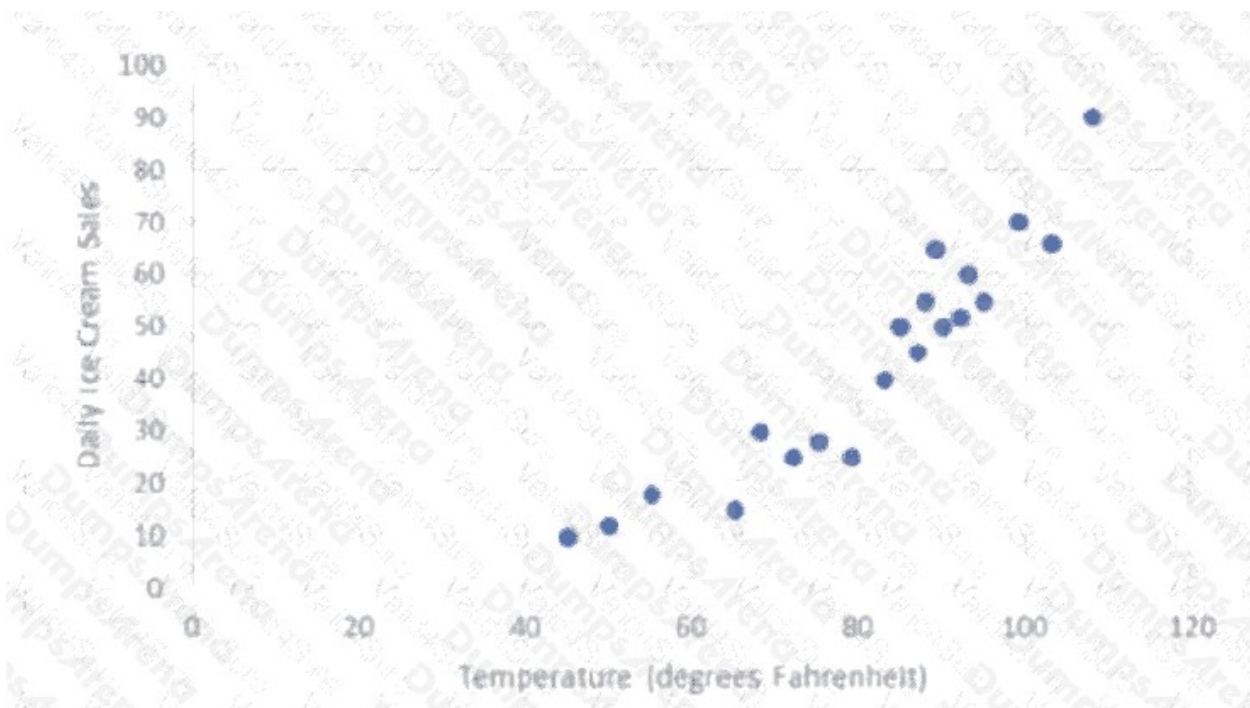
- A. 1/3
- B. 1/2
- C. 1/6
- D. 2/3

ANSWER: B**Explanation:**

The correct answer is **1/2**. For a standard fair six-sided die, the sample space consists of six equally likely results: 1, 2, 3, 4, 5, and 6. The outcomes that satisfy the event “rolling an even number” are 2, 4, and 6, so there are three favorable outcomes. The theoretical probability is calculated as the number of favorable outcomes divided by the total number of equally likely outcomes: $3/6$. Reducing this fraction by dividing both numerator and denominator by 3 gives $1/2$. Thus, an even result occurs on half of the possible die faces. This uses the basic theoretical-probability rule $P(\text{event}) = \text{favorable outcomes} / \text{total outcomes}$, applicable when each outcome in the sample space has the same chance of occurring. For further review of probability as a ratio of favorable equally likely outcomes to all possible outcomes, see [OpenStax Introductory Statistics: Probability Terminology](#) and [Encyclopaedia Britannica: Probability Theory](#).

QUESTION NO: 8

A childâ€™s parents want to analyze the relationship between the daily high temperature and ice cream sales at their childâ€™s ice cream stand. They create the following scatterplot.



What is the estimated value of r , the correlation coefficient, between these variables?

- A. $\hat{r} \sim 0.93$
- B. $\hat{r} \sim 0.39$
- C. 0.39
- D. 0.93

ANSWER: D

Explanation:

0.93 is the appropriate estimate because the scatterplot shows a strong positive linear association between daily high temperature and ice cream sales. The plotted values generally move upward from left to right: higher temperatures are associated with greater sales. This establishes that the correlation coefficient should be positive. In addition, the points appear to cluster relatively closely around an upward-sloping straight-line pattern rather than being widely dispersed. That degree of linear consistency indicates a correlation close to positive one.

The correlation coefficient, r , measures the direction and strength of a linear relationship and ranges from -1 to 1 . Values near 1 represent a strong positive linear relationship, meaning that observations tend to increase together in a pattern that is well approximated by a line. The visual strength of this upward pattern supports an estimate of 0.93 . This coefficient describes association, not proof that temperature alone causes the change in sales; nevertheless, it accurately summarizes the strong positive relationship displayed in the graph. See OpenStax's discussion of correlation at [OpenStax Introductory Statistics](#) and the NIST overview of correlation at [NIST/SEMATECH e-Handbook](#).

QUESTION NO: 9

Rare event count per interval modeled by:

- A. Poisson
- B. Binomial
- C. Normal
- D. Uniform

ANSWER: A

Explanation:

Poisson is the appropriate distribution for a count of events occurring within a specified interval, such as a time period, length of roadway, geographic area, or volume. It applies when events occur independently, the average occurrence rate is stable over comparable intervals, and the outcome is a nonnegative whole-number count. The distribution is defined by its rate parameter, λ (lambda), which represents the expected number of events in one interval. For example, if a help desk receives an average of three calls per minute, a Poisson model can be used to find the probability of receiving a particular number of calls during the next minute. “Rare event count per interval” is a standard description of this use case because the model is especially useful for infrequent occurrences measured over opportunities for occurrence rather than across a fixed number of trials. The probability of observing a count is determined from λ using the Poisson probability mass function. The [NIST/SEMATECH e-Handbook description of the Poisson distribution](#) identifies it as a discrete distribution for event counts in a fixed interval. See also the [OpenStax discussion of discrete probability distributions](#) for foundational context on count-based models.

QUESTION NO: 10

Probability of rolling sum = 12 = ?

- A. 1/36
- B. 1/6
- C. 2/36
- D. 1/12

ANSWER: A

Explanation:

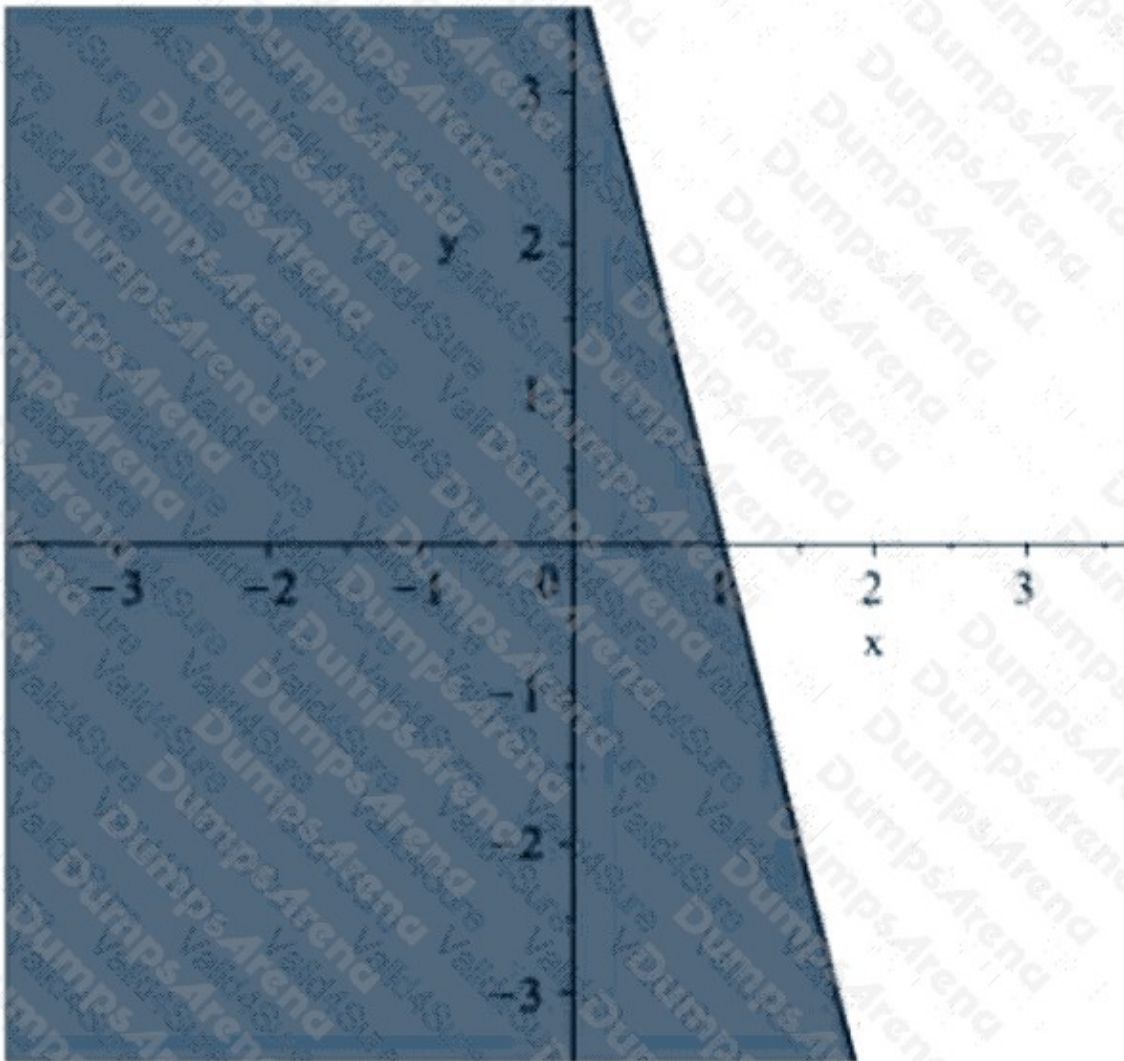
The correct answer is **1/36**. When rolling two fair six-sided dice, each die has six possible results, so there are $6 \times 6 = 36$ equally likely ordered outcomes. “Ordered” means that a first-die result of 4 and a second-die result of 6 is distinct from a first-die result of 6 and a second-die result of 4, even though both produce the same total. To obtain a sum of 12, however, there is only one possible outcome: the first die shows 6 and the second die shows 6, written as (6, 6). The theoretical probability is therefore the number of favorable outcomes divided by the total number of equally likely outcomes: 1/36. This approach follows the fundamental probability rule $P(\text{event}) = \text{favorable outcomes} / \text{total outcomes}$ when each outcome in the sample space is equally likely. A two-die outcome table is also a useful visual tool: only the cell where both dice are 6 has a total of 12. See the [OpenStax discussion of sample spaces and equally likely outcomes](#) and [NIST’s overview of discrete probability distributions](#).

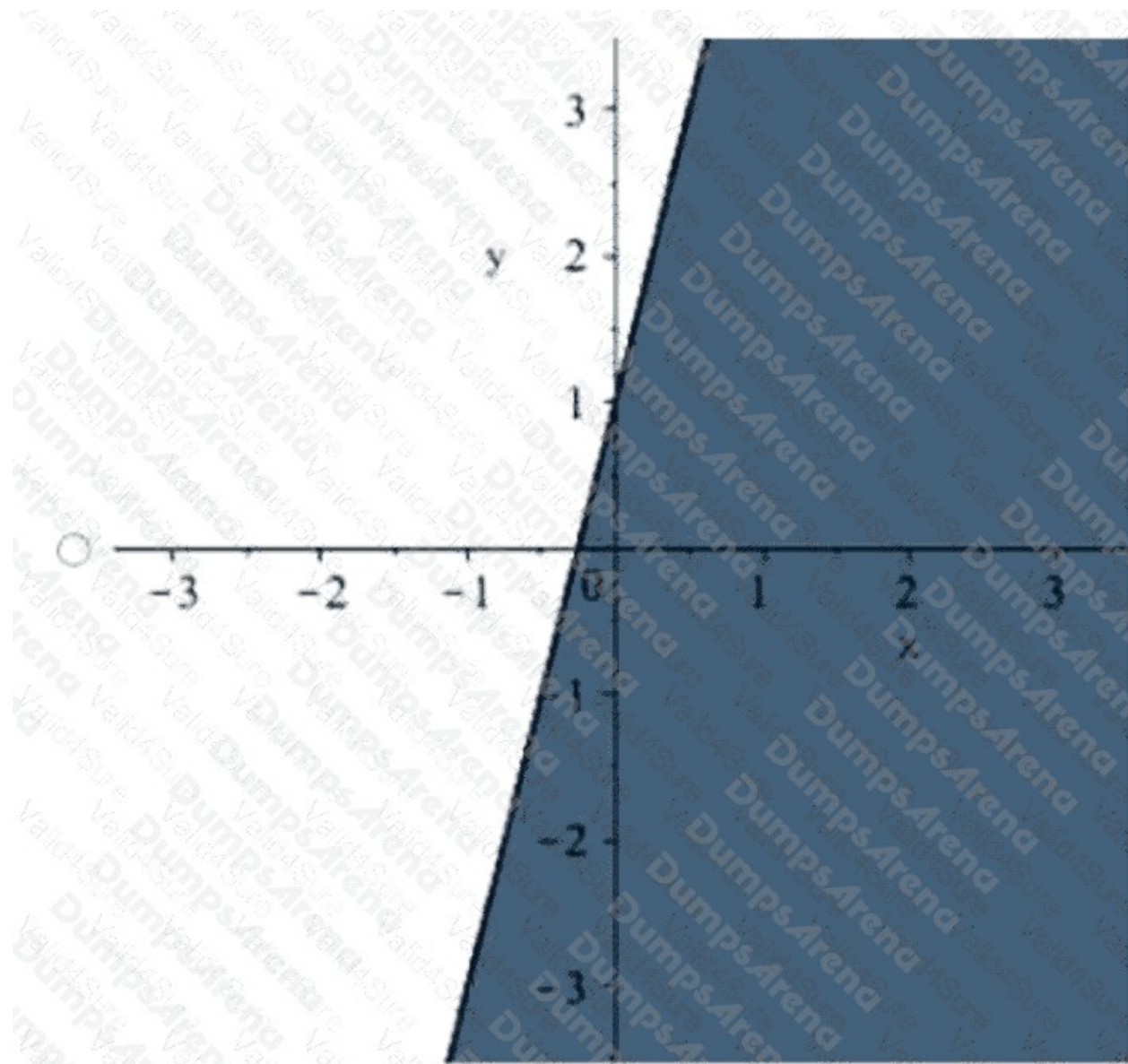
QUESTION NO: 11

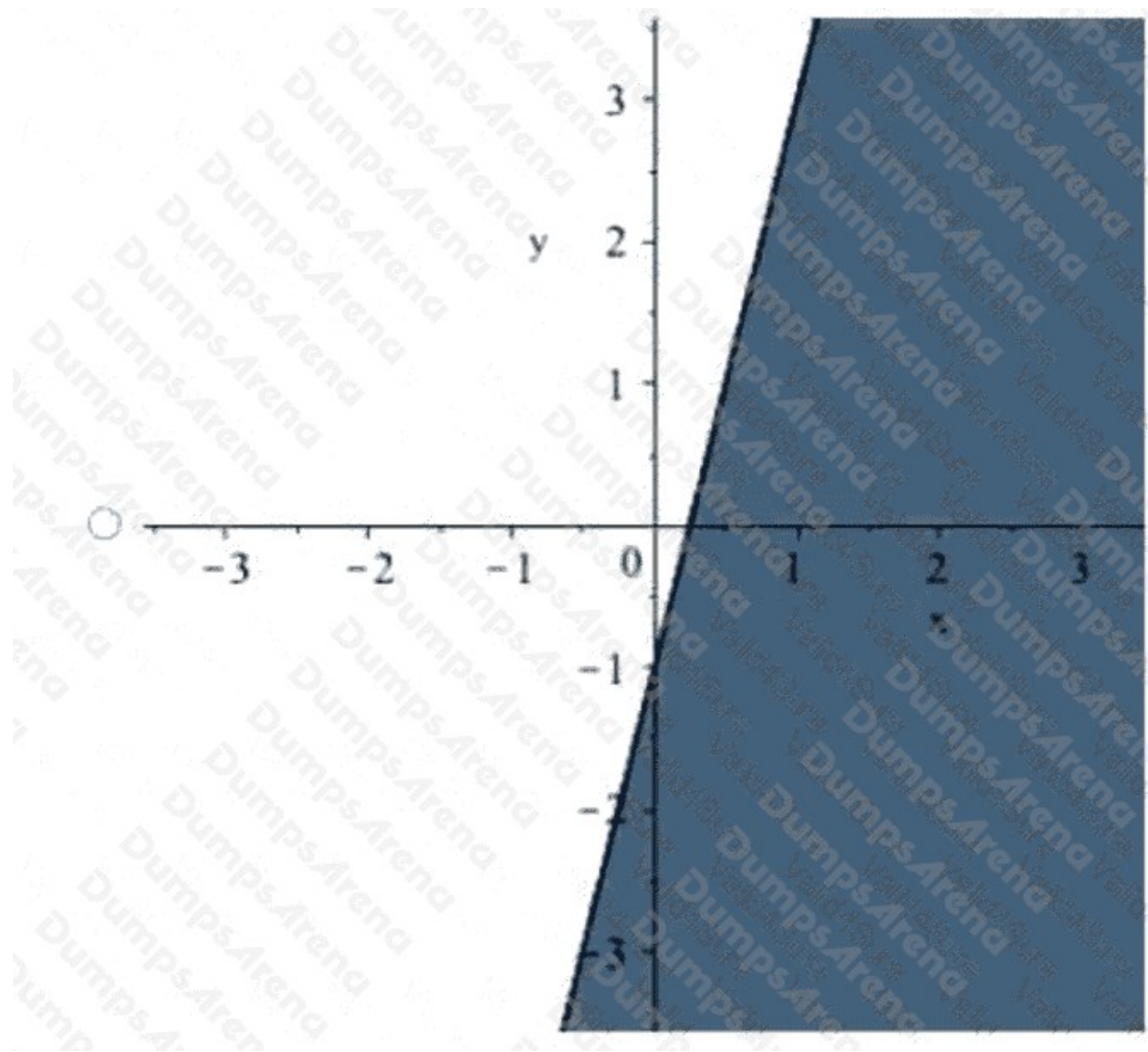
Review the following inequality:

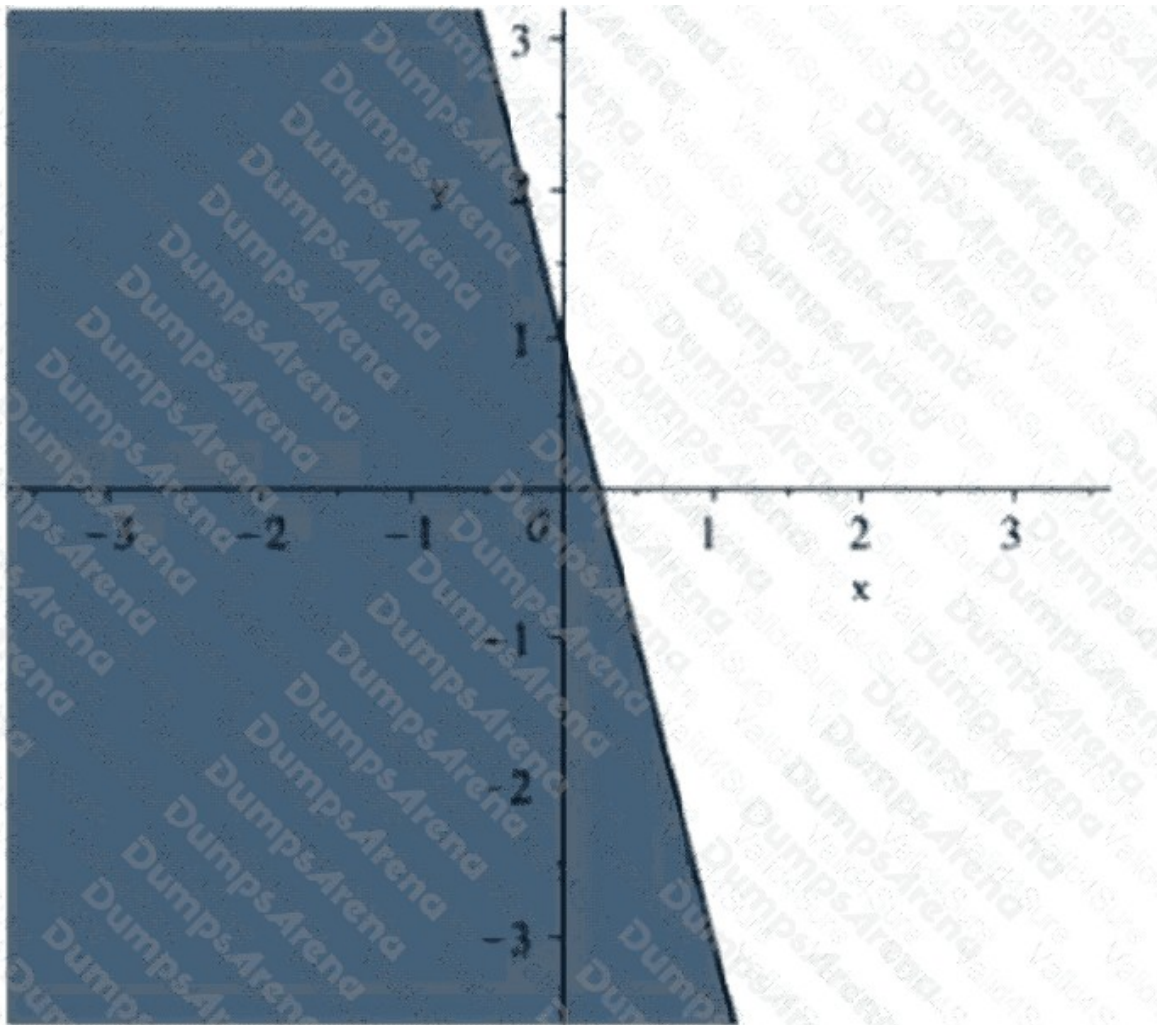
$$y \geq 4x + 1$$

Which shaded region is described by this inequality?









- A. Graph A
- B. Graph B
- C. Graph C
- D. Graph D

ANSWER: D

Explanation:

Graph D correctly represents the inequality $y \leq -4x + 1$. Its boundary is found by replacing the inequality symbol with an equals sign, producing $y = -4x + 1$. In slope-intercept form, the slope is -4 , so the line falls steeply as x increases, and the y -intercept is 1 , meaning it crosses the y -axis at $(0, 1)$. Because the relation includes “equal to,” points on the boundary satisfy the inequality; therefore, the graph must use a solid boundary line. The inequality describes all points whose y -values are less than or equal to the y -value on that line, so the solution region is the side below the line. Testing the origin confirms the appropriate side: substituting $(0, 0)$ gives $0 \leq -4(0) + 1$, or $0 \leq 1$, which is true. Thus, the shaded region must include the origin while remaining on or below the solid, downward-sloping boundary. This follows the standard procedure for graphing linear inequalities: graph the boundary, determine whether it is included, and use a test point to identify the solution side. See [OpenStax: Systems of Linear Inequalities](#) and [Khan Academy: Graphing Inequalities](#).

QUESTION NO: 12

Empirical probability is based on:

- A. Observed frequency of events

- B. Classical calculation
- C. Subjective guess
- D. Theoretical models

ANSWER: A

Explanation:

Observed frequency of events is correct because empirical probability is calculated from actual data collected through observations, experiments, or repeated trials. It uses the relative frequency formula: the number of times an event occurs divided by the total number of observations or trials. For example, if a quality-control team inspects 500 products and finds 15 defective items, the empirical probability of a defect is $15/500$, or 0.03. This estimate is grounded in what happened in the observed sample rather than in assumptions about equally likely outcomes or personal judgment. As more relevant observations are gathered under consistent conditions, the relative frequency often becomes a more stable estimate of the event's long-run probability. This connection between repeated trials and long-run relative frequency is a central idea in introductory probability and statistics. OpenStax describes experimental probability as probability determined through repeated trials and observed outcomes; see [OpenStax Introductory Statistics](#). The NIST/SEMATECH e-Handbook also discusses estimating probabilities from observed frequencies in data: [NIST e-Handbook](#).

QUESTION NO: 13

A hospital analyzes the recovery rates of patients undergoing two different treatments for a specific condition. Treatment X has an overall recovery rate of 80%, while Treatment Y has a recovery rate of 90%. However, when the data is divided by age groups, Treatment X has a higher recovery rate in each age group compared to Treatment Y.

Is Simpson's paradox present in this study?

- A. No, because the study does not involve a confounding variable that affects recovery rates.
- B. No, because the overall recovery rate is higher for Treatment Y.
- C. Yes, because the recovery rates differ between the two treatments.
- D. Yes, because the overall recovery rates are inconsistent with the age group-specific rates.

ANSWER: D

Explanation:

Yes, because the overall recovery rates are inconsistent with the age group-specific rates. Simpson's paradox occurs when a comparison based on combined, or aggregated, data gives an apparent relationship that reverses after the data are separated into meaningful subgroups. Here, the aggregated recovery rates favor Treatment Y: 90% versus 80%. Yet, within every age group, Treatment X has the higher recovery rate. Thus, the conclusion from the overall percentages conflicts with the conclusion reached from every age-specific comparison, which is the defining feature of Simpson's paradox.

Age is the relevant grouping variable because recovery likelihood can vary substantially by age, and the treatments may have been given to patients with different age distributions. For example, Treatment X could have been used more often among older patients, who may have lower recovery rates overall, while Treatment Y could have been used more often among younger patients. Combining these unequal groups can make Treatment Y appear superior overall even though Treatment X performs better within each age category. Interpreting the stratified comparisons is therefore essential when age is associated with both treatment assignment and recovery outcomes. See [Encyclopaedia Britannica's overview of Simpson's paradox](#) and the [NIST/SEMATECH e-Handbook of Statistical Methods](#).

QUESTION NO: 14

Type II error = ?

- A. Accept H_0 when false

- B. Reject H_0 when true
- C. Correct decision
- D. None

ANSWER: A

Explanation:

“Accept H_0 when false” represents the conventional introductory-statistics description of a Type II error. More precisely, a Type II error occurs when the null hypothesis is actually false but the hypothesis test does not reject it. In other words, the test result fails to identify a real effect, difference, or association that exists in the population. This outcome is also called a false negative.

The probability of making this error is denoted by β . Its complement, $1 - \beta$, is statistical power: the probability that a test correctly rejects a false null hypothesis. Power is influenced by factors such as sample size, variability, the chosen significance level, and the true magnitude of the effect. For example, if a treatment genuinely changes an outcome but a study does not produce sufficient evidence to reject the null hypothesis, the study has made a Type II error. This framing is consistent with standard hypothesis-testing terminology used in introductory statistics. See the [NIST/SEMATECH e-Handbook discussion of Type I and Type II errors](#) and the [Penn State STAT 200 hypothesis-testing overview](#).

QUESTION NO: 15

A tech company surveys its employees to assess remote work preferences. The survey reveals that 55% of employees prefer working remotely, while the remaining employees prefer working in the office. Of the employees who prefer remote work, 75% report high productivity. Among those who prefer working in the office, only 50% report high productivity.

Which percentage of employees prefer working in the office and report high productivity?

- A. 22.5%
- B. 27.5%
- C. 45%
- D. 50%

ANSWER: A

Explanation:

22.5% is correct because the question asks for the percentage of all employees who meet two conditions simultaneously: they prefer working in the office and they report high productivity. First, the percentage who prefer office work is the complement of the 55% who prefer remote work: $100\% - 55\% = 45\%$. Next, 50% of that office-preferring group reports high productivity. This is a conditional proportion, so it must be applied to the 45% office-preference group rather than treated as a percentage of all employees. The joint percentage is therefore calculated by multiplying the probability of office preference by the conditional probability of high productivity given office preference: $0.45 \times 0.50 = 0.225$. Converting 0.225 to a percentage gives 22.5%. For example, in a group of 100 employees, 45 would prefer office work, and half of those 45 employees—22.5 employees per 100, expressed as an expected percentage—would report high productivity. This multiplication rule for a joint probability follows the conditional-probability relationship $P(\text{office and high productivity}) = P(\text{office}) \times P(\text{high productivity} \mid \text{office})$. See the [NIST Engineering Statistics Handbook discussion of conditional probability](#) and [OpenStax probability rules](#).

QUESTION NO: 16

Mean = 100, SD = 15, $z = 2$. Value = ?

- A. 130

- B. 115
- C. 125
- D. 110

ANSWER: A

Explanation:

The correct value is **130**. A z-score indicates a raw value's distance from the mean in units of standard deviations. To convert a z-score to its corresponding raw value, use the rearranged z-score formula: $x = \mu + z\sigma$, where μ is the mean, z is the z-score, and σ is the standard deviation. With a mean of 100, standard deviation of 15, and z-score of 2, the calculation is $x = 100 + (2 \times 15)$. Two standard deviations equal 30 points, so adding 30 to the mean yields 130. Thus, 130 is the score located two standard deviations above a mean of 100. This conversion is useful when interpreting standardized scores and locating observations relative to the center of a distribution. For a definition and calculation of z-scores, see [NIST's discussion of the standard score](#). Additional guidance on the relationship between z-scores, means, and standard deviations is available from [Statistics How To](#).

QUESTION NO: 17

A university surveys faculty and students to determine support for a new campus recycling initiative.

The results of the survey are shown in the following 2×2 contingency table:

	Support	Oppose	Total
Faculty	40	10	50
Students	200	300	500

Which statement is true?

- A. Faculty members are slightly more likely than students to support the recycling initiative.
- B. Faculty members are much more likely than students to support the recycling initiative.
- C. Faculty members are slightly less likely than students to support the recycling initiative.
- D. Faculty members are much less likely than students to support the recycling initiative.

ANSWER: B

Explanation:

Faculty members are much more likely than students to support the recycling initiative. The appropriate comparison is the conditional proportion supporting the initiative within each surveyed group, because the question asks how support differs between faculty and students. Among faculty, 40 of 50 respondents support the initiative, giving a support rate of $40/50 = 0.80$, or 80%. Among students, 200 of 500 respondents support it, giving a support rate of $200/500 = 0.40$, or 40%.

The faculty support rate is therefore 40 percentage points higher than the student support rate. It is also twice the student rate, since 80% is two times 40%. This is a substantial practical difference, supporting the description that faculty are much more likely to favor the initiative. A two-way contingency table summarizes the relationship between two categorical variables—in this case, respondent group and support status—and conditional relative frequencies allow meaningful comparisons across groups of different sizes. This approach follows standard guidance for interpreting categorical data using two-way tables; see [OpenStax, Two-Way Tables](#) and [Penn State STAT 200: Two-Way Tables](#).