

AP Statistics Practice Test (Sample)

Study Guide



Everything you need from our exam experts!

Copyright © 2025 by Examzify - A Kaluba Technologies Inc. product.

ALL RIGHTS RESERVED.

No part of this book may be reproduced or transferred in any form or by any means, graphic, electronic, or mechanical, including photocopying, recording, web distribution, taping, or by any information storage retrieval system, without the written permission of the author.

Notice: Examzify makes every reasonable effort to obtain from reliable sources accurate, complete, and timely information about this product.

SAMPLE

Questions

- 1. What does R-squared indicate in a regression model?**
 - A. The accuracy of the predictions made by the model**
 - B. The relationship between independent and dependent variables**
 - C. The proportion of variance in the dependent variable explained by the independent variable(s)**
 - D. The standard deviation of the residuals**
- 2. Which of the following is not a source of caution in regression analysis between two variables?**
 - A. Outliers**
 - B. Multicollinearity**
 - C. Non-linearity**
 - D. All of these are potential problems**
- 3. If the slope of the line predicting SAT score from family income is 6.25 points per \$1000, what is the slope when predicting family income from SAT score?**
 - A. 0.0625**
 - B. 0.037**
 - C. 1.5**
 - D. 15**
- 4. What effect does removing a point in the upper right corner of a scatterplot have on the slope of the line of best fit and the correlation?**
 - A. The slope will increase, and correlation will decrease**
 - B. The slope will decrease, and correlation will increase**
 - C. Both slope and correlation will increase**
 - D. Both slope and correlation will decrease**
- 5. What might a non-random pattern in a residuals plot suggest?**
 - A. The model may be inappropriate**
 - B. The data points are perfectly linear**
 - C. The errors are consistently minimal**
 - D. The relationship is purely random**

- 6. In a properly designed experiment, what is the purpose of random assignment?**
- A. To determine the outcome variable**
 - B. To ensure the manipulation of dependent variables**
 - C. To eliminate bias and control for confounding variables**
 - D. To establish a relationship between two variables**
- 7. Explain the term "outlier."**
- A. A data point that significantly differs from other observations in the dataset**
 - B. An average value calculated from a set of observations**
 - C. A repeated value within a dataset**
 - D. A strategy used to eliminate extreme variations in data**
- 8. In the context of experiments, what are places used for?**
- A. Data collection**
 - B. Blinding**
 - C. Sample selection**
 - D. Randomization**
- 9. How do you find the mode of a data set?**
- A. By calculating the average of all values**
 - B. By identifying the value that appears most frequently in the dataset**
 - C. By finding the middle value when data is ordered**
 - D. By determining the range of the dataset**
- 10. In a regression analysis with $r^2 = 0.72$ for profits and advertising spending, what can be concluded?**
- A. 72% of the variability is explained**
 - B. There is no correlation**
 - C. The linear model is inappropriate**
 - D. The slope is negative**

Answers

SAMPLE

1. C
2. D
3. B
4. B
5. A
6. C
7. A
8. B
9. B
10. A

SAMPLE

Explanations

SAMPLE

1. What does R-squared indicate in a regression model?

- A. The accuracy of the predictions made by the model**
- B. The relationship between independent and dependent variables**
- C. The proportion of variance in the dependent variable explained by the independent variable(s)**
- D. The standard deviation of the residuals**

R-squared, also known as the coefficient of determination, is a key metric in regression analysis that quantifies how well the independent variable(s) account for the variability in the dependent variable. Specifically, it represents the proportion of the total variance in the dependent variable that can be explained by the model's independent variable(s). A higher R-squared value implies that a greater proportion of variance is explained by the model, indicating a potentially better fit. It ranges from 0 to 1, where 0 means that the model explains none of the variability, and 1 indicates that it accounts for all of it. This makes R-squared a valuable indicator of the effectiveness of the model in explaining the variation in the outcome being measured. Understanding R-squared helps in evaluating the strength of the relationship between the independent and dependent variables. While it does provide insight into the relationship, its primary role remains rooted in variance explanation rather than merely indicating the strength or accuracy of predictions. Other options touch on related concepts, but without capturing the essence of R-squared's role in the context of statistical modeling.

2. Which of the following is not a source of caution in regression analysis between two variables?

- A. Outliers**
- B. Multicollinearity**
- C. Non-linearity**
- D. All of these are potential problems**

In regression analysis, caution is necessary when interpreting the relationships between variables, as certain conditions can skew results or lead to misleading conclusions. While each of the mentioned issues—outliers, multicollinearity, and non-linearity—can indeed present concerns in regression analysis, the statement that all of these are potential problems is valid. Outliers can significantly impact the slope of the regression line, influence the correlation coefficient, and overall distort the representation of the data. Multicollinearity refers to the situation in which independent variables in the regression are highly correlated, which can lead to difficulty in estimating the coefficients accurately or in assessing the effect of each variable. Non-linearity indicates that the relationship between the independent and dependent variables is not a straight line, which can render linear regression analysis inappropriate. Therefore, recognizing these issues is essential for proper interpretation and understanding of regression results. Choosing the option that states "all of these are potential problems" highlights the understanding that caution is warranted when addressing any of the identified issues in the context of regression analysis between two variables.

3. If the slope of the line predicting SAT score from family income is 6.25 points per \$1000, what is the slope when predicting family income from SAT score?

A. 0.0625

B. 0.037

C. 1.5

D. 15

To find the slope when predicting family income from SAT score, it's essential to understand the relationship between the two variables and how the slope of a regression line works. The slope of your initial relationship tells you how much the SAT score is predicted to change for a given increase in family income. Specifically, a slope of 6.25 means that for every additional \$1000 of family income, the SAT score is predicted to increase by 6.25 points. When flipping the predictors, the slope will also change in a way that relates to the reciprocal of the initial slope. To derive the new slope, we first express the original slope in terms of points per dollar. Since 6.25 points are gained for every \$1000 of income, we convert this to a per dollar basis by dividing by 1000, which gives us 0.00625 points per dollar. This reflects how much the SAT score increases for each dollar increase in income. Now, to predict family income from SAT score, we consider the inverse relationship. The slope for family income would be the reciprocal of the SAT score's increase per dollar of income. Thus, we take the inverse of 0.00625, resulting in a slope of 160 for income

4. What effect does removing a point in the upper right corner of a scatterplot have on the slope of the line of best fit and the correlation?

A. The slope will increase, and correlation will decrease

B. The slope will decrease, and correlation will increase

C. Both slope and correlation will increase

D. Both slope and correlation will decrease

Removing a point in the upper right corner of a scatterplot typically results in a change to the overall trend represented by the remaining data points. When a high-value point, often seen as an outlier, is removed, it can have a reducing effect on the slope of the line of best fit. This happens because that upper right corner point usually contributes positively to both the slope and the correlation due to its high values on both axes. By eliminating this point, the remaining data may have a lower average x-value relative to y-value, which thus leads to a decrease in the overall slope. As for the correlation coefficient, it measures the strength and direction of a linear relationship between two variables. The removal of the upper corner point can often result in a tighter cluster of points with less variability, thereby increasing the strength of the linear relationship among the remaining points. Consequently, as this higher correlation suggests a more consistent linear relationship, the correlation value tends to increase. Thus, the action of removing a point in the upper right corner decreases the slope of the line of best fit while increasing the strength of the correlation among the remaining data points.

5. What might a non-random pattern in a residuals plot suggest?

- A. The model may be inappropriate**
- B. The data points are perfectly linear**
- C. The errors are consistently minimal**
- D. The relationship is purely random**

A non-random pattern in a residuals plot suggests that the model used to fit the data may be inappropriate. In a well-fitted regression model, the residuals—which are the differences between the observed values and the values predicted by the model—should display a random pattern when plotted against predicted values or against the independent variable. A random pattern indicates that the model has captured the underlying relationship well, and any randomness in the data is expected. However, if there is a discernible pattern in the residuals (such as a curve or clustering), it can indicate that the model is not adequately explaining the data. This might mean that the model is missing key variables, has the wrong functional form, or may require transformation of the variables to better fit the data. Thus, recognizing a non-random pattern in residuals is critical for diagnosing and improving the model's performance.

6. In a properly designed experiment, what is the purpose of random assignment?

- A. To determine the outcome variable**
- B. To ensure the manipulation of dependent variables**
- C. To eliminate bias and control for confounding variables**
- D. To establish a relationship between two variables**

Random assignment is a fundamental technique used in experimental design to ensure that participants are evenly and randomly distributed across different treatment groups. The primary purpose of random assignment is to eliminate bias and control for confounding variables. By randomly assigning subjects to various groups, researchers can minimize the impact of external factors that could influence the results, thereby increasing the internal validity of the experiment. This process helps to ensure that the groups being compared are statistically similar before any treatment or manipulation is applied. As a result, any observed differences in outcomes can be more confidently attributed to the treatment itself rather than to pre-existing differences between the groups. This aspect of random assignment is crucial for establishing cause-and-effect relationships, as it strengthens the conclusion that changes in the independent variable directly lead to changes in the dependent variable, rather than being influenced by other, uncontrolled factors.

7. Explain the term "outlier."

- A. A data point that significantly differs from other observations in the dataset**
- B. An average value calculated from a set of observations**
- C. A repeated value within a dataset**
- D. A strategy used to eliminate extreme variations in data**

The term "outlier" refers to a data point that significantly differs from other observations in the dataset. This could mean that its value is much higher or much lower than most of the data points, which can affect statistical analyses and interpretations. Outliers can arise due to variability in the data or may indicate that there are measurement errors or special conditions affecting those data points. Identifying outliers is crucial in statistics because they can skew results, impact the mean and standard deviation, and influence the conclusions drawn from the data analysis. For example, if an outlier is present in a dataset of heights, it could represent an exceptionally tall or short individual, which would not reflect the central tendency of the group. While the other options provide different statistical concepts, they do not accurately capture the definition of an outlier. Concepts such as averages, repeated values, or strategies to eliminate variations focus on different aspects of data analysis rather than identifying points that stand out significantly from a dataset.

8. In the context of experiments, what are places used for?

- A. Data collection**
- B. Blinding**
- C. Sample selection**
- D. Randomization**

In the context of experiments, places, often referred to as "blinding," serve the essential function of reducing bias. Blinding refers to preventing participants or researchers from knowing specific details about the treatment groups. For example, in a double-blind study, neither the participants nor the researchers know which individuals are receiving the treatment and which are receiving a placebo. This helps ensure that the results are not influenced by the participants' or researchers' expectations or perceptions, thereby leading to more reliable and valid results. The practice of blinding is crucial because it helps maintain the objectivity of the outcomes. Without blinding, there is a risk that biases could affect how results are recorded or interpreted, ultimately impacting the conclusions drawn from the experiment. Therefore, places used for blinding play a significant role in the integrity of experimental results.

9. How do you find the mode of a data set?

- A. By calculating the average of all values
- B. By identifying the value that appears most frequently in the dataset**
- C. By finding the middle value when data is ordered
- D. By determining the range of the dataset

To find the mode of a data set, you identify the value that occurs with the highest frequency. The mode is simply the number that appears most often in the dataset. It can be useful in various statistical analyses, particularly when you are dealing with categorical data where you want to know which category is the most common. For example, if you have a data set of shoe sizes worn by a group of people, the mode would be the shoe size that appears the most within that group. It is important to note that a data set can have one mode, more than one mode (bimodal or multimodal), or no mode at all if all values appear with the same frequency. The other methods listed do not correlate with finding the mode. Calculating the average refers to determining the mean, finding the middle value pertains to the median, and determining the range involves finding the difference between the maximum and minimum values in the dataset. These concepts serve different statistical purposes and should not be confused with finding the mode.

10. In a regression analysis with $r^2 = 0.72$ for profits and advertising spending, what can be concluded?

- A. 72% of the variability is explained**
- B. There is no correlation
- C. The linear model is inappropriate
- D. The slope is negative

When $r^2 = 0.72$, it indicates that 72% of the variability in the dependent variable, which in this case is profits, can be explained by the independent variable, advertising spending. This means that a significant portion of the variation in profits can be accounted for by changes in advertising spending, suggesting a strong relationship between the two variables within the context provided. This metric serves as a gauge of the explanatory power of the regression model, confirming that a substantial amount of the data variability is captured by the model. In essence, the high r^2 value suggests that the regression model is effective in explaining the relationship between advertising expenditures and profits.