

SOCIETY OF ACTUARIES (SOA) PA PRACTICE EXAM (SAMPLE)

STUDY GUIDE



Prepare with structure. Pass with confidence



Copyright © 2026 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 accurate, complete, and timely information about this product from reliable sources.

SAMPLE

Table of Contents

Copyright	1
Table of Contents	2
Introduction	3
How to Use This Guide	4
Questions	5
Answers	8
Explanations	10
Next Steps	16

SAMPLE

Introduction

Preparing for a certification exam can feel overwhelming, but with the right tools, it becomes an opportunity to build confidence, sharpen your skills, and move one step closer to your goals. At Examzify, we believe that effective exam preparation isn't just about memorization, it's about understanding the material, identifying knowledge gaps, and building the test-taking strategies that lead to success.

This guide was designed to help you do exactly that.

Whether you're preparing for a licensing exam, professional certification, or entry-level qualification, this book offers structured practice to reinforce key concepts. You'll find a wide range of multiple-choice questions, each followed by clear explanations to help you understand not just the right answer, but why it's correct.

The content in this guide is based on real-world exam objectives and aligned with the types of questions and topics commonly found on official tests. It's ideal for learners who want to:

- Practice answering questions under realistic conditions,
- Improve accuracy and speed,
- Review explanations to strengthen weak areas, and
- Approach the exam with greater confidence.

We recommend using this book not as a stand-alone study tool, but alongside other resources like flashcards, textbooks, or hands-on training. For best results, we recommend working through each question, reflecting on the explanation provided, and revisiting the topics that challenge you most.

Remember: successful test preparation isn't about getting every question right the first time, it's about learning from your mistakes and improving over time. Stay focused, trust the process, and know that every page you turn brings you closer to success.

Let's begin.

How to Use This Guide

This guide is designed to help you study more effectively and approach your exam with confidence. Whether you're reviewing for the first time or doing a final refresh, here's how to get the most out of your Examzify study guide:

1. Start with a Diagnostic Review

Skim through the questions to get a sense of what you know and what you need to focus on. Your goal is to identify knowledge gaps early.

2. Study in Short, Focused Sessions

Break your study time into manageable blocks (e.g. 30 – 45 minutes). Review a handful of questions, reflect on the explanations.

3. Learn from the Explanations

After answering a question, always read the explanation, even if you got it right. It reinforces key points, corrects misunderstandings, and teaches subtle distinctions between similar answers.

4. Track Your Progress

Use bookmarks or notes (if reading digitally) to mark difficult questions. Revisit these regularly and track improvements over time.

5. Simulate the Real Exam

Once you're comfortable, try taking a full set of questions without pausing. Set a timer and simulate test-day conditions to build confidence and time management skills.

6. Repeat and Review

Don't just study once, repetition builds retention. Re-attempt questions after a few days and revisit explanations to reinforce learning. Pair this guide with other Examzify tools like flashcards, and digital practice tests to strengthen your preparation across formats.

There's no single right way to study, but consistent, thoughtful effort always wins. Use this guide flexibly, adapt the tips above to fit your pace and learning style. You've got this!

Questions

SAMPLE

- 1. When examining bivariate relationships, why is it important to check for extreme outliers?**
 - A. They can indicate potential data entry errors
 - B. They always enhance analysis quality
 - C. They are automatically excluded by software
 - D. They typically improve model accuracy
- 2. Which of the following R commands is used to graph the ROC curve?**
 - A. `plot(roc)`
 - B. `roc(data)`
 - C. `curve(roc)`
 - D. `draw(roc)`
- 3. In the context of decision trees, what does a higher cp value indicate?**
 - A. A more complex tree
 - B. A less complex tree
 - C. No effect on tree complexity
 - D. Increased predictive accuracy
- 4. Which of the following is a key assumption of OLS Regression?**
 - A. The response variable follows a uniform distribution.
 - B. The conditional distribution of the response variable is normal.
 - C. The response variable is always positive.
 - D. The response variable must be discrete.
- 5. What is a key outcome of balancing the class distribution with techniques like undersampling or oversampling?**
 - A. Train the model on a more diverse dataset
 - B. Maintain original class proportions
 - C. Decrease prediction accuracy
 - D. Prioritize majority class representation

- 6. What does Cook's distance measure in a Residuals vs Leverage graph?**
- A. Effectiveness of the model
 - B. Influential points relative to model fitting
 - C. Variance of residuals
 - D. Distribution of selected predictors
- 7. Which limitation is associated with feature selection through regularization?**
- A. It is an easily interpretable technique
 - B. It is dependent on the model used
 - C. It guarantees cross-validation success
 - D. It always retains all variables in the model
- 8. What term is used to describe the sections of the decision tree that are non-critical and can be removed?**
- A. Splitting
 - B. Overfitting
 - C. Pruning
 - D. Branching
- 9. What effect do undersampling and oversampling have on predictions?**
- A. They worsen the predictions for the minority class
 - B. They balance the performance between minority and majority classes
 - C. They make predictions more unreliable
 - D. They have no effect on the predictions
- 10. Which function is used to convert a continuous variable into a factor variable based on specified buckets?**
- A. cut
 - B. factor
 - C. as.factor
 - D. dummyVars

Answers

SAMPLE

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

SAMPLE

Explanations

SAMPLE

1. When examining bivariate relationships, why is it important to check for extreme outliers?

- A. They can indicate potential data entry errors**
- B. They always enhance analysis quality
- C. They are automatically excluded by software
- D. They typically improve model accuracy

In the context of examining bivariate relationships, checking for extreme outliers is crucial because they can indicate potential data entry errors. Outliers may represent values that are significantly different from the rest of the data set, which can skew results and lead to misleading conclusions. Identifying these outliers allows analysts to investigate whether they are indeed valid observations or the result of errors in data collection or entry processes. While outliers can sometimes be valid and carry important information, they should be scrutinized carefully. They can disproportionately influence statistical measures, such as the mean and correlation coefficients, which could ultimately misrepresent the underlying relationship being studied. Therefore, recognizing that outliers can signal inaccuracies in the data encourages more careful validation of the dataset prior to proceeding with further analysis. This careful scrutiny helps ensure that the conclusions drawn from the analysis are based on accurate and representative data.

2. Which of the following R commands is used to graph the ROC curve?

- A. `plot(roc)`**
- B. `roc(data)`
- C. `curve(roc)`
- D. `draw(roc)`

The command to graph the ROC (Receiver Operating Characteristic) curve is indeed achieved by using the function `plot()` on the ROC object that was created using the `roc()` function from the `pROC` package in R. The `roc()` function calculates necessary statistics and generates an object containing the data needed to visualize the ROC curve. Once the ROC object is created, calling `plot()` on it effectively generates the graphical representation of the curve. This approach separates the data preparation from the visualization process, making it clearer and more modular. Essentially, the `roc(data)` function prepares the ROC data that will be plotted, and `plot(roc)` is what you use to visualize this data on a graph. This clear distinction allows for more flexibility in handling ROC analysis and plotting in R. The other options do not correspond to the correct method of graphing the ROC curve in R. For example, `roc(data)` is meant for calculating the ROC statistics, not for plotting. Similarly, `curve(roc)` and `draw(roc)` are not standard functions in R used for creating ROC curves, thereby making those incorrect choices for the task at hand.

3. In the context of decision trees, what does a higher cp value indicate?

- A. A more complex tree
- B. A less complex tree**
- C. No effect on tree complexity
- D. Increased predictive accuracy

In the context of decision trees, a higher cp value, known as the complexity parameter, indicates a less complex tree. The cp value serves as a penalty for the number of splits in the tree: as this value increases, it restricts the tree from growing too complex. Essentially, when the cp value is set to a higher level, the tree will tend to be pruned more aggressively, leading to fewer branches and splits, which simplifies the model. This relationship between cp and tree complexity is crucial because a simpler tree often helps to mitigate overfitting, allowing the model to generalize better to unseen data. Therefore, selecting an appropriate cp value is an important step in building a decision tree that balances bias and variance, ultimately aiming for optimal predictive performance.

4. Which of the following is a key assumption of OLS Regression?

- A. The response variable follows a uniform distribution.
- B. The conditional distribution of the response variable is normal.**
- C. The response variable is always positive.
- D. The response variable must be discrete.

In ordinary least squares (OLS) regression, one of the primary assumptions is that the conditional distribution of the response variable, given the explanatory variables, is normally distributed. This assumption is crucial because OLS regression relies on the normality of the errors (the differences between observed and predicted values) to ensure that the estimates of the coefficients are unbiased, efficient, and have the minimum variance among all linear estimators. Normality of the conditional distribution is important for valid hypothesis testing and for constructing confidence intervals for the regression coefficients. If this assumption holds, it means that for any given set of independent variables, the distribution of the predicted values will be normally distributed around the true regression line, which is foundational for making statistical inferences about the model parameters. Other options do not accurately represent key assumptions underlying OLS regression. For instance, there is no requirement for the response variable to follow a uniform distribution, nor is there a restriction that it must always be positive or discrete. The response variable can be continuous and can take any real value, without necessarily being constrained to specific types of distributions.

5. What is a key outcome of balancing the class distribution with techniques like undersampling or oversampling?

- A. Train the model on a more diverse dataset**
- B. Maintain original class proportions
- C. Decrease prediction accuracy
- D. Prioritize majority class representation

Balancing the class distribution through techniques such as undersampling or oversampling leads to training the model on a more diverse dataset. This diversity is crucial because it helps ensure that the model learns from a more representative sample of both the majority and minority classes. When you balance the classes, you mitigate biases that might exist due to an uneven distribution, which often leads to improved generalization when the model is applied to new, unseen data. In undersampling, you reduce the number of instances in the majority class, thereby giving the minority class a better chance to be represented more fully in the training data. In contrast, oversampling involves duplicating instances in the minority class or generating synthetic examples. Both techniques aim to create a dataset where minority classes are more prominently represented, leading to a more balanced and diverse dataset that enhances the model's ability to learn the characteristics of all classes equally well. Ultimately, this approach contributes to better predictive performance, particularly for the minority class, which often suffers when class imbalance is present. Thus, it allows the model to make accurate predictions across different classes, demonstrating why training on a more diverse dataset is a key outcome of these balancing techniques.

6. What does Cook's distance measure in a Residuals vs Leverage graph?

- A. Effectiveness of the model
- B. Influential points relative to model fitting**
- C. Variance of residuals
- D. Distribution of selected predictors

Cook's distance measures the influence of each data point on the overall regression model. In the context of a Residuals vs Leverage graph, it specifically identifies points that significantly affect the fitted model. High values of Cook's distance indicate that a data point has the potential to disproportionately affect the estimates of the regression coefficients. This means that such points may be influential in determining the shape and position of the regression line, suggesting they should be examined more closely. By analyzing Cook's distance, practitioners can effectively identify which observations may be outliers or leverage points that could unduly influence the results, allowing for better model diagnostics and potentially more robust conclusions.

7. Which limitation is associated with feature selection through regularization?

- A. It is an easily interpretable technique
- B. It is dependent on the model used**
- C. It guarantees cross-validation success
- D. It always retains all variables in the model

Feature selection through regularization is a process where the introduction of penalties on the size of coefficients in a model encourages sparsity, meaning it may push some coefficients to zero. The dependence on the model used refers to the fact that different types of regularization, such as LASSO (which can set some coefficients to exactly zero) or Ridge (which shrinks coefficients but does not set them to zero), will yield different sets of selected features. Therefore, the outcome of feature selection cannot be universally applied across all contexts but rather is influenced by the specific regularization technique chosen and the characteristics of the model being used. Each regularization method addresses feature selection differently, thus emphasizing the importance of the model in determining the final set of selected features. Other options do not accurately capture the limitations associated with regularization in the context of feature selection. Regularization methods can vary in their interpretability depending on the model and the approach used. They do not inherently guarantee cross-validation success, as model performance can still vary with different datasets and splits. Furthermore, regularization methods do not always retain all variables in the model, as they can effectively reduce the number of features by zeroing out the coefficients of certain variables.

8. What term is used to describe the sections of the decision tree that are non-critical and can be removed?

- A. Splitting
- B. Overfitting
- C. Pruning**
- D. Branching

The term that describes the sections of a decision tree that are non-critical and can be removed is known as pruning. Pruning is a technique used in decision tree algorithms to enhance the model's performance by reducing complexity and avoiding overfitting. When a decision tree is built, it might create many branches based on the training data, leading to a model that fits the noise rather than the underlying pattern. Pruning helps address this issue by removing branches that offer little improvement to the model's predictive power. This results in a simpler, more interpretable model that generalizes better to new data. In the context of decision trees, pruning can involve removing branches that do not contribute significantly to the accuracy of predictions, thereby streamlining the decision-making process without sacrificing model effectiveness. This is crucial for improving the model's performance on unseen data and enhancing its applicability in real-world scenarios.

9. What effect do undersampling and oversampling have on predictions?

- A. They worsen the predictions for the minority class
- B. They balance the performance between minority and majority classes**
- C. They make predictions more unreliable
- D. They have no effect on the predictions

The correct answer highlights the utility of both undersampling and oversampling in addressing class imbalance in datasets. When dealing with skewed distributions of data – where one class, often the minority class, is significantly less represented than the others – these techniques help to balance the representation of each class. Oversampling involves increasing the number of instances in the minority class, often by duplicating existing instances or creating synthetic data points. This approach provides the model with more information about the minority class, helping it to better learn its characteristics and improve prediction accuracy for that class. On the other hand, undersampling reduces the number of instances in the majority class to better balance the dataset. This can help to prevent the model from being biased toward the majority class and improves the model's ability to recognize patterns from the minority class. Both methods aim to create a more balanced dataset that enables the predictive model to perform optimally across all classes, enhancing overall predictive accuracy and ensuring the model does not disproportionately favor the majority class. As a result, the prediction performance improves, especially for the minority class, which is crucial in many applications, such as fraud detection or disease diagnosis, where identifying rare events is vital. These methods can lead to improved predictive performance, particularly in scenarios where class imbalance

10. Which function is used to convert a continuous variable into a factor variable based on specified buckets?

- A. cut**
- B. factor
- C. as.factor
- D. dummyVars

The function that is used to convert a continuous variable into a factor variable based on specified buckets is the cut function. This function takes a numeric vector and divides it into intervals or "bins," allowing you to categorize continuous data into discrete groups. When you apply the cut function, you specify breaks (the points at which the data is divided), and it assigns factor levels to the data based on which interval each value falls into. This is particularly useful in statistical analysis and data visualization, where you may want to analyze trends or differences across groups that are defined by ranges of a continuous variable. In contrast, the other options serve different purposes. The factor function is used to create a factor variable from a vector of values but does not inherently create categories from continuous data. Similarly, as.factor is a base R function that converts an input into a factor but does not define the intervals for a continuous variable on its own. DummyVars is utilized for creating dummy variables for categorical encoding, primarily in the context of modeling and is not intended for segmenting continuous data into factor levels.

SAMPLE

Next Steps

Congratulations on reaching the final section of this guide. You've taken a meaningful step toward passing your certification exam and advancing your career.

As you continue preparing, remember that consistent practice, review, and self-reflection are key to success. Make time to revisit difficult topics, simulate exam conditions, and track your progress along the way.

If you need help, have suggestions, or want to share feedback, we'd love to hear from you. Reach out to our team at hello@examzify.com.

Or visit your dedicated course page for more study tools and resources:

<https://soa-pa.examzify.com>

We wish you the very best on your exam journey. You've got this!

SAMPLE