

EVIDENCE-BASED PRACTICE (EBP) II PRACTICE EXAM (SAMPLE)

STUDY GUIDE



SAMPLE STUDY GUIDE WITH LIMITED QUESTIONS

VISIT US ONLINE FOR
THE FULL VERSION STUDY GUIDE

<https://ebp2.examzify.com>



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	15

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. Relative risk reduction is best described as:**
 - A. The magnitude of the association between exposure and disease, showing the likelihood that the exposed group will develop the disease relative to those not exposed
 - B. The difference in incidence between exposed and unexposed
 - C. The reciprocal of ARR
 - D. The absolute risk difference between groups
- 2. The Friedman test is the nonparametric alternative to which parametric test?**
 - A. Independent t-test
 - B. One-way repeated measures ANOVA
 - C. Two-way ANOVA
 - D. Paired t-test
- 3. If a study reports a p-value greater than 0.05, what does this indicate?**
 - A. The result is statistically significant.
 - B. The result is not statistically significant.
 - C. The study has high statistical power.
 - D. The confidence interval is narrow.
- 4. In evaluating the development of a CPR, which criterion ensures that all relevant predictive factors are included during development?**
 - A. Did investigators compare results to a criterion test?
 - B. Were the predictive factors operationally defined?
 - C. Was the sample size large enough to accommodate the number of predictor factors?
 - D. Did investigators include all relevant predictive factors in the development process?
- 5. Which statement accurately defines inter-rater reliability?**
 - A. The same rater tests the same patient on different occasions
 - B. Two or more different raters test the same patient
 - C. The same test is repeated with the same rater
 - D. The patient is tested twice with different index tests

- 6. Prognosis is estimated by?**
- A. Response to questions related to disease-specific factors/demographics
 - B. Medical comorbidities
 - C. Behavioral comorbidities
 - D. Double-blind randomization
- 7. Which statement best describes the role of reports on study quality in evidence synthesis?**
- A. They determine the sequencing of trials in the database
 - B. They influence interpretation of overall findings and confidence in effect estimates
 - C. They replace risk of bias assessment
 - D. They are optional if the meta-analysis is significant
- 8. In assessing credibility of self-reported measures, which criterion involves documenting how items were developed and tested?**
- A. Adequate description of survey development process
 - B. Administration and scoring requirements
 - C. Instrument reliability
 - D. Is instrument valid
- 9. Which level of measurement is defined by ordering without assuming equal intervals between the levels?**
- A. Interval
 - B. Nominal
 - C. Ratio
 - D. Ordinal
- 10. If a study uses more than one administration method for a self-reported instrument, what should researchers do?**
- A. Use only the mode with the highest reliability
 - B. Reexamine measurement properties for each mode
 - C. Ignore mode differences
 - D. Rely on the original validation

Answers

SAMPLE

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

SAMPLE

Explanations

SAMPLE

1. Relative risk reduction is best described as:

- A. The magnitude of the association between exposure and disease, showing the likelihood that the exposed group will develop the disease relative to those not exposed**
- B. The difference in incidence between exposed and unexposed
- C. The reciprocal of ARR
- D. The absolute risk difference between groups

Relative risk reduction hinges on comparing risk between two groups and expressing how much risk is lowered in the treated or exposed group relative to the control group. The best description among the options is the one that talks about the magnitude of the association between exposure and disease, i.e., how likely the exposed group is to develop the disease compared with those not exposed. This captures the idea of a ratio of risks (the relative risk), which underpins how we think about reductions in risk when calculating RRR. The other choices describe the absolute difference in risk (the risk difference or ARR) or its reciprocal, which are not what relative risk reduction measures. In practice, you get RRR by taking 1 minus the relative risk, so recognizing the comparison between groups is the core concept.

2. The Friedman test is the nonparametric alternative to which parametric test?

- A. Independent t-test
- B. One-way repeated measures ANOVA**
- C. Two-way ANOVA
- D. Paired t-test

The main idea here is that the Friedman test serves as the nonparametric counterpart to the one-way repeated measures ANOVA. It is used when you have more than two related or matched conditions measured on the same subjects and you cannot assume that the data are normally distributed. Instead of using raw scores, the Friedman test works with ranks across those related conditions, preserving the within-subjects nature of the design while relaxing the normality assumption. This makes it the appropriate nonparametric alternative for detecting overall differences across multiple related conditions. The other options don't fit because the independent t-test compares two unrelated groups, and the paired t-test compares two related conditions (only two levels). Two-way ANOVA involves two factors (and can include interactions), whereas Friedman handles a single within-subjects factor with multiple levels. For more than two related conditions with nonparametric data, Friedman is the right choice.

3. If a study reports a p-value greater than 0.05, what does this indicate?

- A. The result is statistically significant.
- B. The result is not statistically significant.**
- C. The study has high statistical power.
- D. The confidence interval is narrow.

P-values gauge how compatible the observed data are with the null hypothesis. If the p-value is greater than 0.05, we do not have enough evidence to reject the null at the conventional 5% significance level, so the result is not statistically significant. This means the data don't show a statistically reliable effect at that threshold, not that there is definitely no effect. It could reflect limited sample size or variability rather than a true absence of an effect. It's also not a direct measure of study power or of how precise the estimate is; a high p-value doesn't guarantee high power or a narrow confidence interval.

4. In evaluating the development of a CPR, which criterion ensures that all relevant predictive factors are included during development?

- A. Did investigators compare results to a criterion test?
- B. Were the predictive factors operationally defined?
- C. Was the sample size large enough to accommodate the number of predictor factors?
- D. Did investigators include all relevant predictive factors in the development process?**

Including all relevant predictive factors during development ensures the rule truly reflects the influences that matter for the outcome. When you develop a CPR, you want the model to account for the key drivers that could change risk; omitting important predictors leaves out pieces of the puzzle, leading to incomplete or biased predictions and poorer performance across settings. By deliberately incorporating all relevant factors identified from evidence, clinical reasoning, and prior research, the rule gains content validity and better generalizability. Other options touch on important aspects of validation and measurement—comparing to a criterion test pertains to external validation, defining predictive factors relates to how they are measured, and having a large enough sample size helps stabilize estimates and prevents overfitting. However, none of these guarantees that every relevant predictive factor is included in the development process the way explicitly including all relevant predictive factors does.

5. Which statement accurately defines inter-rater reliability?

- A. The same rater tests the same patient on different occasions
- B. Two or more different raters test the same patient**
- C. The same test is repeated with the same rater
- D. The patient is tested twice with different index tests

Inter-rater reliability measures how consistently different raters score or classify the same patient. When two or more clinicians independently assess the same person and arrive at similar results, the assessment demonstrates good inter-rater reliability, meaning the rating is not heavily dependent on who does the rating. This matters because it shows the instrument or method yields consistent results across observers. If ratings vary widely between clinicians, reliability is low. For contrast, intra-rater reliability would involve the same clinician rating the same patient on different occasions, emphasizing consistency of one observer. Repeating the same test with the same rater checks test-retest reliability, not agreement between different raters. Testing the patient twice with different index tests relates to concordance between tests, not observer reliability.

6. Prognosis is estimated by?

- A. Response to questions related to disease-specific factors/demographics**
- B. Medical comorbidities
- C. Behavioral comorbidities
- D. Double-blind randomization

Prognosis is estimated from information that directly reflects how the disease tends to progress in an individual, combined with who the person is. Disease-specific factors such as stage or severity and relevant biological markers, along with demographic details like age and overall health, provide the best basis for forecasting outcomes and guiding treatment decisions. Medical and behavioral comorbidities can influence prognosis by altering overall health or complicating management, but they're not the primary determinants used to estimate disease trajectory. Double-blind randomization is a research design feature and does not inform a patient's prognosis. In practice, prognosis comes from integrating the disease's characteristics with the patient's profile, not from trial design elements.

7. Which statement best describes the role of reports on study quality in evidence synthesis?

- A. They determine the sequencing of trials in the database
- B. They influence interpretation of overall findings and confidence in effect estimates**
- C. They replace risk of bias assessment
- D. They are optional if the meta-analysis is significant

Understanding study quality helps us judge how trustworthy the overall findings are. When included studies have flaws in design or conduct, their results may be biased, so the combined estimate from a synthesis can overstate or understate the true effect. Reports that describe study quality tell us how much confidence to place in the pooled result and often guide judgments about the certainty of evidence. They also support sensitivity analyses—checking whether excluding lower-quality studies changes the conclusion. The other points don't fit because study quality reports don't determine the order of trials in a database, they don't replace a formal risk of bias assessment (they're part of that process, not a replacement), and quality considerations are important even if the meta-analysis shows a significant result.

8. In assessing credibility of self-reported measures, which criterion involves documenting how items were developed and tested?

- A. Adequate description of survey development process**
- B. Administration and scoring requirements
- C. Instrument reliability
- D. Is instrument valid

Transparency in how a survey is developed and tested is what lends credibility to self-reported measures. Documenting the development process shows how items were generated, refined, and evaluated before use—often detailing the literature basis, expert reviews, pilot testing, cognitive interviews, and revisions. This kind of description helps readers judge whether the items adequately cover the intended domain and whether potential biases were addressed during creation, which supports trust in the measure's content and construct relevance. Administration and scoring requirements focus on how the survey is given and how responses are transformed into data; reliability concerns whether the measure yields consistent results; validity concerns whether the instrument actually measures what it intends to. While those aspects matter, the requirement to document the development and testing process specifically targets the justification and transparency of the items' creation, which underpins the instrument's credibility from the outset.

9. Which level of measurement is defined by ordering without assuming equal intervals between the levels?

- A. Interval
- B. Nominal
- C. Ratio
- D. Ordinal**

The key idea is that you can rank observations, but you cannot assume equal distances between consecutive levels. This lets you say which is higher or lower, but not by how much one exceeds another. Examples include class ranking or Likert-type scales (strongly agree to strongly disagree). Because you're just ordering, the precise difference between adjacent points isn't defined, so arithmetic on the values isn't meaningful in the way it is for other levels. Nominal data are categories with no inherent order, so they don't fit the idea of ranking. Interval data have ordered levels with equal intervals between them, which means you can discuss differences and averages in a meaningful way, but there's no true zero. Ratio data also have equal intervals and a true zero, introducing even more mathematical possibilities. The description given matches ordinal measurement.

10. If a study uses more than one administration method for a self-reported instrument, what should researchers do?

- A. Use only the mode with the highest reliability
- B. Reexamine measurement properties for each mode**
- C. Ignore mode differences
- D. Rely on the original validation

When a self-report instrument is given through more than one administration method, responses can be affected by the mode itself, so you need to verify that the instrument works the same way across modes. This means reexamining its measurement properties for each mode—looking at reliability, validity, and how well the item responses relate to the underlying construct in each mode. You'd also assess measurement equivalence across modes, such as whether items operate differently for different formats (differential item functioning) and whether a unified factor structure holds (invariance testing). If you find differences, you shouldn't pool data blindly; you might adjust scoring, use mode-specific norms, or analyze modes separately and report the limitations. Simply picking the mode with the highest reliability ignores potential validity issues across modes, ignoring mode differences risks biased conclusions or misinterpretation, and relying on the original validation assumes the same properties across modes, which may not be true.

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://ebp2.examzify.com>

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

SAMPLE