

ETHICS OF ARTIFICIAL INTELLIGENCE (AI) PRACTICE TEST (SAMPLE)

STUDY GUIDE



**SAMPLE STUDY GUIDE
WITH LIMITED QUESTIONS**

**VISIT US ONLINE FOR
THE FULL VERSION STUDY GUIDE**

<https://ethicsofai.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. Which view argues that meaning is determined by use?**
 - A. Structuralist View.
 - B. Associationist View.
 - C. Formalist View.
 - D. Cognitivist View.

- 2. What does a cosine similarity score of 0 indicate about two vectors?**
 - A. No similarity
 - B. Opposite meanings
 - C. Identical meaning
 - D. Moderate similarity

- 3. What are context-dependent embeddings?**
 - A. Dynamic vector representations of words that change based on surrounding context, modern models follow this
 - B. Static vector representations that do not change with context
 - C. A method only used for image data
 - D. A technique for compressing models

- 4. What is the main takeaway from the Chinese Room argument regarding AI understanding of language?**
 - A. A system with fluent outputs may still lack understanding.
 - B. A look-up table cannot produce fluent language.
 - C. A system with perception and motor capacities guarantees understanding.
 - D. Language understanding requires consciousness.

- 5. What does The Associationist View state about word meaning?**
 - A. The dictionary defines the meaning of a word.
 - B. Meaning is determined by grammar alone.
 - C. Meaning is fixed by semantics.
 - D. The meaning of a word is defined by the way it is used (Wittgenstein, 1932).

- 6. Which of the following do word embeddings capture that Markov models don't?**
- A. Short-range word sequences only
 - B. Long-range dependencies, analogies, and rich, contextual meaning
 - C. Only surface form without meaning
 - D. Identical to Markov model output
- 7. What conclusion arises from these discussions about language tasks and understanding?**
- A. Passing language tasks does not prove genuine understanding or mental states.
 - B. Language fluency proves consciousness.
 - C. Symbol manipulation is identical to thinking.
 - D. Minds cannot be studied scientifically.
- 8. Why might embedding models place eggplant and aubergine close in vector space?**
- A. Because there is a universal tag for synonyms.
 - B. Because synonyms always share exact spelling.
 - C. Because words with similar contexts tend to have high similarity in embedding space.
 - D. Because the model memorizes every sentence verbatim.
- 9. A cosine similarity score of -1 indicates what relation between two terms?**
- A. Synonyms
 - B. Unrelated
 - C. Antonyms
 - D. Contradictions
- 10. In NLP, what does a cosine similarity score of 1 signify?**
- A. Extremely similar contexts
 - B. No similarity
 - C. Opposite meanings
 - D. Somewhat related

Answers

SAMPLE

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

SAMPLE

Explanations

SAMPLE

1. Which view argues that meaning is determined by use?

- A. Structuralist View.
- B. Associationist View.**
- C. Formalist View.
- D. Cognitivist View.

Meaning being determined by use means that what a word means comes from how it's used in real-life practice and the responses it tends to evoke, not from some fixed inner essence or formal structure. The Associationist View captures this by linking meaning to the network of associations built up through experience with words in various contexts. When you use a word, you trigger a web of connections—to objects, actions, sensations, social situations, and prior outcomes—so the word's meaning is effectively the sum of those learned associations shaped by usage over time. This explains why meanings can shift with different contexts or communities of use, and why language learning hinges on forming those associations through practice. In contrast, structuralist approaches look to relationships within language systems themselves, formalist views focus on form and rules rather than actual use, and cognitivist perspectives emphasize internal mental representations—none of which centers on usage as the source of meaning in the way associationism does.

2. What does a cosine similarity score of 0 indicate about two vectors?

- A. No similarity
- B. Opposite meanings**
- C. Identical meaning
- D. Moderate similarity

Cosine similarity measures how much two vectors point in the same direction, ignoring how long they are. A score of 1 means they point the same way, exactly aligned. A score of -1 means they point in opposite directions. A score of 0 means the vectors are perpendicular to each other, with no directional alignment between their features. In other words, there's no shared orientation in the space the vectors live in, which is what we mean by no similarity in their directions. Opposite meanings would correspond to a cosine similarity of -1, not 0. Identical meaning would be 1, and moderate similarity would be a positive value between 0 and 1.

3. What are context-dependent embeddings?

- A. Dynamic vector representations of words that change based on surrounding context, modern models follow this**
- B. Static vector representations that do not change with context
- C. A method only used for image data
- D. A technique for compressing models

Context-dependent embeddings are dynamic vector representations of words that vary with surrounding context. Unlike static embeddings, where each word has a single fixed vector, contextualized embeddings are computed for each occurrence by incorporating the other words in the sentence (and sometimes the broader text) through mechanisms like self-attention in transformers. This means the same word can have different representations in different sentences, which helps resolve nuances and multiple meanings of a word, such as "bank" in "river bank" versus "financial bank." Modern models like BERT or GPT produce these contextual embeddings, reflecting how language usage changes with context. The other options don't fit because one describes static embeddings that ignore context, another wrongly limits the concept to image data, and the last describes model compression, which is unrelated to how word representations are formed.

4. What is the main takeaway from the Chinese Room argument regarding AI understanding of language?

- A. A system with fluent outputs may still lack understanding.**
- B. A look-up table cannot produce fluent language.
- C. A system with perception and motor capacities guarantees understanding.
- D. Language understanding requires consciousness.

The main idea is that producing fluent language output can be done without actually understanding the language. In the Chinese Room thought experiment, a person who doesn't know Chinese follows a precise set of rules to transform symbols and generate convincing Chinese replies. To an external observer, the responses read as fluent understanding, but the person inside the room has no grasp of meaning. The takeaway is that syntactic symbol manipulation can yield convincing language without semantic comprehension, so fluent output alone does not prove true understanding. That's why the option stating that a system with fluent outputs may still lack understanding is the best fit. The other ideas misstate the point: a simple lookup mechanism could, in principle, produce fluent responses, so the claim isn't about look-up tables versus fluency; having perception or motor abilities does not by itself guarantee understanding; and while some argue about the role of consciousness, the Chinese Room primarily targets the distinction between appearing to understand and actually understanding, not a definitive claim about consciousness being required.

5. What does The Associationist View state about word meaning?

- A. The dictionary defines the meaning of a word.
- B. Meaning is determined by grammar alone.
- C. Meaning is fixed by semantics.
- D. The meaning of a word is defined by the way it is used (Wittgenstein, 1932).**

Meaning, in this view, comes from how a word is used in practice. The Associationist perspective sees words as tied to learned associations with situations, actions, and responses; the sense of a word unfolds as we encounter it in different contexts and rely on those contexts to guide how we use it again. In other words, what a word means isn't fixed by a dictionary entry or by grammar alone, but by the way it functions in real language—the actions, expectations, and social purposes it enables. Wittgenstein's remark that meaning is the use of a word in the language captures this idea: meaning is rooted in use, not in an abstract definition. This emphasis on use and context helps explain how words can shift meaning across situations while still aligning with the same overall familiar usage.

6. Which of the following do word embeddings capture that Markov models don't?

- A. Short-range word sequences only
- B. Long-range dependencies, analogies, and rich, contextual meaning**
- C. Only surface form without meaning
- D. Identical to Markov model output

Word embeddings are designed to encode words as dense vectors in such a way that semantic relationships and context are reflected in their geometry. This means that words with related meanings or usages end up with similar embeddings, and relationships between words can be captured by simple vector operations. Because of this, embeddings can reflect long-range patterns of meaning and how words relate to each other across many contexts, not just immediate neighbors. Markov models, on the other hand, generate predictions based on a fixed and typically short history of preceding words. They rely on local, short-range dependencies and don't build a space that encodes broader semantic relationships or analogies. That's why they miss the richer, non-local structure that embeddings capture. So the option describing long-range dependencies, analogies, and rich contextual meaning best captures what word embeddings provide beyond what Markov models offer.

7. What conclusion arises from these discussions about language tasks and understanding?

- A. Passing language tasks does not prove genuine understanding or mental states.
- B. Language fluency proves consciousness.
- C. Symbol manipulation is identical to thinking.
- D. Minds cannot be studied scientifically.**

Performing well on language tasks does not prove genuine understanding or mental states. The discussions draw a distinction between outward performance and inner experience: a system can manipulate symbols and generate fluent, plausible responses without actually grasping meaning or having conscious experiences. This is the intuition behind separating syntax from semantics—the rules-driven way a program handles symbols versus what it feels like to understand. That's why language fluency alone isn't evidence of consciousness; fluent output can be produced by algorithms without any subjective experience. It's also why symbol manipulation is not the same as thinking—thinking involves intentionality, beliefs, and awareness that mere symbol processing typically lacks. Finally, these discussions do not deny that minds can be studied scientifically; cognitive science and neuroscience investigate mental processes empirically.

8. Why might embedding models place eggplant and aubergine close in vector space?

- A. Because there is a universal tag for synonyms.
- B. Because synonyms always share exact spelling.
- C. Because words with similar contexts tend to have high similarity in embedding space.**
- D. Because the model memorizes every sentence verbatim.

Word embeddings rely on the idea that words used in similar contexts have similar meanings. When two words share many similar contexts, their vectors in the embedding space become close to each other. Eggplant and aubergine are just two names for the same thing, so they tend to appear in the same kinds of sentences—recipes, grocery lists, vegetable sections, cooking terms—which makes their contextual patterns align. That alignment pulls their vectors together, reflecting their semantic similarity. The other options don't capture how embeddings work: there isn't a universal tagging system for synonyms in these models, synonyms don't require identical spelling, and embeddings don't memorize every sentence verbatim; they generalize from patterns in the data.

9. A cosine similarity score of -1 indicates what relation between two terms?

- A. Synonyms
- B. Unrelated
- C. Antonyms**
- D. Contradictions

Cosine similarity measures how closely two word vectors point in the same direction. A score of -1 means the vectors point exactly opposite to each other, signaling opposite meanings. In word embeddings, synonyms share similar meanings and therefore align in direction, giving high positive scores; unrelated terms tend to have directions that don't align, giving scores near zero. Antonyms, being opposite in meaning, often fall in opposite directions, producing negative similarity, and in the extreme case, -1. So the best answer is antonyms.

10. In NLP, what does a cosine similarity score of 1 signify?

- A. Extremely similar contexts**
- B. No similarity
- C. Opposite meanings
- D. Somewhat related

Cosine similarity in NLP measures how close the directions of two embedding vectors are, ignoring their magnitudes. A score of 1 means the vectors point in exactly the same direction, which means the two representations capture the same meaning or occur in essentially the same contexts. In other words, they are maximally aligned semantically. That's why the phrase describing extreme similarity of contexts is the best fit. Scores near 0 indicate no linear relationship, while scores near -1 would imply opposite meanings; values between 0 and 1 reflect varying degrees of relatedness, with higher values showing closer similarity.

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://ethicsofai.examzify.com>

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

SAMPLE