

AAISM DOMAIN 1: AI GOVERNANCE, PROGRAM MANAGEMENT PRACTICE TEST (SAMPLE)

STUDY GUIDE



SAMPLE STUDY GUIDE WITH LIMITED QUESTIONS

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

<https://aaismdomain1.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	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. Which stakeholder group is primarily responsible for ensuring data used in AI meets governance and quality standards?**
 - A. Data engineers
 - B. Data stewards/owners
 - C. AI ethicists
 - D. Privacy experts
- 2. Which term describes a data repository that can aggregate structured and unstructured data while preserving access controls from the source?**
 - A. Data Exploration and Training Platform
 - B. Vector Database
 - C. Data Lake
 - D. AI System Production
- 3. What term describes deliberately feeding incorrect data to an AI to generate incorrect results?**
 - A. Prompt injection
 - B. Adversarial Inference
 - C. Model drift
 - D. Data poisoning
- 4. What metrics demonstrate AI program ROI beyond financial returns?**
 - A. Only revenue growth.
 - B. Productivity gains, decision speed, risk reduction, customer impact, governance maturity improvements, and regulatory compliance.
 - C. Only cost savings.
 - D. Only market share.
- 5. Which mechanisms involve data validation, cleaning, and anomaly detection to prevent data poisoning?**
 - A. Data quality in AI security
 - B. Explainability in AI
 - C. Human-in-the-loop in AI
 - D. Data validation/cleaning and anomaly detection mechanisms

- 6. Which data quality attribute means data is free from errors and representative of real-world situations?**
- A. Timeliness
 - B. Accuracy
 - C. Validity
 - D. Completeness
- 7. What is the purpose of business continuity planning for AI systems?**
- A. Improve UI aesthetics.
 - B. Ensure resilience and recovery after AI system disruptions.
 - C. Increase training data.
 - D. Reduce model complexity.
- 8. What role would likely be responsible for interpreting anomalous AI outputs during model development?**
- A. Data engineers and scientists
 - B. AI ethicists
 - C. Privacy experts
 - D. Data stewards/owners
- 9. Which standard is described as a voluntary framework focusing on ethical use, transparency, and accountability in AI governance?**
- A. ISO/IEC 42001
 - B. ISO/IEC 23053
 - C. OECD AI Principles
 - D. NIST RMF
- 10. Explainability in AI governance is used to assess what aspect of model decisions?**
- A. Explainability scores indicate how transparent a model's decisions are and support audits.
 - B. Explainability scores measure hardware efficiency.
 - C. Explainability scores are optional and rarely used.
 - D. Explainability scores replace data privacy.

Answers

SAMPLE

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

SAMPLE

Explanations

SAMPLE

1. Which stakeholder group is primarily responsible for ensuring data used in AI meets governance and quality standards?

- A. Data engineers
- B. Data stewards/owners**
- C. AI ethicists
- D. Privacy experts

Data stewards/owners are accountable for data governance and data quality across the organization. They define and enforce the standards that data must meet—such as accuracy, completeness, consistency, timeliness, and proper metadata—ensuring the data used by AI is reliable and well-managed. This role also covers data definitions, lineage, cataloging, access policies, and ongoing quality monitoring, which together establish trust in AI outputs. Data engineers, while they build and maintain the data pipelines that supply data, typically execute and operationalize those standards rather than own them. AI ethicists focus on ethical considerations in AI applications, such as fairness and transparency, not the formal governance of data quality. Privacy experts concentrate on protecting personal data and compliance with privacy laws, which is essential but addresses a specific aspect of governance rather than the overall data quality governance for AI.

2. Which term describes a data repository that can aggregate structured and unstructured data while preserving access controls from the source?

- A. Data Exploration and Training Platform
- B. Vector Database
- C. Data Lake**
- D. AI System Production

A data lake is the repository that can hold both structured and unstructured data at scale while preserving governance and access controls from the source. It stores data in its native format, ranging from tables and CSVs to PDFs, images, and logs, enabling flexible schema-on-read analytics and data science workflows. Because data lakes are designed to integrate with security and governance mechanisms—like identity and access management, encryption, and policy-based controls—they can maintain or enforce access rules as data is ingested, stored, and consumed, helping preserve the origin's security posture. The other options don't fit as well. A data exploration and training platform focuses on tools and interfaces for analyzing data and building models rather than serving as a central, diverse data repository with integrated access controls. A vector database specializes in storing high-dimensional vectors for similarity search and is not typically used to aggregate all data types with preserved source-level access controls. An AI system production refers to deploying models and serving AI capabilities, not to storing and governing large, heterogeneous datasets.

3. What term describes deliberately feeding incorrect data to an AI to generate incorrect results?

- A. Prompt injection
- B. Adversarial Inference
- C. Model drift
- D. Data poisoning**

Deliberately feeding incorrect data to an AI to produce incorrect results is data poisoning. The idea is to taint the information the model learns from, so its outputs become unreliable, biased, or aligned with the attacker's goals. This typically happens during training or data collection, when manipulated data shifts the model's understanding or creates vulnerabilities that show up later in deployment. Since the attack targets the model's learned parameters and knowledge, the effects can persist and be hard to undo without remediation like data cleansing or retraining. Prompt injection, in contrast, aims to influence the model's behavior at inference time through the prompt itself, not by altering the training data. Model drift refers to natural changes in data patterns over time that affect performance, not deliberate manipulation. Adversarial inference isn't the standard term for this kind of data manipulation and is not the right descriptor for feeding bad data to corrupt training.

4. What metrics demonstrate AI program ROI beyond financial returns?

- A. Only revenue growth.
- B. Productivity gains, decision speed, risk reduction, customer impact, governance maturity improvements, and regulatory compliance.**
- C. Only cost savings.
- D. Only market share.

ROI from an AI program comes from a mix of outcomes, not just dollars. The most complete way to show value beyond financial returns is to track a range of metrics that reflect how AI boosts operations, speed, risk management, and stakeholder value. Productivity gains measure how much more work gets done or how much human effort is saved. Decision speed tracks how quickly insights translate into action. Risk reduction captures fewer incidents, lower error rates, and stronger safeguards. Customer impact looks at how AI improves customer experience, satisfaction, retention, or loyalty. Governance maturity and regulatory compliance gauge how well the organization is standardizing processes, improving model governance, audit readiness, and staying aligned with regulations. Together, these metrics paint a fuller picture of ROI by showing tangible improvements in efficiency, capability, and risk management, not just revenue or cost figures. Relying on only revenue growth misses efficiency and risk benefits; focusing only on cost savings omits speed, decision quality, and governance gains; and measuring only market share overlooks internal improvements and compliance.

5. Which mechanisms involve data validation, cleaning, and anomaly detection to prevent data poisoning?

- A. Data quality in AI security
- B. Explainability in AI
- C. Human-in-the-loop in AI

D. Data validation/cleaning and anomaly detection mechanisms

Data validation, cleaning, and anomaly detection are all about protecting the data that trains and updates an AI system. In data poisoning, an attacker tries to insert manipulated data to distort the model's behavior. Validation checks ensure data records conform to expected formats, types, ranges, and cross-feature consistency, so bad entries don't slip through. Cleaning goes further by removing or correcting noisy or obviously faulty data, reducing the chance of poisoned examples influencing learning. Anomaly detection looks for data points that don't fit the normal patterns or distributions, flagging suspicious samples for review or rejection before they affect the model. Used together, these mechanisms strengthen the data pipeline, making it harder for poisoned data to contaminate training. Other options focus on different aspects, like explaining how models decide their outputs or adding human oversight, rather than the data-layer defenses described here.

6. Which data quality attribute means data is free from errors and representative of real-world situations?

- A. Timeliness
- B. Accuracy**
- C. Validity
- D. Completeness

Accuracy is the data quality attribute that means data values are correct and true to reality. It ensures data is free from errors and faithfully reflects real-world measurements or facts. For example, a recorded sales amount should match the actual transaction, and a customer's address should correspond to where they live. When data is accurate, analyses and decisions based on it align with what's really happening. Timeliness discusses how current the data is, completeness covers whether all required data is present, and validity checks that data follows defined formats or rules.

7. What is the purpose of business continuity planning for AI systems?

- A. Improve UI aesthetics.
- B. Ensure resilience and recovery after AI system disruptions.**
- C. Increase training data.
- D. Reduce model complexity.

Business continuity planning for AI systems is about resilience and rapid recovery when disruptions occur. It involves identifying which AI services are critical, setting recovery targets, and building redundant infrastructure, backup data pipelines, and clear incident response and rollback procedures. In the AI context, disruptions can come from data outages, failures in model hosting, broken data streams, model drift, software bugs, or cyber incidents. A solid plan provides standby models, failover paths, version control, and tested recovery runbooks to minimize downtime and preserve the integrity of decisions. The other options don't address keeping AI-enabled operations running or restoring them after interruptions—UI aesthetics, more training data, or simply reducing model complexity don't inherently ensure continuity or quick recovery.

8. What role would likely be responsible for interpreting anomalous AI outputs during model development?

- A. Data engineers and scientists**
- B. AI ethicists
- C. Privacy experts
- D. Data stewards/owners

Interpreting anomalous AI outputs during model development is a technical debugging task centered on understanding both the data pipeline and the model itself. Data engineers and data scientists have the hands-on role of inspecting data quality, feature representations, model configurations, and evaluation results. They trace unexpected outputs to sources such as data preprocessing issues, misaligned features, or model parameter choices, then adjust the data flow, features, or the model to fix the problem. This focused, technical interpretation and remediation during development is precisely what they do. AI ethicists are more concerned with fairness, accountability, and the societal impact of AI behavior, while privacy experts concentrate on data protection and compliance. Data stewards/owners manage data governance and stewardship responsibilities. While important, these roles don't typically perform the real-time technical interpretation and debugging of anomalous outputs that occurs in model development.

9. Which standard is described as a voluntary framework focusing on ethical use, transparency, and accountability in AI governance?

- A. ISO/IEC 42001**
- B. ISO/IEC 23053
- C. OECD AI Principles
- D. NIST RMF

The question is testing recognition of a formal standard that provides a voluntary governance framework for AI focused on ethics, transparency, and accountability. The standard that fits this description is a management-system standard specifically for AI governance. It offers a structured, auditable approach organizations can adopt to ensure responsible AI use, with clear policies, governance roles, and controls that promote ethical practices, transparency in how AI decisions are made, and accountability for outcomes. Being voluntary, it's intended to be adopted and possibly certified against, rather than mandated. Why the other options aren't the best fit: guidelines like OECD AI Principles set high-level ethical expectations but are not a formal standard with a governance management system and certification path. The NIST RMF is a risk-management framework for information systems in general, not tailored to AI governance. ISO/IEC 23053 provides an overview of AI standardization rather than a governance framework.

10. Explainability in AI governance is used to assess what aspect of model decisions?

- A. Explainability scores indicate how transparent a model's decisions are and support audits.
- B. Explainability scores measure hardware efficiency.
- C. Explainability scores are optional and rarely used.
- D. Explainability scores replace data privacy.

Explainability in AI governance centers on how transparent a model's decisions are and how easily humans can understand and trace the reasoning behind them. This is why the best answer states that explainability scores indicate how transparent decisions are and support audits. In governance, stakeholders—from developers to regulators—need to know why a model produced a given output, what factors influenced it, and whether the reasoning is fair or biased. Explainability scores quantify this transparency, making it easier to verify, hold the model accountable, troubleshoot issues, and ensure compliance during audits and reviews. They help reveal which inputs or features most influenced a decision and whether there are hidden biases or risky behaviors. The other options don't fit the governance goal. Hardware efficiency is a separate concern from how decisions are explained. Treating explainability scores as optional and rarely used conflicts with governance practices that often mandate transparency. Thinking explainability scores replace data privacy confuses the purpose—explanations should enhance understanding while respecting privacy, not substitute for it. So, explainability is fundamentally about making model decisions transparent and auditable.

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://aaismdomain1.examzify.com>

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

SAMPLE