



Certified AI Assurance Professional (CAIAP)

Ten Sample Questions

These ten questions are written at the difficulty of the live exam. Six carry a figure, a data table, a configuration or a code listing; four are plain scenario questions, as on the live paper. Every one is a scenario a working AI assurance professional would meet, and none of them appears in the live question bank.

The live exam is multistage adaptive: how you score on one stage selects the difficulty of the next, so no two candidates sit the same paper. It is a performance-based exam: each form can carry up to 5 to 12 performance-based questions alongside the multiple-choice questions, which means you build, sequence or complete something rather than only choosing from a list. Answers and full explanations follow each question here; the live exam does not disclose item-level answers.

Question 1

Domain 1: AI Governance Foundations

DATA

Emma Wilson matches each governance role to its responsibility.

Role	Responsibility
AI governance board	(to choose)
Model risk owner	(to choose)
Assurance and audit function	(to choose)
Data protection lead	(to choose)

Table 1.1 Roles and their responsibilities.

Which set of responsibilities is correct?

- ☐ A Sets policy and risk appetite, owns a single model's risk, gives independent assurance, and advises on privacy, but with ownership and oversight swapped so no one is truly independent.
- ☐ B Sets policy and risk appetite, owns a single model and its controls, provides independent assurance, and advises on data protection, matching each role to its actual responsibility.
- ☐ C Owns a single model, sets the policy and risk appetite, advises on data protection, and provides assurance, so the board audits itself and independence is lost across the board.
- ☐ D Provides independent assurance, sets policy, owns every model at once, and advises on privacy, so the audit function both sets policy and signs off on its own decisions.

Correct answer: B

The board sets policy and risk appetite, the model risk owner owns a specific model and its controls, the assurance and audit function provides independent oversight, and the data protection lead advises on privacy. Only option B keeps these responsibilities separate so no one reviews their own work, which the others break by mixing ownership and oversight.

Question 2

Domain 1: AI Governance Foundations

CONFIG

James Carter reviews an organization that is deploying models with no governance in place.

```
ai system inventory : none
accountable owner   : unassigned
risk classification : not done
policy sign-off     : skipped
```

Listing 2.1 The current state.

What is the soundest first step?

- ☐ A Deploy every model faster to build momentum, on the assumption that a governance program can always be added later once systems are already in production use.
- ☐ B Ban all AI use across the organization outright, since removing every system is the only workable way to manage the risk that an ungoverned model could pose.
- ☐ C Build an AI system inventory, assign an accountable owner to each, classify risk, and require policy sign-off, so governance rests on a known population of systems.
- ☐ D Buy a single monitoring tool and treat that purchase alone as the whole program, since tooling by itself will surface and manage every governance gap on its own.

Correct answer: C

Governance begins with knowing what you have, so an inventory with an accountable owner, a risk classification, and policy sign-off gives the program a known population of systems to manage. Deploying faster defers the problem, an outright ban is not workable, and a single tool does not by itself establish accountability.

Question 3

Domain 2: AI Risk Assessment and Controls

DATA

Olivia Davis maps each AI risk to the control that addresses it.

Risk	Control
Model output can be biased	(to choose)
Training data may leak	(to choose)
Third-party model is opaque	(to choose)
No record of decisions	(to choose)

Table 3.1 Risks and controls.

Which set of controls is correct?

- A** Bias testing across groups, data minimization and access controls, vendor due diligence with documentation, and decision logging, pairing each risk with its control.
- B** Decision logging, bias testing, vendor due diligence and data minimization, ordered so that logging is what prevents the training data from leaking in the first place.
- C** Vendor due diligence, decision logging, data minimization and bias testing, arranged so that due diligence is the measure that removes bias from the model output.
- D** Data minimization, vendor due diligence, decision logging and bias testing, arranged so that minimization is the control that makes an opaque vendor model transparent.

Correct answer: A

Biased output is checked by bias testing across groups, data leakage is reduced by minimization and access controls, an opaque vendor model calls for due diligence and documentation, and a missing record is fixed by decision logging. Only option A pairs each risk with the control that directly addresses it; the others map controls to risks they do not solve.

Question 4

Domain 2: AI Risk Assessment and Controls

CONFIG

Liam Anderson lists the stages of a risk assessment, shown out of order.

- A. Assign and track treatment for each risk
- B. Identify the AI system and its context of use
- C. Sign off residual risk against risk appetite
- D. Analyze likelihood and impact of each risk
- E. Identify what could go wrong and who is affected

Listing 4.1 Stages to order.

What is the correct order of these stages?

- A** D, then B, then E, then C, then A, analyzing risks before the system or its harms are identified and signing off before any treatment has been assigned.
- B** B, then D, then A, then E, then C, analyzing impact before the harms are even identified and signing off residual risk before the harms are known at all.
- C** E, then C, then B, then D, then A, signing off residual risk near the start before the system context and its likelihood and impact have been established.
- D** B, then E, then D, then A, then C, scoping the system, identifying harms, analyzing likelihood and impact, assigning treatment, then signing off residual risk.

Correct answer: D

A risk assessment first scopes the system and its use (B), identifies the harms and who is affected (E), analyzes likelihood and impact (D), assigns and tracks treatment (A), and finally signs off the residual risk against the risk appetite (C). The other orders analyze or sign off risks before the system and its harms have even been identified.

Question 5

Domain 4: Regulation, Standards and Frameworks

DATA

Ava Thompson matches each EU AI Act risk tier to how the framework treats it.

Risk tier	Treatment
Unacceptable-risk use	(to choose)
High-risk use	(to choose)
Limited-risk use	(to choose)
Minimal-risk use	(to choose)

Table 5.1 Tiers and their treatment.

Which set of treatments is correct?

- ☐ A Prohibited outright, light transparency only, full conformity obligations, and no obligations, with the high-risk and limited-risk duties swapped between the two tiers.
- ☐ B Prohibited outright, subject to strict obligations, subject to transparency duties, and largely unrestricted, matching each risk tier to how the framework treats it.
- ☐ C Allowed with notice, prohibited, unrestricted, and strict obligations, arranged so an unacceptable-risk use is merely flagged with a transparency notice to users.
- ☐ D Strict obligations, prohibited outright, unrestricted, and transparency only, arranged so the highest tier faces controls while a banned use is simply prohibited.

Correct answer: B

The EU AI Act prohibits unacceptable-risk uses, places strict obligations on high-risk uses, applies transparency duties to limited-risk uses, and leaves minimal-risk uses largely unrestricted. Only option B keeps each tier with its correct treatment; the others swap the duties or treat a banned use as merely something to disclose.

Question 6

Domain 6: AI Lifecycle Assurance and Auditing

DATA

An AI system is assured as shown below.

Observation	Result
Independent review before release	Never
Controls tested, not just described	No
Monitored for drift in production	No
Re-assessed after a model change	No

Table 6.1 How assurance is done today.

What is the core weakness, and what should be fixed?

- ☐ A Nothing is wrong, because a model that passed review once needs no further monitoring, testing, or re-assessment for the rest of its time in production use.
- ☐ B The model is too small, so a larger model would pass assurance on its own and remove any need to test controls or monitor the system after it is released.
- ☐ C Assurance exists only on paper, so test controls rather than describe them, review independently before release, and monitor for drift and re-assess after each change.

- D** The policy is too short, so lengthening the written policy alone would make the controls effective without any independent testing or ongoing monitoring at all.

Correct answer: C

Controls that are described but never tested, with no independent review and no monitoring, mean assurance exists only on paper. The fix is to test controls for real, review independently before release, and monitor for drift with a re-assessment after each change. A bigger model or a longer policy does not make untested controls effective.

Question 7 Domain 3: Responsible and Trustworthy AI

A model denies loan applications, and applicants are given no way to understand or challenge the decision. What is the soundest response?

- A** Provide a meaningful explanation of the main factors behind each decision and a route to human review, so affected people can understand and contest the outcome.
- B** Give no explanation at all to any applicant, since revealing how the model works would always let people game the system and undermine every future decision.
- C** Replace the model with a coin flip for fairness, on the assumption that a purely random decision cannot be biased against any individual or group of applicants.
- D** Tell applicants only that the computer decided and cannot be questioned, since an automated decision is by definition final and needs no human review at all.

Correct answer: A

Responsible use of a high-impact automated decision requires a meaningful explanation of the main factors and a route to human review so affected people can understand and contest it. Withholding all explanation removes accountability, a coin flip is not a decision, and declaring the outcome final denies people any recourse.

Question 8 Domain 4: Regulation, Standards and Frameworks

A team asks how the NIST AI Risk Management Framework is meant to be used. Which description is soundest?

- A** It is a binding law with fines, so following it is mandatory in every country and replaces the need for any internal AI governance policy of your own.
- B** It is a certification you buy once, after which no further risk work, mapping, measurement, or management of AI systems is ever needed again for compliance.
- C** It is a fixed checklist of banned uses, so meeting it means simply avoiding a short list of prohibited applications and nothing more beyond that step.
- D** It is a voluntary framework whose functions help organizations govern, map, measure, and manage AI risk, applied continuously across the system lifecycle.

Correct answer: D

The NIST AI RMF is a voluntary framework whose core functions, govern, map, measure, and manage, help organizations address AI risk continuously across the lifecycle. It is not a binding law, a one-time certification, or a fixed checklist of banned uses, so only option D describes how it is meant to be applied.

Question 9

Domain 5: Data Governance and Privacy for AI

Personal data was collected and used to train a model with no lawful basis established. What does privacy law require?

- ☐ A Nothing is required, because any data that can be found on the public internet is automatically free to use for training a model without further consideration.
- ☐ B A lawful basis, purpose limitation, and data minimization are required, and individuals need transparency and a way to exercise their rights over the data.
- ☐ C Only a longer privacy policy is required, since disclosing the practice in dense legal text alone satisfies every privacy obligation for training data use.
- ☐ D Only encryption at rest is required, since protecting the stored data from outside access resolves the lawful-basis and purpose-limitation questions on its own.

Correct answer: B

Using personal data to train a model still requires a lawful basis, purpose limitation, and data minimization, along with transparency and a way for individuals to exercise their rights. Public availability does not remove these duties, and a longer policy or encryption alone does not supply a lawful basis or limit the purpose of use.

Question 10

Domain 6: AI Lifecycle Assurance and Auditing

A team keeps no records of training data, model versions, or key decisions, so an audit cannot verify what was done. What practice best fixes this?

- ☐ A Keep no records at all, because a clean absence of documentation means an auditor has nothing to criticize and the system will therefore pass every review.
- ☐ B Recreate records from memory only when an auditor asks, since reconstructing the history after the fact is just as reliable as recording it as work happens.
- ☐ C Maintain traceable records of data sources, model versions, and key decisions across the lifecycle, so assurance and audit can verify what was done and why it was done.
- ☐ D Store only the final model file and nothing else, on the assumption that the shipped artifact alone proves how every earlier decision was made and approved.

Correct answer: C

Assurance and audit depend on traceable records of data sources, model versions, and key decisions kept throughout the lifecycle, so reviewers can verify what was done and why. Keeping nothing leaves the work unverifiable, reconstructing from memory is unreliable, and the final model file alone does not show the decisions behind it.