

Globally Recognised Certification Body
190+ countries · 2,50,000+ Certified Professionals

AI GOVERNANCE PROFESSIONAL INTERVIEW PREPARATION GUIDE



Basic AI Governance Questions

Interviewers will often open with foundational questions to gauge your conceptual understanding of AI governance before moving into more technical territory. These questions test your ability to define key terms, distinguish between frameworks, and articulate why governance matters in an organisational context. Being crisp and confident here sets a strong tone for the rest of the interview.

Q1: What is AI governance and why does it matter?

Sample Answer: AI governance is the set of policies, processes, roles, and accountability structures that guide how an organisation develops, deploys, and monitors AI systems. It matters because AI introduces unique risks — from biased outputs to opaque decision-making — that traditional IT governance frameworks were not designed to address. Effective AI governance ensures systems are trustworthy, compliant, and aligned with organisational values.

Q2: How would you distinguish AI governance from general IT governance?

Sample Answer: IT governance focuses on the availability, security, and performance of technology systems. AI governance goes further by addressing model transparency, algorithmic fairness, explainability, and ethical use. AI systems make probabilistic decisions that can change over time due to model drift, requiring continuous monitoring that traditional IT controls do not typically anticipate.

Q3: What are the core pillars of a responsible AI programme?

Sample Answer: Most responsible AI programmes are built on five pillars: fairness and non-discrimination, transparency and explainability, accountability and human oversight, privacy and data protection, and robustness and safety. These pillars are reflected across major frameworks including the EU AI Act, NIST AI RMF, and ISO 42001.

Q4: What is the difference between AI risk management and AI compliance?

Sample Answer: Compliance is about meeting minimum regulatory or contractual obligations. Risk management is broader — it involves identifying, assessing, and mitigating risks even where no specific rule exists. In AI governance, you need both: compliance ensures legality, while risk management ensures responsible and sustainable deployment.

ISO 42001 Interview Questions

ISO 42001 is the world's first international standard for AI management systems. Candidates for senior AI governance, compliance, or GRC roles are increasingly expected to demonstrate familiarity with its structure, requirements, and practical application. The questions below reflect what experienced interviewers ask when assessing your ISO 42001 knowledge.

Q5: What is ISO 42001 and what does it govern?

Sample Answer: ISO 42001 is an international standard that specifies requirements for establishing, implementing, maintaining, and continually improving an Artificial Intelligence Management System (AIMS). It applies to any organisation that develops, provides, or uses AI products and services, and it aligns closely with other ISO management system standards like ISO 27001 and ISO 9001 — meaning organisations with existing management systems can leverage a familiar structure.

Q6: How does ISO 42001 address AI risk?

Sample Answer: ISO 42001 requires organisations to conduct a context-of-use analysis to identify AI-specific risks and opportunities. This includes assessing the potential impact of AI systems on individuals, society, and the environment. It mandates a risk treatment plan, documented controls, and regular review — similar in structure to ISO 27001's risk management approach, but tailored to AI-specific concerns such as bias, explainability, and data quality.

Q7: What is an AIMS and what are its key components?

Sample Answer: An Artificial Intelligence Management System (AIMS) is the structured framework an organisation implements to manage its AI activities in accordance with ISO 42001. Key components include leadership commitment and AI policy, a documented scope of the management system, risk and impact assessments, controls drawn from Annex A, defined roles and responsibilities, and a process for continual improvement including internal audits and management review.

Q8: What role does Annex A play in ISO 42001?

Sample Answer: Annex A provides a reference set of AI-specific controls that organisations can select based on their risk assessment. These controls cover areas including AI policy, data governance, human oversight, system performance monitoring, and stakeholder engagement. Organisations document their control selections and any exclusions in a Statement of Applicability — a concept borrowed directly from ISO 27001.

NIST AI RMF Interview Questions

The NIST AI Risk Management Framework (AI RMF 1.0) has become a cornerstone reference for AI governance practitioners in the United States and increasingly globally. Interviewers at regulated firms, government contractors, and large enterprises will probe your familiarity with its four core functions and how they translate into operational practice.

Q9: Describe the four core functions of the NIST AI RMF.

Sample Answer: **GOVERN:** Establishes accountability, culture, and organisational commitment to AI risk management. **MAP:** Identifies and categorises AI risks in context. **MEASURE:** Analyses and assesses those risks using quantitative and qualitative methods. **MANAGE:** Prioritises and treats risks, with ongoing monitoring and response.

Q10: How does the GOVERN function differ from the other three?

Sample Answer: GOVERN is the enabling function — it creates the organisational conditions under which MAP, MEASURE, and MANAGE can operate effectively. Without GOVERN, the other three functions lack the authority, resources, and accountability structures needed to produce meaningful outcomes. GOVERN encompasses leadership tone, AI policy, workforce accountability, and the organisational culture around responsible AI.

Q11: What are AI RMF Profiles and how are they used?

Sample Answer: Profiles are customisations of the AI RMF tailored to a specific organisational context, sector, or use case. They allow organisations to align the framework's outcomes to their particular risk tolerance, regulatory environment, and AI portfolio. NIST has published sector-specific profiles — for example, for generative AI — that provide more targeted guidance.

Q12: How would you apply the MAP function to a new AI system?

Sample Answer: I would begin by characterising the AI system's intended use, the affected stakeholders, and the deployment environment. I would then identify potential risks — including performance failures, misuse, and discriminatory outputs — and classify them by likelihood and severity. This forms the risk inventory that feeds into the MEASURE function.

NIST AI RMF Interview Questions (continued)

These questions continue the NIST AI RMF series, focusing on measurement, cross-framework integration, implementation challenges, and how the framework applies to emerging AI technologies such as generative AI.

Q13: What metrics might you use in the MEASURE function?

Sample Answer: Quantitative metrics could include model accuracy, false positive/negative rates disaggregated by demographic group, data drift indicators, and system uptime. Qualitative assessments might include expert red-teaming results, stakeholder feedback, and third-party audits. The goal is a multi-dimensional view of trustworthiness.

Q14: How does NIST AI RMF relate to ISO 42001?

Sample Answer: Both frameworks address AI risk but from different perspectives. NIST AI RMF is a voluntary, outcomes-oriented framework focused on risk management processes. ISO 42001 is a certifiable management system standard with auditable requirements. They are complementary: an organisation can use NIST AI RMF as the operational methodology within an ISO 42001 AIMS structure.

Q15: How would you use the MANAGE function to respond to an AI incident?

Sample Answer: The MANAGE function provides the structure for prioritising and treating identified risks, including incident response. Upon an AI incident, I would activate the pre-defined response plan, contain the impact by suspending or limiting the system, conduct a root cause analysis, and implement corrective actions. The incident would be logged in the risk register, and lessons learned would feed back into the MAP and MEASURE functions to prevent recurrence.

Q16: How does the NIST AI RMF apply to generative AI systems?

Sample Answer: NIST has published a Generative AI Profile (NIST AI 600-1) that extends the AI RMF to address risks specific to large language models and other generative systems — including hallucination, data provenance, intellectual property concerns, and misuse for disinformation. The profile maps these risks to the four core functions and provides targeted suggested actions, making it an essential reference for organisations deploying generative AI at scale.

AI Risk & Compliance Questions

Risk and compliance questions are central to virtually every AI governance interview. Hiring managers want to know that you can translate abstract risk concepts into practical controls, policies, and audit-ready documentation. These questions assess your ability to operate in a regulated environment, work with legal and technical teams, and maintain a defensible compliance posture.

Q17: How do you conduct an AI risk assessment?

Sample Answer: An AI risk assessment follows a structured process: identify the system's purpose and use context, catalogue potential risks (technical, operational, ethical, legal), assess likelihood and impact, identify existing controls and gaps, and produce a risk register with treatment recommendations. I would align this process to a recognised framework such as NIST AI RMF or ISO 42001 and involve cross-functional stakeholders including legal, data science, and business owners.

Q18: What is the EU AI Act and how does it classify AI systems?

Sample Answer: The EU AI Act is the world's first comprehensive legal framework for AI. It classifies AI systems into four risk tiers: unacceptable risk (prohibited), high risk (subject to strict conformity requirements), limited risk (transparency obligations), and minimal risk (no specific obligations). High-risk systems — such as those used in recruitment, credit scoring, or law enforcement — must meet requirements around data governance, transparency, human oversight, and accuracy before deployment.

Q19: How would you design an AI model inventory?

Sample Answer: An AI model inventory is a centralised register of all AI systems in use across an organisation. Each entry should capture the system's purpose, owner, data inputs, model type, deployment environment, risk classification, applicable regulations, approval status, and monitoring schedule. The inventory supports audit readiness, incident response, and regulatory reporting, and should be reviewed at least quarterly or upon material changes to any system.

Q20: What is model drift and why is it a compliance concern?

Sample Answer: Model drift occurs when a model's performance degrades over time as the real-world data it encounters diverges from its training data. From a compliance perspective, drift can cause a previously validated and approved model to produce unreliable or discriminatory outputs — potentially in violation of fairness obligations or regulatory requirements. Continuous monitoring and pre-defined revalidation triggers are essential controls.

AI Risk & Compliance Questions (continued)

These questions continue the AI risk and compliance series, covering vendor due diligence, incident management, AI auditing, and the governance of high-risk AI systems under emerging regulatory frameworks.

Q21: How do you approach third-party AI vendor due diligence?

Sample Answer: Third-party AI vendor due diligence should mirror and extend standard vendor risk management. I would assess the vendor's AI governance programme, model transparency documentation, data processing practices, security controls, bias testing results, and contractual commitments around auditability and incident notification. For high-risk AI use cases, I would require model cards, system cards, or equivalent technical documentation.

Q22: What is an AI incident and how should it be managed?

Sample Answer: An AI incident is any event where an AI system produces outputs or behaves in a way that causes or risks causing harm — whether through bias, error, misuse, or unexpected failure. Incident management requires a pre-defined response plan covering detection, containment, investigation, remediation, and stakeholder notification. For regulated sectors, incidents may also trigger mandatory reporting obligations to supervisory authorities.

Q23: How would you structure an AI audit programme?

Sample Answer: An AI audit programme should cover three layers: technical audits of model performance and data quality, process audits of governance controls and documentation, and compliance audits against applicable regulations and standards. Audits should be risk-based — higher-risk systems audited more frequently — and should include both internal reviews and periodic independent assessments. Findings should feed into a continuous improvement cycle.

Q24: What is a conformity assessment under the EU AI Act and who conducts it?

Sample Answer: A conformity assessment is the process by which a high-risk AI system is evaluated against the requirements of the EU AI Act before it can be placed on the market or put into service. For most high-risk systems, providers can conduct a self-assessment. For certain categories — such as biometric identification systems — third-party notified body assessment is required. The assessment must be documented and a CE marking applied upon successful completion.

Scenario-Based Interview Questions

Scenario Q25: A business unit wants to deploy a generative AI tool that was not reviewed by IT or legal. What do you do?

Sample Answer: First, I would engage the business unit to understand the intended use case, the data being processed, and the vendor involved — without being adversarial. I would then conduct a rapid risk triage to determine whether the tool handles personal data, generates customer-facing content, or influences material decisions. Based on the risk profile, I would either fast-track a formal review, implement interim controls, or escalate to leadership if deployment needs to pause. The incident would also prompt a review of the AI intake and approval process to close the gap that allowed this to happen.

Scenario Q26: An AI system used in recruitment is flagged for potential bias against a protected group. How do you respond?

Sample Answer: This is a high-priority incident requiring immediate action. I would convene a cross-functional team including HR, legal, the data science team, and the relevant business owner. My first step would be to suspend or limit use of the system pending investigation. I would then conduct a bias audit using disaggregated performance data, review the training data for historical bias, and assess whether any adverse employment decisions were made. Findings would be reported to senior leadership and potentially to regulators depending on jurisdiction. Remediation would include model retraining, outcome review for affected candidates, and policy updates.

Scenario Q27: Your organisation is preparing for an ISO 42001 certification audit. Where do you start?

Sample Answer: I would begin with a gap assessment against the standard's requirements — mapping existing policies, controls, and processes to the clauses of ISO 42001 and Annex A. From there, I would prioritise gaps by risk and effort, develop a remediation roadmap, and ensure leadership has formally committed to the AIMS scope and objectives. I would also conduct internal audits and a management review before the certification audit to surface and address any non-conformities in advance.

Scenario Q28: Your board has asked for a briefing on the organisation's AI risk exposure. How do you prepare it?

Sample Answer: I would structure the briefing around three areas: what AI systems we operate and their risk classifications, what controls and governance structures are in place, and where the material gaps or emerging risks lie. I would translate technical risk concepts into business language — framing risks in terms of regulatory exposure, reputational impact, and operational resilience. I would include a heatmap of AI systems by risk tier, a summary of open audit findings, and a clear set of recommended actions with owners and timelines. The goal is to give the board the information they need to exercise meaningful oversight — not a technical deep-dive.

Shadow AI Questions

Q29: What is Shadow AI and what risks does it create?

Sample Answer: Shadow AI refers to AI tools and applications used by employees without the knowledge or approval of IT, legal, or governance functions. It mirrors the earlier challenge of shadow IT but with amplified risks. These risks include data leakage when employees input sensitive or personal data into consumer AI tools, model outputs being used in consequential decisions without validation, regulatory non-compliance (particularly under GDPR or the EU AI Act), and an incomplete AI inventory that undermines the organisation's ability to manage its AI risk posture. The proliferation of free and freemium generative AI tools has made shadow AI a near-universal challenge.

Q30: How would you build a Shadow AI detection and response programme?

Sample Answer: Detection requires a combination of technical controls — such as DNS filtering, DLP tools that flag uploads to known AI platforms, and browser extension monitoring — and cultural controls such as mandatory disclosure mechanisms and manager training. Response should be graduated: not punitive for first-time, good-faith use, but firm where data exposure or policy breach has occurred. Crucially, any shadow AI programme must be accompanied by a fast, low-friction approved AI pathway so employees have a legitimate alternative.

Q31: How do you balance restricting Shadow AI with maintaining employee productivity?

Sample Answer: The key is to be an enabler, not a blocker. If the only AI governance output employees experience is restrictions and denials, shadow AI will increase. I would invest in building an approved AI catalogue with clearly communicated safe-use guidance, streamline the intake process so business units can get tools reviewed quickly, and run awareness campaigns that frame governance as supporting — not opposing — innovation. Governance friction should be proportional to risk: low-risk tools with no personal data processing should have a lighter-touch approval pathway.

Q32: What policy controls would you put in place to address Shadow AI?

Sample Answer: I would implement an Acceptable Use Policy for AI that explicitly defines what is permitted, what requires approval, and what is prohibited. This should be supported by a mandatory AI tool disclosure process, training modules on data handling risks, and clear consequences for wilful non-compliance. The policy should be reviewed at least annually and updated as the AI tool landscape evolves — which it does rapidly.

AI Ethics & Responsible AI Questions

Q33: What does algorithmic fairness mean and how do you measure it?

Sample Answer: Algorithmic fairness refers to the absence of unjustified discriminatory effects in AI outputs across demographic groups. There are multiple technical definitions — demographic parity, equalised odds, individual fairness — and they can conflict with each other. Measurement requires disaggregated performance analysis: testing model outputs across protected characteristics to identify differential error rates. The appropriate fairness metric depends on the use case and the legal context, which is why involving legal and ethics experts alongside data scientists is essential.

Q34: What is explainability in AI and when is it required?

Sample Answer: Explainability refers to the ability to describe, in terms understandable to a given audience, why an AI system produced a particular output. It is required — to varying degrees — in regulated contexts such as credit decisions (GDPR's right to explanation), medical AI, and EU AI Act high-risk systems. Techniques include LIME, SHAP, and attention visualisation, but the right level of explainability depends on the audience: regulators need audit trails, end users need plain-language explanations, and data scientists need technical interpretability.

Q35: How do you operationalise an AI ethics policy?

Sample Answer: An AI ethics policy only has value if it is embedded in processes, not just documented on paper. Operationalisation requires: ethics review as a mandatory stage in the AI development lifecycle, trained reviewers who can apply principles to real systems, clear escalation paths for ethical concerns, and a feedback loop from incidents back to policy. I would also advocate for diverse representation in ethics review panels to prevent homogeneous blind spots.

Q36: How would you handle a situation where a commercially valuable AI system is found to produce biased outputs?

Sample Answer: "This is the central tension in responsible AI — commercial pressure versus ethical obligation. My approach would be to document the bias findings rigorously, assess the severity of harm to affected individuals, and present leadership with a clear risk-benefit analysis. I would not recommend continuing deployment as-is if the bias creates legal exposure or material harm. Short-term, I would implement mitigating controls — human review of high-stakes decisions, output filtering — while the model is remediated. The business case for getting this right is strong: the reputational and regulatory cost of a bias incident far outweighs a delayed launch."

Human Oversight Questions

Q37: What is the difference between human-in-the-loop and human-on-the-loop oversight?

Sample Answer: Human-in-the-loop means a human reviews and approves each AI output before it is acted upon — providing the highest level of oversight but also the highest operational cost. Human-on-the-loop means the AI system operates autonomously but a human monitor can intervene if anomalies are detected. The appropriate model depends on the stakes involved: high-risk decisions such as medical diagnoses or loan rejections warrant stronger in-the-loop oversight, while lower-stakes automated workflows may function adequately with on-the-loop monitoring.

Q38: How do you prevent human oversight from becoming a rubber stamp?

Sample Answer: Meaningful oversight requires that human reviewers have genuine ability and authority to question or override AI outputs — not just click an approve button. I would design oversight roles with adequate time per review, provide reviewers with context about the AI system's limitations and known failure modes, track override rates as a metric (very low rates may indicate rubber-stamping), and include oversight effectiveness as a control in internal audits. Training and cognitive bias awareness for reviewers are also critical.

Q39: What oversight obligations does the EU AI Act impose on high-risk AI systems?

Sample Answer: For high-risk AI systems, the EU AI Act requires providers to design systems with built-in human oversight measures — including the ability to stop, pause, or override the system. Deployers must assign human overseers with the necessary competence and authority to interpret outputs. The system must also produce audit logs sufficient to enable post-hoc oversight and investigation. These obligations must be documented and verifiable at the point of conformity assessment.

Q40: How does automation bias undermine human oversight?

Sample Answer: Automation bias is the tendency of humans to defer to automated or AI-generated recommendations, even when they should apply independent judgement. It is one of the most significant threats to meaningful human oversight. Mitigations include training reviewers to recognise and resist this tendency, providing contrarian decision-support (e.g. showing what the decision would be without the AI), randomised audits of override behaviour, and ensuring review interfaces do not present AI outputs as a default to confirm rather than a recommendation to evaluate.

Final Preparation: Tips for Interview Success

Knowing the right answers is only part of your preparation. How you structure and deliver your responses, how you demonstrate commercial awareness, and how you show that your governance expertise translates into real organisational value are equally important. Use these final tips to ensure you walk into your interview with confidence, clarity, and competitive edge.

Structure Every Answer

Use STAR (Situation, Task, Action, Result) for behavioural questions and a clear Define–Apply–Connect structure for conceptual questions. Opening with a crisp definition and closing with a real-world implication shows both depth and practical grounding.

Know the Regulatory Landscape

AI governance is a fast-moving field. Demonstrate that you follow developments: the EU AI Act implementation timeline, NIST AI RMF updates, ISO 42001 certification momentum, and sector-specific guidance from the FCA, ICO, or equivalent regulators in your jurisdiction.

Show Cross-Functional Thinking

The best AI governance practitioners speak the language of legal, data science, product, and the business. Illustrate that you can work across silos — drafting policy that lawyers can approve, risk registers that data scientists find credible, and briefings that board members find actionable.

Close With Curiosity

Prepare two or three incisive questions for your interviewers: How mature is the current AI governance programme? What is the board's appetite for AI risk? How is the team responding to the EU AI Act? Questions like these signal genuine engagement and strategic thinking.

40+

Questions Covered

Across all major AI governance domains

6

Core Frameworks

ISO 42001, NIST AI RMF, EU AI Act, GDPR, responsible AI principles, shadow AI

100%

Practical Focus

Every question includes a structured sample answer you can adapt

📌 **You are ready.** AI governance is one of the fastest-growing professional disciplines of this decade. The organisations hiring right now need people who combine technical literacy, regulatory awareness, and practical judgement. This guide has equipped you with all three. Go show them what you know.



Globally Recognised Certification Body
190+ countries · 2,50,000+ Certified Professionals



CERTIFIED AI GOVERNANCE PROFESSIONAL

ABOUT GSDC CERTIFICATION



EBOOK

Extensive and exclusive Ebook created by world's experts to help you with understanding core concepts.



LEARNING MATERIALS

Get access to learning materials such as videos, ebooks, templates, and practice exams, which will help you clear the certification exam.



CREATED BY EXPERTS

GSDC certifications are created and authored by world's leading experts in the field.

LEARNING OBJECTIVE

- Gain a strong understanding of AI governance principles and practices
- Apply AI governance concepts to real-world organizational needs
- Develop skills to identify and manage AI-related risks
- Build effective AI governance practices for responsible AI adoption

Enroll now with the code **LEARN20** To avail **20%** discount

Enroll Now

www.gsdcouncil.org 