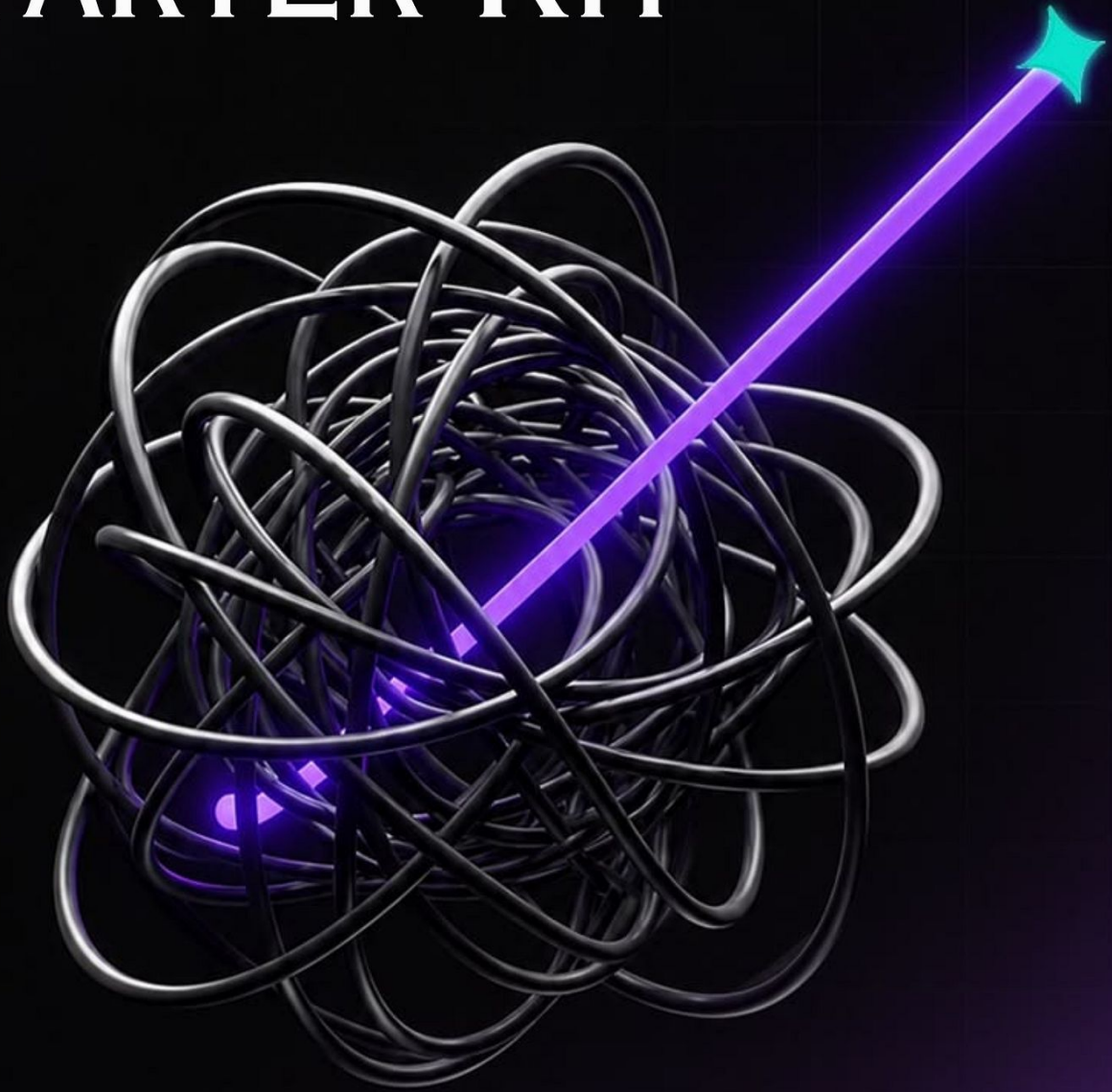


Globally Recognised Certification Body
190+ countries · 2,50,000+ Certified Professionals

AI GOVERNANCE STARTER KIT



What Is AI Governance?

AI governance is the set of policies, processes, and oversight structures that ensure AI systems are developed and used safely, fairly, and accountably. It gives organisations a way to manage the risks that come with AI — and to build trust with the people affected by it.

Think of it as the rulebook and referee for how AI is built, deployed, and monitored in your organisation.

→ It sets clear accountability

AI governance defines who is responsible for AI decisions — from development teams to senior leadership — so there is always someone answerable when things go wrong.

→ It manages risk before it escalates

By identifying potential harms early — such as biased outputs, data misuse, or regulatory breaches — governance helps organisations act proactively rather than reactively.

→ It keeps AI aligned with your values

Governance ensures that AI systems reflect the organisation's ethical principles and legal obligations, not just technical capabilities or commercial pressures.

→ It builds trust — internally and externally

Customers, regulators, and employees are more likely to accept AI when they can see it is being overseen responsibly and transparently.

The Core Pillars of AI Governance

A reference guide to the six foundations every organisation needs to govern AI responsibly.

1. Policy & Rules

Clear written policies define how AI can and cannot be used across the organisation.

2. Roles & Accountability

Named individuals and teams are responsible for AI decisions at every level.

3. Risk Management

Potential harms are identified, assessed, and mitigated before AI systems go live.

4. Transparency & Explainability

AI outputs can be understood and explained to the people affected by them.

5. Human Oversight

People remain in the loop to review, challenge, and override AI decisions where needed.

6. Monitoring & Measurement

AI systems are continuously tracked against agreed standards and performance indicators.

AI Governance Policy: What It Should Cover

A good AI governance policy doesn't need to be lengthy — it needs to be clear, actionable, and understood by everyone it applies to. Here are the seven elements every policy should include.

1. Purpose & Scope

State *why* the policy exists and *who* it applies to. Include all staff, contractors, and third parties using AI tools on behalf of the organisation. Name the AI systems or categories of use it covers.

2. Acceptable Use Rules

Define what AI *can* be used for — drafting, summarising, analysis, code generation, etc. Set clear boundaries around sensitivity levels and approval requirements for new use cases.

3. Prohibited Uses

Explicitly list what AI must *not* be used for — such as making autonomous decisions about individuals, generating deceptive content, or processing personal data without authorisation.

4. Data Handling

Specify what data can and cannot be entered into AI tools. Address personal data, confidential information, and client data. Reference relevant data protection obligations (e.g. GDPR).

5. Human Oversight Requirements

Identify decision types that require human review before action is taken. State who is responsible for reviewing AI outputs and what sign-off is needed for high-risk use cases.

6. Incident Reporting

Set out how staff should report AI-related errors, harms, or unexpected outputs. Name the reporting channel, timeframe, and who investigates. Link to the broader incident management process.

7. Review Cadence

State how often the policy will be reviewed — at minimum annually, or after significant AI-related incidents or regulatory changes. Name the policy owner responsible for keeping it current.

Roles & Accountability in AI Governance

Clear role definitions prevent gaps, duplication, and accountability gaps — ensuring that someone is always responsible when AI decisions need to be made or reviewed.



AI Governance Lead

Owns the overall AI governance framework, coordinates across functions, and ensures policy is implemented and kept current.



AI Risk Owner

Identifies, assesses, and monitors risks associated with AI systems, escalating concerns and maintaining the risk register.



Data Protection Officer

Ensures AI use complies with data protection law, advises on personal data handling, and leads on privacy impact assessments.



Business Unit AI Champion

Acts as the local point of contact for AI governance within a department, supporting adoption and flagging operational concerns.



Legal / Compliance Advisor

Reviews AI use cases for legal and regulatory compliance, provides sign-off on high-risk deployments, and tracks emerging regulation.



Technical AI Owner

Accountable for the technical integrity, security, and performance of AI systems, including vendor management and model documentation.



Board / Executive Sponsor

Provides strategic direction and top-level accountability for AI governance, ensuring it receives the resources and attention it requires.

AI Inventory: Know What You Have

You cannot govern what you cannot see. Before you can manage risk, ensure compliance, or maintain oversight, you need a clear picture of every AI system operating across your organisation. An AI System Inventory is the essential first step — a single, structured record that brings visibility to what is often scattered, informal, or simply unknown.

What to Record for Each AI System

For every AI tool, platform, model, or embedded feature in use, capture the following fields:

System Name

The official name of the AI tool or system, with a brief description of what it does and how it is used in your organisation.

Purpose / Use Case

The specific business function it supports — e.g. recruitment screening, customer service, fraud detection, content generation.

Owner

The named individual or team accountable for this system's use, performance, and governance within your organisation.

Data Inputs

What data the system receives and processes — including whether it handles personal, sensitive, or regulated data.

Risk Level

An initial classification of High, Medium, or Low risk based on factors such as consequential decision-making, autonomy, and sensitive data use.

Regulatory Category

Any applicable legal or regulatory frameworks — such as UK GDPR, the EU AI Act, or sector-specific rules — that govern this system's use.

Vendor / Source

Whether the system is third-party licensed, built internally, or accessed via API — and who the supplier is where applicable.

Review Date

When the entry was last reviewed and when the next scheduled review is due, to ensure the inventory stays current as systems evolve.

 Start simple. Even a shared spreadsheet with these eight fields is a meaningful step forward. The goal is visibility — not perfection on day one.

AI Risk Management: How It Works

An **AI Risk Register** is a living document that records every known risk associated with your AI systems — what the risk is, how serious it could be, who is responsible for it, and what is being done about it. It turns risk management from a vague intention into a practical, trackable process. Think of it as the central logbook that keeps your AI governance programme honest and accountable.

The register works by following a simple five-step cycle. Each risk moves through these steps — and returns to the beginning as circumstances change, new systems are deployed, or controls are tested.

01

Identify the Risk

Name the specific risk — for example, a hiring algorithm that may disadvantage certain applicants, or a chatbot that could leak sensitive data. Be concrete. Vague risks are hard to manage.

02

Assess Likelihood and Impact

Rate how likely the risk is to occur and how serious the consequences would be. This gives you an **inherent risk rating** — your starting point before any controls are in place.

03

Assign an Owner

Every risk needs a named person who is accountable for it. Without clear ownership, risks drift. The owner is responsible for ensuring controls are in place and that the risk is reviewed on schedule.

04

Apply Controls

Put measures in place to reduce the risk — such as bias audits, access restrictions, human oversight checkpoints, or contractual safeguards with vendors. After controls, reassess to get the **residual risk rating**.

05

Monitor and Review

AI systems change over time — and so does the risk landscape. Schedule regular reviews to check whether controls are still working, whether new risks have emerged, and whether the residual rating remains acceptable.

- ❏ The register is only useful if it is kept up to date. Build reviews into your governance calendar — not just when something goes wrong.

Assessing Risk in Individual AI Systems

Before deploying any AI system, you should conduct a structured risk assessment. This is a focused evaluation of that specific system — not AI in general. It asks seven practical questions. Work through each one, document your answers, and identify any gaps that need to be addressed before go-live.

1 What does the system do and decide?

- 1 Describe exactly what the system does in plain terms. Does it generate content, score applicants, flag transactions, recommend actions, or make autonomous decisions? Clarify where it has discretion and where a human has the final say. The clearer you are here, the easier everything else becomes.

2 Who is affected?

- 2 Identify every group touched by this system — directly or indirectly. This includes employees, customers, applicants, patients, or members of the public. Consider whether any groups are more vulnerable to harm, and whether the system could affect people who have no knowledge it is being used.

3 What data does it use?

- 3 List the data inputs the system relies on. Is any of it personal data? Is it sensitive — such as health, financial, or demographic information? Where does the data come from, how current is it, and how might gaps or biases in the data affect outputs? Poor data produces poor — and potentially discriminatory — results.

4 How transparent is it?

- 4 Can the system explain why it reached a particular output? If it makes a decision that affects someone, can that decision be meaningfully explained to them? If the answer is no, that is itself a risk — especially in regulated sectors or where legal rights are involved.

5 What could go wrong?

- 5 Think through realistic failure modes. What happens if the system produces incorrect outputs at scale? Could it be manipulated or gamed? What are the consequences of errors — are they minor inconveniences or serious harms? This is where you identify your key risks and rate them by likelihood and severity.

6 What controls are in place?

- 6 Document the safeguards that reduce each identified risk — for example, regular accuracy testing, bias audits, data access restrictions, or contractual obligations on suppliers. For each control, assess whether it is robust enough, who is responsible for it, and how you will know if it stops working.

7 Is human oversight built in?

- 7 Identify where humans are involved in reviewing, challenging, or overriding the system's outputs. Is there a clear escalation path if something looks wrong? Are the people doing oversight sufficiently informed to exercise genuine judgement — or are they simply rubber-stamping automated decisions?

Transparency & Explainability in AI

Two terms come up repeatedly in AI governance: **transparency** and **explainability**. They are related but distinct.

Transparency means being open about the fact that an AI system is being used — what it does, what data it relies on, and who is responsible for it. **Explainability** means being able to describe, in plain terms, why a system produced a particular output or decision. Together, they allow people to understand, challenge, and trust AI systems — and they are increasingly required by law.

If you cannot explain how a system works or justify the decisions it makes, that is itself a governance risk. Here is how to address it in practice.

Practical Guidance

Document how systems work

For every AI system in use, maintain clear documentation covering: what the system does, what inputs it uses, how it reaches outputs, and what its known limitations are. This does not need to be highly technical — it needs to be accurate and accessible to the people responsible for oversight.

Be able to explain decisions to affected people

If an AI system produces a decision that affects someone — a rejection, a score, a recommendation — you should be able to explain the basis for that decision in plain language. Where legal rights are involved, this is often a legal requirement, not just good practice. Build explanation capability into the system from the start, not as an afterthought.

Maintain audit trails

Keep records of what decisions were made, when, by which system, using what data, and with what human involvement. Audit trails allow you to investigate problems, demonstrate accountability to regulators, and identify patterns of error or bias over time. Agree in advance how long records are retained and who has access.

Communicate AI use to stakeholders

Be open with employees, customers, and other affected parties about where AI is being used and what role it plays. This includes privacy notices, staff briefings, and board-level reporting. Stakeholders who discover AI use unexpectedly — especially in high-stakes contexts — lose trust quickly. Proactive communication prevents that.

Human Oversight: Keeping People in the Loop

Human oversight means that a person — not just an algorithm — remains responsible for decisions that matter. In AI governance, it is the principle that AI systems should inform, assist, or recommend, but that human judgement should be retained wherever the stakes are significant. Oversight is not about slowing things down. It is about ensuring that accountability stays with people who can be held responsible, who understand context, and who can recognise when something has gone wrong.

The following five principles set out how to make human oversight real and effective — not just a formality on paper.

When Oversight Is Required

Oversight is required whenever an AI system produces an output that affects a person's rights, opportunities, or wellbeing — hiring decisions, credit assessments, medical triage, disciplinary actions, and similar high-stakes contexts. It is also required when a system operates in a regulated environment, when outputs are novel or uncertain, or when the consequences of error are serious or hard to reverse. If in doubt, apply oversight.

The threshold for removing it should be high.

What Meaningful Oversight Looks Like

Meaningful oversight is not simply having a human present. It means the person reviewing an AI output has enough time, information, and authority to genuinely evaluate it — and to override it if necessary. They should understand what the system was asked to do, what data it used, and what its known limitations are. Oversight without that context is oversight in name only.

How to Avoid Rubber-Stamping

Rubber-stamping happens when reviewers approve AI outputs automatically, without real scrutiny — often because they feel pressure to move quickly, lack the information needed to challenge the system, or assume the AI is more reliable than their own judgement. Counter this by setting realistic review times, providing reviewers with plain-language explanations of how outputs were reached, tracking override rates, and making it culturally safe to push back on AI recommendations.

Escalation Paths

Not every reviewer will have the authority or expertise to resolve every edge case. Define clear escalation routes in advance: who should a reviewer contact if they are uncertain about an output, if they suspect a systemic error, or if the AI recommendation conflicts with their professional judgement? Escalation paths should be easy to use, well-known within the organisation, and free of any stigma around raising concerns.

Documenting Oversight Decisions

Keep a record of where human oversight was applied, who reviewed the output, what decision was reached, and — crucially — whether the AI recommendation was followed or overridden. If it was overridden, record why. These records serve multiple purposes: they demonstrate accountability to regulators, allow you to spot patterns of AI error, and provide evidence that oversight is genuine rather than procedural. Decide in advance how long records are retained and who can access them.

Shadow AI: The Hidden Risk

Shadow AI refers to the use of AI tools by employees without the knowledge or approval of their organisation. Just as "shadow IT" described unsanctioned software and devices, shadow AI is the same problem in a new form — and it carries serious governance risks that are easy to overlook precisely because the activity is invisible to those responsible for managing it.

What Is Shadow AI?

Shadow AI occurs when staff use consumer AI tools — chatbots, writing assistants, image generators, code tools — for work purposes without going through any approval process. The tools may be free, easy to access, and genuinely useful. Employees are not usually acting in bad faith. They are trying to get things done. But the organisation has no visibility into what data is being shared, what outputs are being relied upon, or what obligations may be triggered.

Why Employees Use Unauthorised Tools

Approved tools are sometimes slow, limited, or hard to access. Consumer AI products are fast, capable, and free. When the official route feels like friction, people find workarounds. Shadow AI grows where there is a gap between what employees need and what the organisation provides — and where there is no clear policy explaining why certain tools require approval.

The Risks It Creates

Data Leakage

When employees paste documents, client data, or internal communications into a consumer AI tool, that data may be used to train third-party models or stored on external servers. Once shared, it cannot be recalled.

Compliance Breaches

Processing personal data through unapproved systems may violate GDPR, sector regulations, or contractual obligations — without anyone in the organisation being aware it is happening.

Unvetted Outputs

AI tools used outside any governance framework produce outputs with no quality checks, no audit trail, and no accountability. Errors or biases in those outputs may flow directly into decisions.

Five Steps to Detect and Manage Shadow AI

1 Build Awareness

Most employees using shadow AI are unaware of the risks. Start with clear, jargon-free communication explaining what shadow AI is, why it matters, and what the rules are. Avoid a tone that treats staff as suspects — the goal is informed behaviour, not fear.

2 Publish an Approved Tool List

Maintain a regularly updated list of AI tools that have been reviewed and approved for specific use cases. Make it easy to find and easy to use. If the approved list does not meet staff needs, that is a signal to expand it — not to enforce restrictions more harshly.

3 Create Safe Reporting Channels

Employees who discover shadow AI use — including their own past use — should be able to report it without fear of punishment. Reporting channels should be accessible, confidential where appropriate, and clearly separated from disciplinary processes.

4 Enforce Policy Consistently

Policy is only credible if it is applied. Set out clearly what is and is not permitted, what happens when rules are breached, and how decisions are made. Inconsistent enforcement — where some teams are held to account and others are not — undermines the entire framework.

5 Conduct Regular Audits

Periodically review what AI tools are actually in use across the organisation, not just what has been approved. Network monitoring, procurement records, and staff surveys can all surface shadow AI. Use audit findings to improve the approved tool list and close the gaps that drive unsanctioned use.

AI Incident Management: What to Do When Things Go Wrong

An AI incident is any situation where an AI system behaves in an unexpected, harmful, or unintended way — for example, producing incorrect outputs that affect a decision, exposing personal data, or acting outside the boundaries it was designed to operate within. Incidents range from minor errors to serious failures with legal consequences.

Having a clear, documented response process matters for two reasons. First, it limits harm — a fast, coordinated response reduces the damage an incident causes. Second, it demonstrates to regulators, customers, and staff that your organisation takes AI governance seriously and has the mechanisms in place to act when something goes wrong.

The Six-Step Response Process

01

Detect & Report

Identify that something has gone wrong and log it immediately. Anyone who notices an AI-related problem should be able to report it quickly and without fear of blame. Record the what, when, where, and who.

02

Classify Severity

Assess how serious the incident is. Consider the nature of the harm, how many people are affected, and whether regulatory obligations are triggered. Severity determines how urgently you need to act and who needs to be informed.

03

Contain the Impact

Take immediate steps to stop the problem from getting worse. This might mean pausing the AI system, restricting access, or alerting affected parties. Speed matters here — containment limits the damage.

04

Investigate Root Cause

Once the immediate situation is under control, work out what actually went wrong and why. Was it a data problem, a model failure, a process gap, or a human error? Understanding the cause is essential to fixing it properly.

05

Remediate & Recover

Put the fix in place. This may involve correcting outputs, updating the system, notifying regulators or affected individuals, or making changes to processes. Document every action taken and by whom.

06

Review & Learn

After the incident is resolved, hold a structured review. What could have been caught earlier? What needs to change to prevent recurrence? Update policies, training, or technical controls as needed.

Incident Severity Levels: A Quick Reference

Use this framework to classify AI incidents consistently and trigger the right response. Severity determines who is notified, how fast action is required, and what escalation path applies.

Level	Description	Example	Response Timeframe
Level 1 Critical	Serious harm to individuals, a confirmed regulatory breach, or systemic failure. Immediate executive escalation and external notification required.	AI model produces discriminatory loan decisions affecting hundreds of applicants; personal data exposed at scale.	Immediate. Contain within hours. Notify regulators within 72 hours where required.
Level 2 High	Significant performance failure or material data exposure. Poses serious risk but harm is not yet widespread or irreversible.	AI recommendation system fails silently, resulting in incorrect medical triage flags for a subset of patients.	Urgent. Senior management escalation and response plan within 24 hours.
Level 3 Medium	Limited impact with potential to escalate if not addressed. No immediate harm, but the issue requires a managed, documented response.	Chatbot repeatedly gives outdated policy information; a small number of users receive incorrect guidance.	Managed. Governance team review and remediation plan within 5 business days.
Level 4 Low	Minor anomaly or near-miss with no immediate user impact. Useful as an early warning signal.	AI output flagged internally as unusual but within acceptable bounds; no user-facing effect detected.	Routine. Log, monitor, and review in the next scheduled governance cycle.

 **Tip:** When in doubt, classify up. It is always easier to downgrade an incident after review than to recover from a delayed response to a serious one.

Measuring AI Governance: Key Metrics to Track

A governance programme you cannot measure is a governance programme you cannot improve. Tracking the right metrics helps leadership understand whether controls are working, where gaps exist, and how to make the case for continued investment. You do not need to track everything — start with a small set of meaningful indicators and build from there.

● Risk & Compliance Metrics

- **% of AI systems with completed risk assessments** — Are you actually reviewing the systems you operate?
- **Number of open high-risk items** — How many unresolved issues could cause serious harm or regulatory exposure?
- **Policy compliance rate by business unit** — Which teams are following governance requirements and which are not?
- **Regulatory audit outcomes** — Are findings improving, stable, or getting worse over time?
- **Time from deployment to governance registration** — Are new AI systems being captured before they go live?

● Oversight & Process Metrics

- **Human override rate** — How often are AI decisions reviewed or overturned by humans? Very low rates may signal rubber-stamping.
- **Incident response time** — How quickly are governance incidents contained and resolved after detection?
- **% of high-risk systems with active oversight mechanisms** — Are your riskiest systems actually being monitored?
- **Number of AI projects paused or blocked on governance grounds** — A healthy signal that governance has real authority.
- **Average time to close governance incidents** — A measure of process efficiency and follow-through.

● Culture & Awareness Metrics

- **AI governance training completion rate** — Are staff actually completing required training, or is it being skipped?
- **Shadow AI reports per quarter** — How many unapproved AI tools are being identified? Rising numbers may reflect better detection, not worse behaviour.
- **Stakeholder confidence scores** — Do internal teams and external partners trust the governance programme?
- **Governance programme maturity score** — Benchmarked against industry standards to track progress over time.

① **Practical tip:** Choose 3–5 metrics to start. Report them consistently to senior leadership on a fixed cadence. Metrics only drive improvement when someone is accountable for the numbers.

AI Governance & Regulation: What You Need to Know

Organisations deploying AI are increasingly subject to a growing set of regulatory requirements and international standards. You do not need to be a legal expert — but you do need to know which frameworks apply to you and what they require in practice.

EU AI Act

What it is: The world's first comprehensive AI law, classifying AI systems into risk tiers — unacceptable, high, limited, and minimal risk — with requirements that scale with the level of risk.

What it means for you: If you operate high-risk AI (e.g. in hiring, credit, healthcare, or law enforcement), you will need conformity assessments, human oversight mechanisms, technical documentation, and registration in an EU database before deployment.

UK AI Regulation

What it is: The UK has taken a principles-based, sector-led approach — rather than one overarching AI law, existing regulators (ICO, FCA, CQC, etc.) apply AI-specific guidance within their domains, guided by five cross-sector principles.

What it means for you: You need to understand which sector regulator oversees your AI use cases and comply with their published AI guidance, while demonstrating safety, transparency, fairness, accountability, and contestability.

GDPR & Data Protection

What it is: The General Data Protection Regulation governs how personal data is processed across the EU and UK. Article 22 specifically addresses automated decision-making, giving individuals the right not to be subject to solely automated decisions with significant effects.

What it means for you: Any AI system that makes or significantly influences decisions about individuals — such as loan approvals, job screening, or benefits eligibility — must provide human review on request, clear explanations, and a meaningful right to challenge the outcome.

ISO 42001

What it is: An international standard published in 2023 that specifies requirements for establishing, implementing, maintaining, and continually improving an AI Management System (AIMS) within an organisation.

What it means for you: ISO 42001 gives you a structured, certifiable framework for AI governance — covering risk management, responsible use policies, and accountability structures — and can serve as evidence of good practice to regulators, customers, and partners.

NIST AI RMF

What it is: The US National Institute of Standards and Technology's AI Risk Management Framework, a voluntary framework organised around four functions — Govern, Map, Measure, and Manage — designed to help organisations identify and address AI-related risks.

What it means for you: Even if you are not a US organisation, the NIST AI RMF is widely referenced as a practical, non-prescriptive guide to building AI risk management processes and is increasingly cited in procurement and due diligence requirements.

③ Where to start: Identify which frameworks apply to your organisation based on jurisdiction, sector, and the nature of your AI use cases — then prioritise compliance with legally binding requirements (EU AI Act, GDPR) before adopting voluntary standards.

Getting Started: A 90-Day Governance Plan

AI governance does not need to be built overnight. This phased plan gives your organisation a practical starting point — moving from initial discovery through to embedded, ongoing governance in three structured phases.

Phase 1 — Days 1–30: Foundation

- Inventory all AI systems in use across the organisation
- Identify shadow AI and unsanctioned tools
- Assign governance roles and confirm accountability (RACI)
- Draft or review your AI Governance Policy
- Conduct initial risk classifications for each system

Phase 2 — Days 31–60: Structure

- Complete risk assessments for high-risk AI systems
- Set up human oversight processes for critical decisions
- Establish incident reporting and escalation procedures
- Populate your AI Risk Register
- Address shadow AI findings with clear use policies

Phase 3 — Days 61–90: Embed

- Define and baseline governance KPIs
- Deliver staff awareness and training on AI governance
- Report findings and roadmap to senior leadership
- Schedule ongoing governance reviews and audit cadence
- Identify priorities for the next quarter

📌 Use this plan as a starting framework — adapt the pace and priorities to your organisation's size, sector, and existing governance maturity.

Common AI Governance Mistakes to Avoid

Even well-intentioned governance programmes can fall into predictable traps. These are the six most common mistakes organisations make — and how to avoid them.

1. Treating governance as a one-time exercise

AI governance is not a project with an end date. AI systems evolve, new tools are adopted, and regulations change. **Fix:** Schedule quarterly reviews of your risk register, policy documents, and approved tool list as standing agenda items — not one-off tasks.

2. Assigning governance to one person with no support

Placing AI governance solely on a single individual creates a bottleneck and a single point of failure. **Fix:** Distribute accountability using a RACI model. Governance needs champions in Legal, IT, HR, and business units — not just one overwhelmed owner.

3. Focusing only on technical controls and ignoring culture

Policies and checklists only work if people understand and follow them. Technical guardrails alone will not prevent misuse. **Fix:** Pair every technical control with staff awareness. Make AI governance training part of onboarding and annual compliance cycles.

4. Waiting for regulation before acting

Organisations that delay governance until a law forces their hand are already behind. **Fix:** Use existing voluntary frameworks — NIST AI RMF, ISO 42001 — as your starting point now. Early movers build capability; late movers face rushed, costly remediation.

5. Governing AI in isolation from data and privacy teams

AI governance and data governance are deeply intertwined. Treating them as separate functions leads to duplicated effort and coverage gaps. **Fix:** Involve your Data Protection Officer and data governance leads from the start. Align AI risk assessments with existing DPIA processes wherever possible.

6. Letting the approved tool list go stale

An outdated approved tool list provides false assurance and encourages shadow AI as staff work around tools that no longer meet their needs. **Fix:** Review and update your approved tool list at least every six months. Make it easy for teams to request new tools through a clear, fast intake process.

AI Governance Glossary: Key Terms Explained

A quick reference guide to the essential terms used throughout this Starter Kit. Definitions are written in plain language for practitioners across all functions.

Term	Definition
AI Governance	The policies, processes, roles, and controls an organisation puts in place to ensure AI systems are used responsibly, safely, and in line with legal and ethical standards.
AI Risk Register	A living document that records the AI systems in use, the risks each one poses, and the controls in place to manage those risks. It is reviewed and updated regularly.
RACI Matrix	A responsibility assignment tool that clarifies who is Responsible, Accountable, Consulted, and Informed for each governance task — preventing gaps and overlaps in ownership.
Shadow AI	AI tools and applications that employees use without formal approval or oversight from IT, Legal, or Compliance. Shadow AI creates hidden risks around data security and regulatory compliance.
Human Oversight	The requirement that a human reviews, monitors, or has the ability to intervene in decisions made or supported by an AI system — particularly where those decisions affect people.
Explainability	The ability to describe, in plain language, how an AI system reached a particular output or decision. Explainability is essential for accountability and for building trust with affected individuals.
High-Risk AI	An AI system used in a context where errors could cause significant harm to individuals — such as hiring, credit decisions, healthcare, or law enforcement. High-risk systems require stronger controls and documentation.
Incident	Any event where an AI system behaves unexpectedly, produces harmful outputs, or fails in a way that causes or could cause harm to people or the organisation. Incidents should be logged and investigated.
Residual Risk	The level of risk that remains after controls and mitigations have been applied. No system is entirely risk-free; residual risk must be formally accepted by an appropriate owner.
AI Inventory	A structured record of all AI tools and systems in use across the organisation, including who owns them, what they do, and where they sit in business processes.
Conformity Assessment	A formal evaluation to check whether an AI system meets the requirements set by a regulation or standard — similar in concept to an audit. Required for high-risk AI systems under the EU AI Act.
Model Card	A short document that describes what an AI model does, what data it was trained on, its known limitations, and the contexts it is and is not suitable for. Model cards support transparency and informed use.

Building a Governance Culture: Beyond the Checklist

Templates, registers, and checklists give you the infrastructure for AI governance. But sustainable governance is ultimately a **cultural achievement**. The organisations that get this right don't just follow the rules — they build environments where responsible AI use becomes the norm. Here are five principles to take you from compliance to culture.

1. Leadership Sets the Tone

Governance programmes reflect the behaviours of senior leaders. When executives ask about AI risks, champion ethical use, and allocate real resources to oversight, the rest of the organisation follows. Culture travels top-down.

2. Make It Easy to Do the Right Thing

If governance processes are slow or complicated, people will work around them. Streamline approval workflows, provide clear guidance, and remove friction. Good governance should be the path of least resistance — not an obstacle course.

3. Reward Transparency Over Perfection

Employees need to feel safe raising concerns, logging incidents, and flagging risks — without fear of blame. A culture that punishes mistakes drives problems underground. Celebrate honesty and learning, not just clean records.

4. Embed Governance Into Existing Workflows

Governance that lives in a separate system gets ignored. Integrate risk checks into project kick-offs, procurement decisions, and product reviews. When governance is part of how work gets done, it stops feeling like extra work.

5. Treat Governance as a Capability, Not a Cost

Organisations that invest in AI governance build trust with customers, regulators, and employees. They move faster because they have clear decision-making frameworks. Strong governance is a **competitive advantage** — not a tax on innovation.

AI Governance: Where to Start

Effective AI governance doesn't require perfection — it requires action. These five steps will take you from uncertainty to a programme you can stand behind.

1 Know Your AI Systems

Build your AI inventory. You can't govern what you can't see. Start by mapping every AI tool, model, and system in use across your organisation.

2 Assign Clear Ownership

Use a RACI matrix to define who is responsible, accountable, consulted, and informed for each AI system. Governance without ownership is just paperwork.

3 Assess and Manage Risk

Maintain a living risk register. Identify what could go wrong, score it, and put mitigations in place — before problems become incidents.

4 Keep Humans in the Loop

Build meaningful oversight into your AI workflows. Ensure that consequential decisions are reviewed by people with the authority and context to act on them.

5 Measure and Improve

Track your governance KPIs. Review incidents, audit AI systems regularly, and treat governance as a capability that grows stronger over time.

You don't need to solve every problem at once. Pick up the tools, take the first step, and build from there. Responsible AI governance is a journey — and you're already on your way.



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.gsdccouncil.org 