

A PRACTICAL FRAMEWORK

AI GOVERNANCE

How I evaluate, classify, and document AI applications — from risk assessment to the EU AI Act, using the same governance discipline built for data.



Olivier Bouchez

Data & AI Governance Leader | 30 years in Data

OPEN TO WORK

Same discipline, new failure modes

1

AI systems change after go-live

A dataset is governed once and stays governed. A model keeps learning from new inputs and drifts — governance has to be continuous, not a one-time sign-off.

2

Risk is graduated, not binary

Not every AI use case carries the same stakes. The EU AI Act grades risk from prohibited to minimal — governance effort should scale with that grading, not apply uniformly.

3

The accountability chain gets longer

Provider, deployer, and the business owner using the output all carry obligations. Classic Data Owner / Steward roles still apply — they just need an AI-specific extension.

4

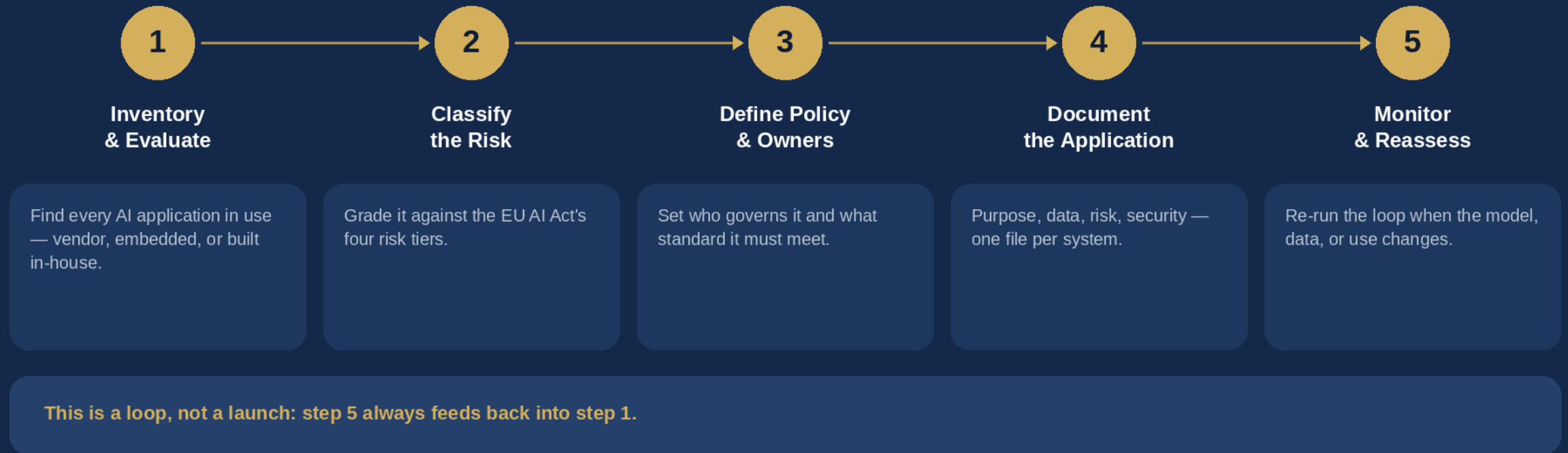
It's now a compliance topic, not just best practice

The EU AI Act made data governance, human oversight, and documentation legal requirements for high-risk systems — with real penalties for getting it wrong.

THE FRAMEWORK

The AI Governance Lifecycle

Five steps, run continuously — not a one-off checklist. Every AI application in scope moves through this loop, and gets revisited whenever it changes.



You can't govern what you haven't found

1

Build the AI Inventory

Discover before you classify

- Vendor and SaaS tools with embedded AI (CRM scoring, HR screening, marketing personalisation).
- Internally built models and scripts — including the ones data teams built for themselves.
- “Shadow AI”: generative AI tools employees already use without an approval process.

2

Evaluate Against Fixed Criteria

Same questions, every application

- Purpose: what decision or output does it produce, and for whom?
- Autonomy: does it recommend, or does it decide and act?
- Reach: how many people are affected, and how reversible is the outcome?

Why this comes first

You cannot classify a risk level, assign an owner, or write documentation for an AI application you don't know exists. The inventory is the same instinct as a data catalog — applied to models instead of datasets.

STEP 2 — MEASURE THE RISK

Not every AI application is governed the same way



Governance effort scales with the tier

Minimal risk needs a light touch. High-risk needs the full documentation set on slide 8.

High-risk in practice (Annex III)

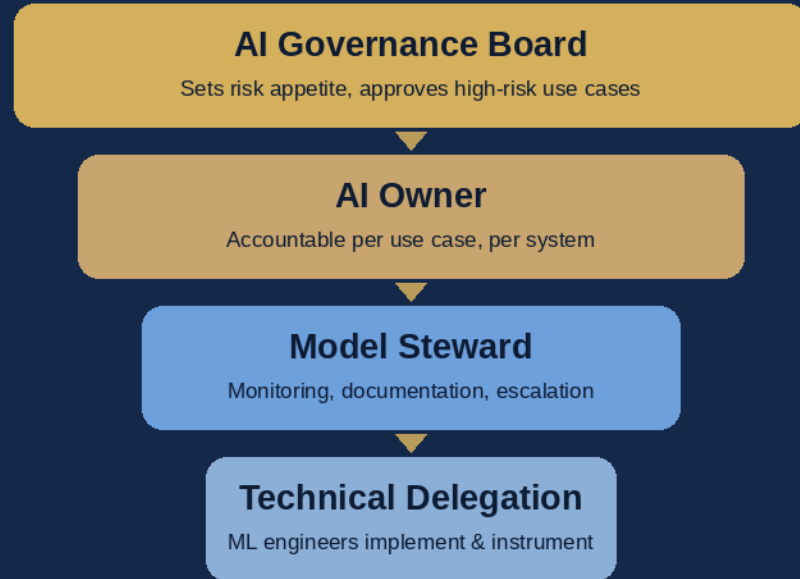
- Credit scoring & creditworthiness evaluation
- Eligibility for essential public or private services
- Recruitment: candidate screening & ranking
- Employee performance monitoring
- Emotion recognition & biometric categorisation
- Risk pricing in health and life insurance

For the Consumer Data domain

Any AI touching consumer eligibility, scoring, or profiling defaults to high-risk until proven otherwise — document the assessment either way.

STEP 3 — DEFINE THE POLICY

Who decides, at each level



The Board owns the policy layer

Same principle as data governance: the policy defines what standards and guidelines must contain — it doesn't write them itself.

What the policy must define

- Approval gates — who signs off before an AI system goes live, by risk tier
- Risk appetite — which use cases the organisation won't pursue at all
- Escalation path — who is notified when a model drifts, fails, or is misused
- Review cadence — how often each system is reassessed, tied to its risk tier
- Vendor requirements — what a third-party AI tool must disclose before approval

Decide the training

AI literacy for anyone using or affected by these systems — a requirement in its own right under the EU AI Act.

Data quality gets four AI-specific extensions

1

Training Data Quality

Representativeness, completeness, and provenance of the data used to train and validate — the same rigour applied to any Consumer Data source, now with a bias lens.

2

Model Performance & Drift

Accuracy and robustness at go-live are a starting point, not a guarantee. Standards must define how and how often performance is re-measured against live data.

3

Fairness & Non-Discrimination

Outcomes tested across protected attributes — measuring whether the system performs consistently across groups, not just on average.

4

Explainability & Human Oversight

Can a human actually understand the output and override it? A model nobody can explain cannot be governed, no matter how accurate it is.

One file per AI system, six sections

● Purpose & Legal Basis

What the system is for, who uses it, and the legal ground it relies on — including its EU AI Act risk classification.

● Risk & Impact Assessment

Who is affected and how; for high-risk systems, a Fundamental Rights Impact Assessment covering the people, not just the system.

● Human Oversight

Who can review, challenge, or override an output — and how that path is actually exercised, not just documented on paper.

● Training Data

Sources, provenance, representativeness, and known limitations of the data used to train, validate and test the model.

● Security Assessment

Access control, adversarial robustness, and incident response — the same discipline as any production system, plus model-specific threats.

● Monitoring & Change Log

Performance and drift metrics over time, plus every material change to data, model, or scope since go-live.

The EU AI Act, in the articles that matter operationally

Art. 5 Prohibited practices

Manipulative, exploitative, or social-scoring AI — banned outright, no documentation exempts it.

Art. 6 / Annex III High-risk classification

Defines which use cases are high-risk by default: employment, credit, essential services, and more.

Art. 10 Data & data governance

Training, validation and test data must be relevant, representative, and error-checked — this is where classic data governance meets the Act directly.

Art. 14 Human oversight

High-risk systems must be designed so a human can understand, monitor, and intervene in their outputs.

Art. 27 Fundamental Rights Impact Assessment

Deployers of certain high-risk systems must assess the impact on the people affected — not just the technical risk.

Art. 50 Transparency obligations

Users must be told they're interacting with AI; deepfakes and synthetic content must be disclosed, even in a lighter form for satire and art.

Phased application after entry into force

6 months — prohibited practices banned

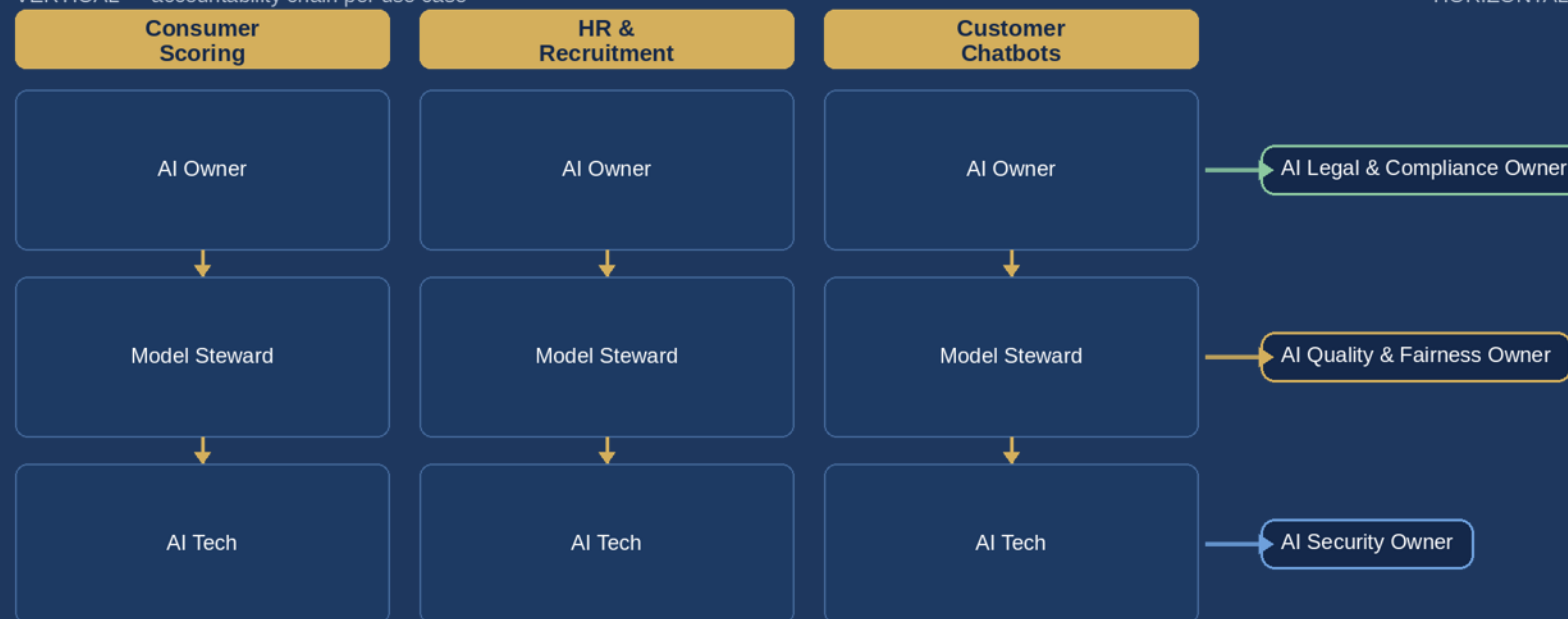
12 months — GPAI obligations apply

24 months — high-risk (Annex III) systems

36 months — high-risk (Annex I) systems

Same matrix logic as data — applied per use case

VERTICAL — accountability chain per use case



HORIZONTAL — cross-cutting AI governance functions

Cross-functional delegation: Data Scientist · ML Engineer · MLOps · Legal / DPO · Security Architect

AI governance is an extension, not a new department

1

Reuse the catalog

The AI inventory lives next to the data catalog, not in a separate tool. If a model consumes a Consumer Data domain, that link is visible in the same place people already look.

2

Reuse the roles

The Data Owner for a domain becomes the AI Owner for the AI applications built on it. No new hierarchy to explain — just an added responsibility on a role people already know.

3

One dashboard for both

A Power BI view tracks approvals, risk tiers, and review dates the same way it already tracks data quality KPIs — one place the sponsor checks, not two.

The AI governance vocabulary, defined simply

● High-Risk AI System

An AI system in a use case listed under Annex III, or a safety component under Annex I law — subject to the Act's full requirements.

● General-Purpose AI (GPAI)

A model trained at scale that performs a wide range of tasks and can be integrated into many downstream systems — e.g. a foundation model.

● Model Drift

The gradual decline in a model's accuracy as real-world data shifts away from what it was trained on — the reason monitoring never stops.

● Fundamental Rights Impact Assessment (FRIA)

A deployer's assessment of how a high-risk system affects the people subject to it, not just its technical performance.

● Human Oversight

The design and operational ability for a person to understand, monitor, and — when needed — override an AI system's output.

● Deepfake

AI-generated or manipulated image, audio, or video content that appears authentic — subject to disclosure obligations under Article 50.

IN CLOSING

Governance that keeps up with the technology

AI governance isn't a separate discipline invented from scratch — it's data governance, extended to cover models that keep learning and a law that made documentation mandatory.

About Olivier Bouchez

30 years in data management and governance across Europe. Most recently Consumer Data Owner at Dun & Bradstreet, now extending that governance discipline to AI applications, risk classification, and EU AI Act compliance.

olivier.bouchez@hotmail.com

linkedin.com/in/olivier-bouchez-bisnode

OPEN TO WORK

Data & AI Governance Leadership roles

Open to conversations across Belgium and Europe.