



Table of contents

Executive Summary	3
1. Introduction	5
2. The EU AI Act: what it requires	6
3. International frameworks – global approaches	12
4. AI governance framework – lifecycle	15
5. The Three Lines Model	18
6. Board oversight and AI risk appetite	20
7. Challenges	21
8. Recommendation for CROs	23
Appendices	24

Executive summary

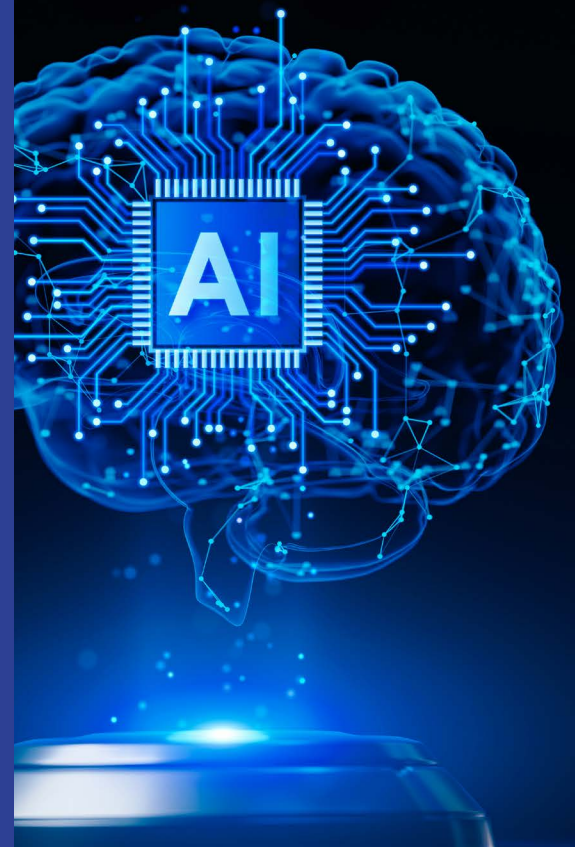
Artificial intelligence (AI) is becoming integral to insurance and reinsurance¹ business models, supporting underwriting, pricing, claims management, fraud analytics and customer interaction across the value chain. This creates material opportunities for efficiency and innovation, but also amplifies conduct, operational, model and reputational risks, particularly where decisions directly affect customer outcomes or rely on complex data and third-party ecosystems.

Good AI governance should be embedded through sound risk management principles rather than focusing primarily on specific technologies, regulations or individual AI use cases. A principles-based approach anchored in established risk management practices provides greater resilience, ensuring that governance remains effective even as AI systems, architectures and business applications change over time.

The EU Artificial Intelligence Act (EU AI Act) introduces an overarching framework that classifies AI systems into different risk tiers and imposes proportionate obligations. For insurers, this is particularly relevant for AI used in risk assessment and pricing in relation to natural persons in the case of life and health insurance, which is classified as high risk and therefore subject to stringent requirements on risk management, data governance, documentation, transparency and human oversight. However, even where AI systems are not formally classified as high risk, existing regulation (including Solvency II, DORA and GDPR) requires insurers to maintain robust governance, data quality, accountability and oversight across all material AI applications. Determining the level of governance and oversight an AI use case requires an **internal risk classification** and should not be dependent on regulation only.

For CROs, the implication is clear: AI risk must be treated as an extension of existing risk categories rather than a standalone topic. The risk management function has a critical role in this transition, ensuring that AI risks are identified, assessed, monitored and reported through existing enterprise risk management processes, while providing independent challenge over material AI use cases and associated control frameworks.

Insurers need to embed AI risk management within their enterprise risk framework, (model) risk governance and conduct oversight, ensuring clear accountability across the AI lifecycle, from development through to deployment, monitoring and decommissioning. In practice, insurers will often act both as provider and deployer, requiring a clear allocation of responsibilities between system design compliance and real-world use and outcomes.



¹ In this paper, the term “insurers” is used broadly to encompass both insurers and reinsurers, unless otherwise specified.

The governance challenge is not only conceptual but operational: who owns the AI inventory, who classifies the systems, who assesses use cases, who validates the outputs, and who escalates when controls fail? Experience shows that these questions are harder to answer than regulatory frameworks itself suggest. And of course, the adoption of AI within risk management to assess, manage and monitor the related AI risks can further help to keep an effective and efficient risk management framework.

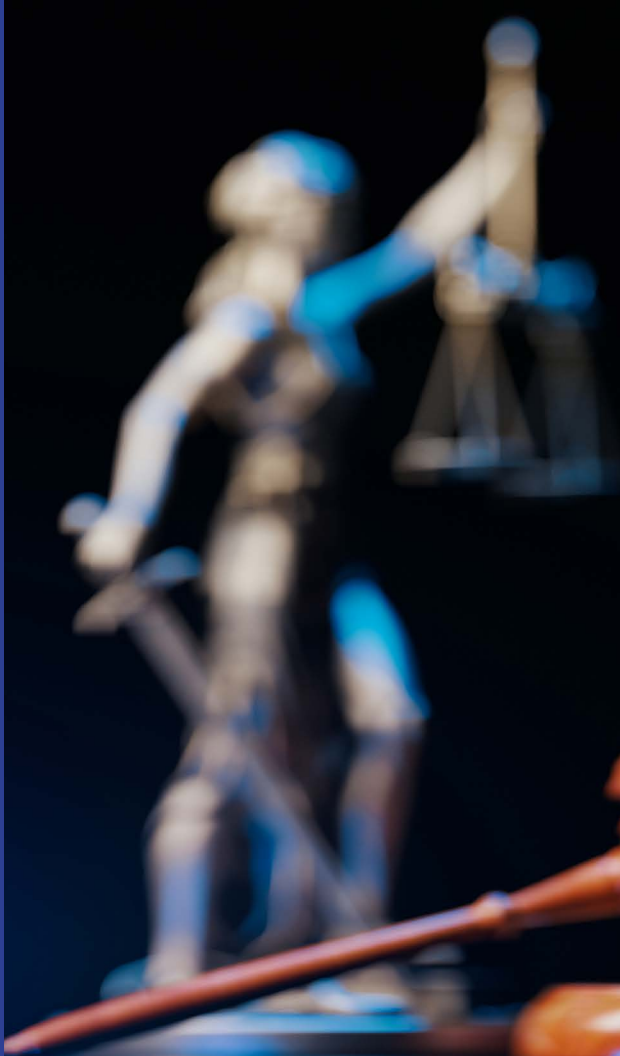
At the same time, insurers face a set of material execution challenges. These include the speed of AI adoption relative to control frameworks, unclear ownership of AI strategy, increasing reliance on third-party providers, the emergence of “shadow AI” use, and the difficulty of ensuring meaningful human oversight and robust validation for complex and non-deterministic models. Traditional validation and control approaches may not be sufficient for newer AI techniques, requiring enhanced (automated) testing, monitoring and governance approaches.

Looking forward, the development of agentic AI further increases this complexity. While the **underlying principles and existing risk management frameworks remain intact**, the practical application shifts from governing individual AI models to overseeing end-to-end processes, where multiple systems interact dynamically. This requires stronger focus on process-level oversight, orchestration risk, and end-to-end validation, particularly where autonomous agents can initiate, adapt or execute business processes without continuous human intervention.

Against this backdrop, insurers should prioritise a structured and pragmatic approach. Immediate actions include establishing a complete AI inventory, performing gap analyses against EU AI Act requirements and other upcoming relevant requirements, defining a clear risk appetite, strengthening third-party AI governance, and building independent validation capabilities. At the same time, insurers must invest in governance structures, board-level understanding, and monitoring capabilities that enable continuous oversight and timely escalation of risks.

Ultimately, competitive advantage will not be determined by the adoption of AI alone, but by the ability to deploy it securely, responsibly and at scale, within a robust governance framework across the AI lifecycle and across the three lines model that protects customers, supports regulatory compliance and maintains trust.

Beyond governing AI, the risk management function can also benefit from AI. Advanced analytics, machine learning and generative AI can enhance risk identification, monitoring, control testing, scenario analysis, event scanning and reporting, enabling risk functions to become more forward-looking, data-driven and scalable.



Note: this is a pre-publication version of a broader practise paper to be published at a later date on AI risk management. This paper focuses mainly on AI regulation and risk management governance.

1. Introduction

Artificial intelligence (AI) has moved into the core of insurance business models, supporting risk selection, pricing, claims handling, fraud detection and customer interaction across the value chain. Properly governed, these applications can improve efficiency, underwriting performance and customer experience, while at the same time increasing the scale and speed at which errors, bias or system failures can materialise and creating stronger dependencies on data, models, technology providers and complex digital processes.

From a risk perspective, AI is not a new standalone risk type, but a force multiplier within existing categories: it increases operational risk through automation and system interdependencies, conduct risk through opaque decision logic and potential discrimination, AI model risk through complex and adaptive algorithms, outsourcing and third-party risk through reliance on external models and cloud infrastructure, and reputational risk where adverse outcomes become visible to customers, regulators or the public.

Recent developments in agentic AI further amplify these dynamics. Rather than supporting isolated decision points, agentic systems can autonomously orchestrate and execute end-to-end processes, interacting with multiple models, data sources and systems. While the underlying risk management principles remain unchanged, this evolution shifts the focus from individual model governance to end-to-end process oversight, where risks arise from interactions between components, non-deterministic behaviour and dynamic decision pathways. For generative AI, these dynamics include the risk of plausible but inaccurate outputs, often described as hallucination, which can create customer, conduct, operational and reputational risk if not sufficiently controlled. This increases complexity in validation, monitoring and accountability, requiring insurers to assess not only individual models but also the integrity, control and outcomes of the full process chain.

At the same time, AI adoption introduces increasing dependency and sovereignty risks. Insurers are becoming structurally reliant on a limited number of global technology providers, including cloud platforms and large language model developers. These dependencies create concentration risks, exposure to vendor-specific failures, and potential constraints arising from geopolitical developments, regulation or access restrictions. The increasing convergence of these dependencies across core business processes effectively creates implicit single points of failure from a business continuity perspective. In particular, reliance on external foundation models may raise questions around data usage, costs, model transparency, portability, and the ability to ensure continuity of critical business processes.

The EU AI Act, in combination with Solvency II, DORA, GDPR, recent supervisory opinions and other frameworks, sets demanding expectations for the governance of AI in insurance. Even where systems are not formally classified as high risk, material AI use cases must in practice meet high standards of data quality, documentation, transparency, human oversight and accountability. These requirements must be embedded into existing risk management frameworks rather than treated as an isolated innovation topic, with increased emphasis on regulatory compliance, internal risk classification, end-to-end governance, third-party oversight and operational resilience in an (agentic) AI-driven environment.

This paper provides a practical overview of emerging AI governance expectations for insurers. It first outlines the regulatory landscape, the EU AI Act and selected international frameworks, before describing the key components of an AI governance framework, the role of the three lines model, Board oversight and risk appetite considerations. The paper concludes with the implementation challenges and a set of priority focus areas for CROs and risk functions.



2. The EU AI Act: what it requires

2.1 Structure and scope of the EU AI Act

The EU AI Act (Regulation (EU) 2024/1689) establishes a harmonised legal framework for AI systems placed on the EU market or used in the EU, aiming to ensure a high level of protection of health, safety and fundamental rights, while supporting innovation and the functioning of the internal market. The EU AI Act is a product-safety regulation, which means that product related concepts such as market placement and putting into service, conformity assessment, technical documentation and post-market monitoring are part of AI governance.

The framework is regarded as risk-based, with four categories of risk and corresponding regulatory requirements arising from the categorisation.

- **Unacceptable risk** → prohibited (Article 5)
- **High-risk systems** → subject to mandatory requirements on governance, documentation, risk management, human oversight and conformity requirements (Title III)
- **Limited risk** → transparency obligations (mainly Article 50)
- **Minimal risk** → largely unregulated under the EU AI Act

For AI systems that function only as support for high-risk purposes, Article 6(3) of the EU AI Act provides specific exceptions. However, these exceptions do not apply where the AI system profiles natural persons.

For insurers, the EU AI Act is particularly relevant because AI is increasingly used across the insurance value chain: underwriting, pricing, claims handling, fraud detection, customer segmentation, distribution, customer communications and operational automation. EIOPA has noted that AI is becoming increasingly important across pricing, underwriting, claims management and fraud detection, and that insurance-sector legislation already provides a governance and risk-management foundation for the use of such tools. The AI Act makes specific references to existing financial-sector regulatory requirements where providers are financial institutions. For example, in article 17(4) of the EU AI Act regarding a quality management system and article 19 (2) of the EU AI Act regarding automatically generated logs. On supervision, the EU AI Act stipulates in article 74 (6) that the relevant market surveillance authority

for the AI Act shall be the relevant national authority responsible for the financial supervision of those institutions. This means that EU AI Act compliance needs to be part of existing financial-sector governance structures or can rely on existing financial-sector controls when these meet the EU AI Act's specific requirements.

The EU AI Act has extra-territorial reach. It can apply not only to EU-based providers and deployers, but also to non-EU providers whose AI systems are placed on the EU market or whose outputs are used in the EU. This is important for international insurance groups that rely on centrally developed AI models, third-party AI vendors or global platforms deployed across European entities.

An AI system is defined according as a machine-based system designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments.

2.2 What counts as high risk for insurers?

First, the EU AI Act categorises certain AI systems as high risk where they are intended to be used for risk assessment and pricing in relation to natural persons in the case of life and health insurance.

This means that models used to determine eligibility, underwriting outcomes, premium levels or other material customer outcomes in life and health insurance may fall within the high-risk regime.

Examples may include:

- AI-supported underwriting engines for life insurance
- Health-risk assessment models
- Automated or semi-automated premium-setting tools for health insurance
- AI systems that materially influence whether a customer receives coverage, on what terms or at what price
- But potentially also, an AI system used for candidate selection

Not all insurance AI use cases will be classified as high risk. Where AI is used for internal process efficiency, administrative workflows, generic customer-service support or operational analytics,

this generally falls outside the high-risk category, although it may still be subject to risk oversight and compliance requirements, for example transparency, data protection, outsourcing, operational resilience or sectoral governance expectations.

Classification is also rarely straightforward in practice. Insurers frequently encounter borderline cases – for example, AI-assisted underwriting tools that inform but do not formally decide, or pricing engines that combine actuarial and machine learning components. A defensible classification requires documented rationale. Besides this, EIOPA's AI governance opinion deliberately excludes AI systems classified as high risk or prohibited under the EU AI Act, focusing instead on how existing Solvency II, IDD and DORA principles apply to other AI uses in insurance.

A crucial point for insurers is that high-risk classification under the EU AI Act depends primarily on the intended purpose of the AI system. The technical model can also be relevant in determining whether a tool qualifies as an AI system. It is also noteworthy that the European Commission recently started a consultation on guidelines for defining which systems should be considered high-risk under the EU AI Act. Contrary to the EU legislator's intent, these draft guidelines appear to extend the applicability of the EU AI Act's high-risk obligations for insurers, as the European Commission proposes a broad interpretation of systems that contribute to the relevant use cases. The above-mentioned exceptions for high-risk use cases should also be read narrowly. The concepts of "risk assessment" and "pricing" seem to be interpreted beyond systems that directly determine individual access to insurance or individual premium outcomes, potentially capturing broader risk profiling, group pricing and lifecycle-support tools.

2.3 Provider vs deployer: why the distinction matters

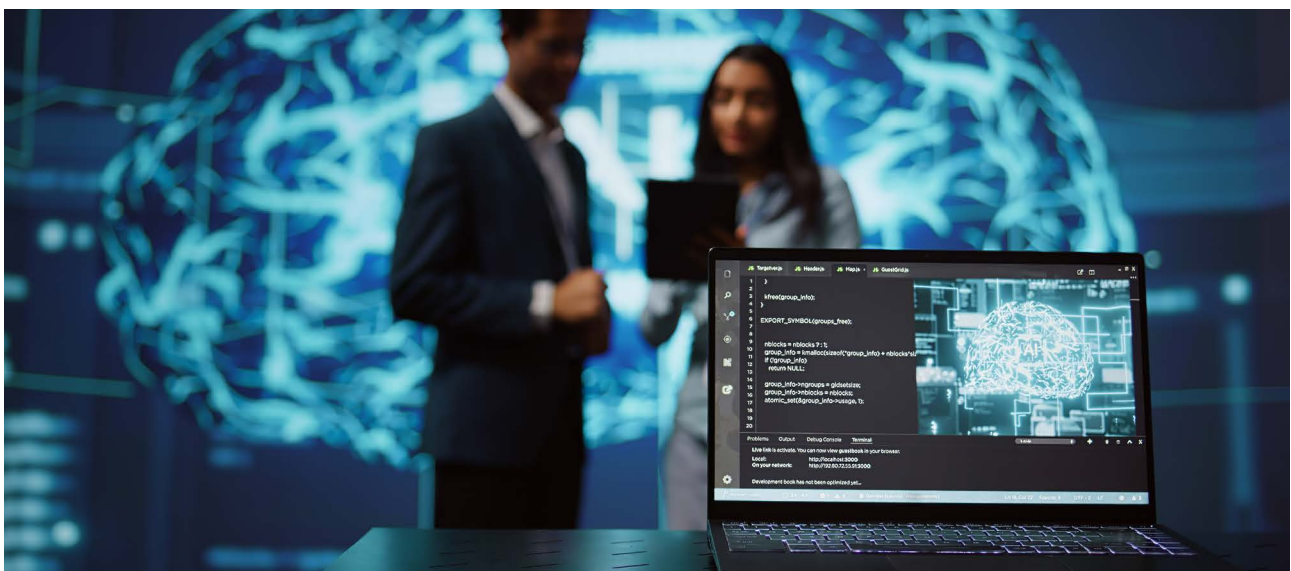
As stated before, coming from the product regulation nature of the EU AI Act, the latter assigns obligations to different actors in the AI value chain. For insurers, the most important distinction is between providers and deployers.

A provider is the natural or legal person that develops, or has developed, an AI system and places it on the market or puts it into service under its own name or trademark. A deployer is the natural or legal person that uses an AI system under its authority during a professional activity.

This distinction matters because insurers may play different roles depending on the operating model:

- An insurer that buys a third-party AI underwriting solution and uses it in its business will typically be a deployer.
- An insurer that develops its own AI pricing model centrally and puts it into service across the group will most likely qualify as a provider.
- An insurer that materially modifies a third-party AI system or changes its intended purpose may become subject to provider obligations.
- A group function that develops a model for local insurance entities may need to assess whether it acts as provider, internal service provider, or part of a wider deployer arrangement.

Providers face full lifecycle requirements, particularly for high-risk systems, while deployers focus on appropriate use, oversight, and monitoring. Insurers may perform both roles depending on their operating model. Misclassification can lead to compliance gaps and governance weaknesses.



In practice we expect insurers to be mainly deployers, but exceptions will exist. For insurance groups operating across multiple European markets, this allocation of roles should be resolved at group level and documented clearly. A central AI hub that develops AI systems for local entities may in practice assume provider-like obligations.

2.4 What high-risk classification requires

High-risk AI systems are subject to a demanding compliance regime. The core requirements include risk management, data governance, technical documentation, logging, transparency, human oversight, accuracy, robustness and cybersecurity.

For insurers, these requirements should not be seen as a separate compliance silo. They overlap strongly with existing expectations for lifecycle management, data governance, operational risk management, outsourcing, actuarial governance, product oversight and customer fairness. EIOPA has emphasised that existing sectoral legislation – including Solvency II, IDD and DORA – already establishes broad, technology-neutral governance and risk-management principles relevant to AI. As stated above, the EU AI Act itself acknowledges that existing financial regulation and measures could, where appropriate, support controls to demonstrate compliance with the AI Act.

Risk management system

Providers of high-risk AI systems must establish, implement, document and maintain a risk management system. This should identify and analyse known and reasonably foreseeable risks, estimate and evaluate risks that may emerge when the system is used as intended or under reasonably foreseeable misuse, and adopt appropriate mitigation measures.

For insurers, this requirement should be embedded into enterprise risk management and life cycle risk governance. A practical insurance implementation would normally include:

- AI use-case risk assessment;
- Own risk classification of systems by customer, financial, operational and regulatory impact;
- Pre-implementation validation;
- Assessment of bias, fairness and explainability;
- Ongoing performance monitoring;
- Trigger-based review after model drift, data changes, incidents or regulatory developments;
- Reporting and escalation routes to risk, compliance, actuarial and senior management bodies.

Existing practices, measures and controls used by insurers to comply with financial regulation could, where appropriate, also be used to comply with certain obligations under the EU AI Act.

Data governance

High-risk AI systems must be supported by appropriate data governance and data management practices. This is especially important where models are trained, validated or operated on personal, behavioural, health, biometric, claims or underwriting data.

For insurers, data governance is not only a technical issue. It is central to customer fairness and non-discrimination. Poor-quality, biased or unrepresentative data can result in unfair pricing, inappropriate exclusions, wrongful claims decisions or indirect discrimination through proxy variables. EIOPA's AI governance opinion highlights data governance, record-keeping, fairness, explainability, human oversight and cybersecurity as relevant supervisory expectations under existing insurance-sector legislation.

Technical documentation

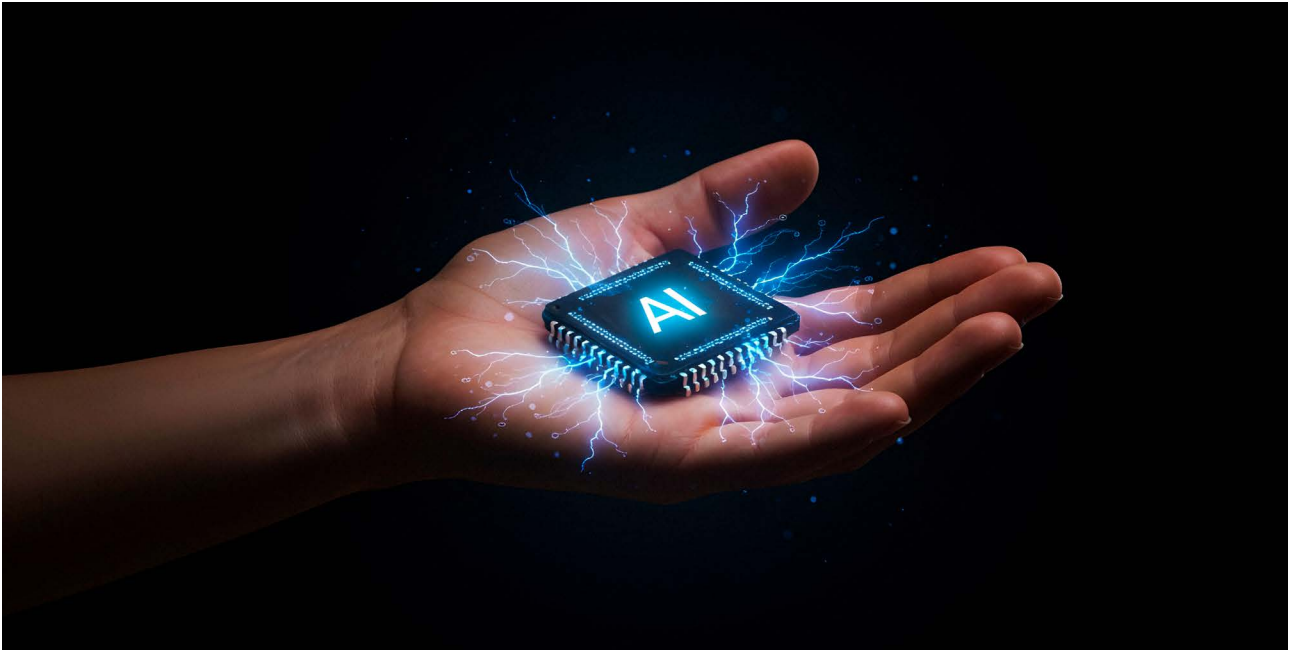
Providers of high-risk AI systems must maintain technical documentation demonstrating compliance with the Act. This documentation should explain the system's intended purpose, design, development, data, performance, risk controls, validation, limitations and oversight arrangements.

For insurers, technical documentation should be designed to serve multiple audiences: model developers, risk management, actuarial teams, compliance, internal audit, senior management and supervisors. A good documentation standard should explain not only how the model works, but also why it is appropriate for the intended insurance use case.

Logging and audit trails

High-risk systems must enable automatic logging of events throughout their lifecycle, to support traceability, monitoring, post-market surveillance and incident investigation.

For insurers, logging is essential where AI outputs influence customer outcomes. Firms should be able to reconstruct which model version was used, what input data was applied, what output was produced, whether a human reviewed the result and what final decision was taken. This is particularly important for underwriting, claims, complaint handling and pricing governance. Article 19 (2) of the EU AI Act specifically addresses the applicability of specific requirements coming from financial regulation.



Transparency and human oversight

High-risk AI systems must be sufficiently transparent to enable deployers to understand and use them appropriately. Human oversight must be designed to prevent or minimise risks to health, safety and fundamental rights.

For insurers, the key question is not whether a human is “in the loop” in a formal sense, but whether human oversight is effective. Effective oversight requires that employees understand the model’s purpose, limitations, confidence levels, escalation triggers and override procedures. Human review should be meaningful, documented and proportionate to the impact of the decision.

In addition, separate transparency obligations apply to certain AI systems beyond the high-risk category. For example, AI systems designed to interact directly with individuals may need to inform users that they are interacting with AI. Article 50 transparency obligations become particularly relevant for chatbots, virtual assistants and AI-generated content.

2.5 Fundamental rights impact assessment

Deployers of certain high-risk AI systems are required to perform a Fundamental Rights Impact Assessment, or FRIA. The purpose is to assess the potential impact of the AI system on fundamental rights before deployment, including risks relating to discrimination, privacy, access to essential services, consumer protection and vulnerable groups, i.e. decisions that materially affect individuals’ rights, opportunities, and life conditions.

For insurers acting as deployers, the FRIA applies only in narrowly defined circumstances, namely where insurers deploy high-risk AI systems for risk assessment and pricing in life and health insurance, as these decisions are considered to have a comparable societal impact to creditworthiness assessments. In these cases, the assessment should consider:

- Which individuals or groups may be affected;
- Whether vulnerable customers may be disproportionately impacted;
- Whether protected characteristics or proxy variables are used;
- Whether outputs could create unfair exclusion or price discrimination;
- Whether human oversight and complaints mechanisms are adequate;
- Whether customers can understand and challenge significant decisions.

The FRIA should not be treated as a one-off legal document. It should link to the insurer’s product approval process, AI validation, conduct risk framework, data protection impact assessment, outsourcing assessment and operational risk management process.

2.6 Conformity assessment and registration

As product regulation, the EU AI Act obliges providers to carry out a conformity assessment before a high-risk AI system is placed on the market or put into service, to ensure compliance. Certain (related to Annex III) high-risk AI systems must also be registered in the EU database.

For insurers, this creates a need for a formal go-live control gate. High-risk AI systems should not move into production unless the firm has completed classification, role mapping, risk assessment, technical documentation, data governance checks, human oversight design, validation, FRIA where required, and conformity or registration steps where applicable.

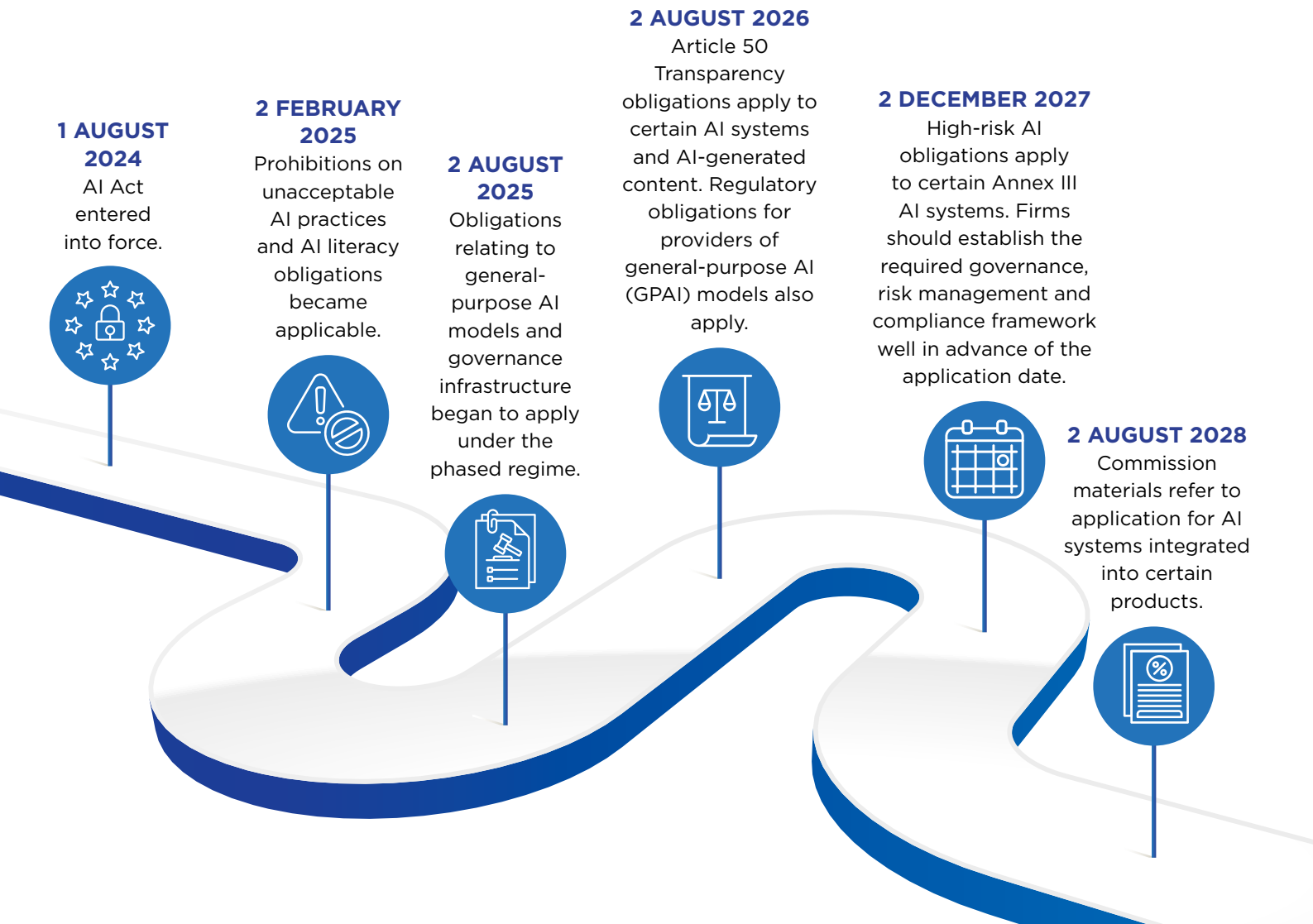
Where third-party vendors provide AI systems, insurers should ensure that contracts provide access to the information needed to satisfy deployer obligations. Vendor due diligence should cover not only performance and security, but also EU AI Act readiness, data lineage, model governance, testing evidence, incident reporting and audit rights.

2.7 Key compliance timeline

The EU AI Act entered into force in 2024, with obligations applying in phases. Key requirements relating to prohibited AI practices, AI literacy, general-purpose AI models and transparency obligations apply before the implementation of the high-risk AI requirements. Obligations for AI systems classified as high risk under Annex III apply from 2 December 2027, while obligations for AI systems integrated into certain regulated products apply from 2 August 2028 (see figure 1).

Regardless of exact phasing, the required control framework is substantial and should be built well before the final application date. Given the speed of development around AI systems and adoption within the insurance industry, a sound risk and compliance management framework should be established in line with business demand and not rely on regulatory requirements.

Figure 1: Key milestones for insurers include:



2.8 Penalties

The EU AI Act contains significant administrative fine powers. The most serious infringements, including breaches of prohibited AI practices, may attract fines of up to €35 million or 7% of total worldwide annual turnover, whichever is higher. Fines vary by infringement type (prohibited use, non-compliance, misinformation).

For insurers, enforcement risk is only one part of the exposure. AI failures may also create conduct risk, litigation risk, reputational damage, supervisory intervention, customer remediation, data protection findings, operational incidents and board accountability concerns.

2.9 Interaction with Solvency II and EIOPA expectations

The EU AI Act does not replace Solvency II, IDD, DORA, GDPR or national insurance supervision. Instead, it overlays these frameworks. For insurers, this means that AI governance should be integrated into the existing system of governance rather than operated as a new standalone framework.

EIOPA's 2025 Opinion on AI governance and risk management clarifies how existing insurance-sector legislation should be interpreted in the context of AI. It is addressed to national supervisors and covers insurance undertakings and intermediaries using AI in the insurance value chain. It does not create new legal requirements and excludes prohibited and high-risk AI systems under the AI Act to avoid regulatory overlap.

The relevant Solvency II governance principles include effective risk management, internal control, fit and proper governance, actuarial oversight, outsourcing control and prudent management. EIOPA's Opinion specifically links AI governance to existing expectations around data governance, record-keeping, fairness, cyber security, explainability and human oversight.

For CROs, the implication is that the EU AI Act should be embedded into:

- The risk management system;
- The model risk management framework;
- Product approval and review processes;
- Conduct and customer fairness frameworks;
- Outsourcing and third-party risk management;
- Data governance and privacy controls;
- Operational and technical resilience;
- Board and senior management reporting.

For CROs, the primary objective should not be regulatory compliance in isolation but ensuring that AI adoption remains consistent with the organisation's risk appetite, operational resilience objectives and long-term customer outcomes, and enabling AI adoption throughout the company.



3. International frameworks – global approaches

3.1 Why international frameworks matter for insurers

Global insurers operate across jurisdictions with different regulatory philosophies. The EU AI Act is legally binding and risk based. The United States relies heavily on voluntary standards and agency-specific supervisory approaches. Singapore has developed financial-sector-specific responsible AI principles and practical toolkits. The United Kingdom has adopted a principles-based and outcomes-focused approach through existing regulators rather than a single comprehensive AI statute. International standards such as ISO/IEC 42001 or NIST AI risk management framework provide management structures that can be used across jurisdictions.

For international insurance groups, the issue is therefore not which framework to choose, but how to build an AI governance framework that can map to several regimes at once. The most efficient approach is to establish a global AI risk management baseline, then add jurisdiction-specific requirements where needed.

3.2 NIST AI risk management framework

The NIST AI Risk Management Framework, or AI RMF 1.0, was released in January 2023. It is voluntary, non-sector-specific and use-case agnostic. Its purpose is to help organisations designing, developing, deploying or using AI systems manage AI risks and promote trustworthy and responsible AI.

The NIST AI RMF is structured around four core functions:

1. Govern

Establish organisational policies, accountability and risk management structures.

2. Map

Understand the context, intended purpose, stakeholders and risks of the AI system.

3. Measure

Assess, analyse and track AI risks and trustworthiness characteristics.

4. Manage

Prioritise, respond to and monitor risks throughout the AI lifecycle

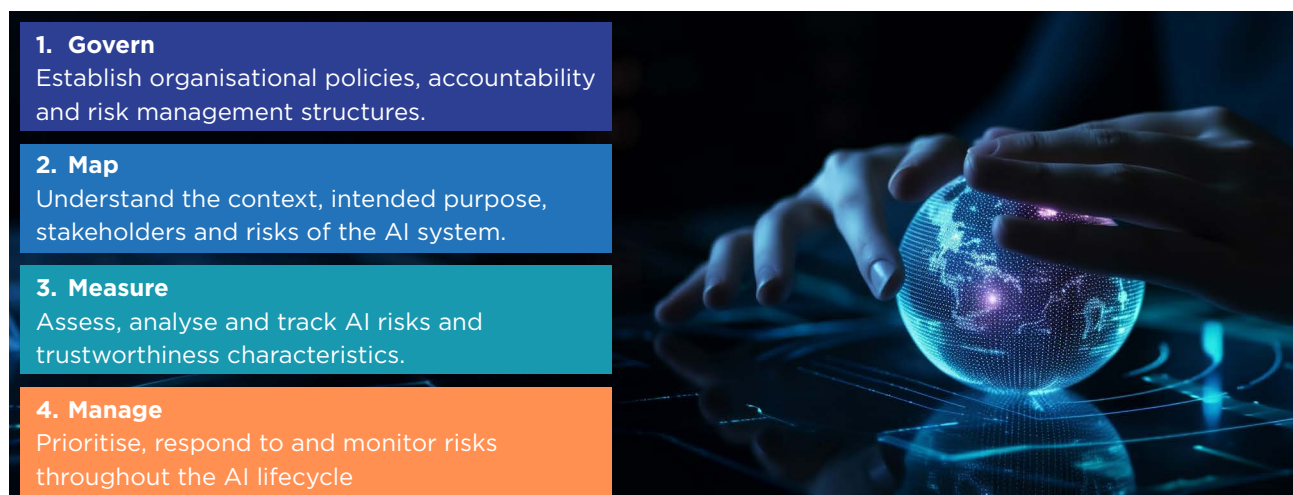
For insurers, NIST is useful because it provides a practical risk management language that can be applied across all AI use cases, including those outside the scope of the EU AI Act's high-risk provisions. It is particularly valuable for model inventory, risk tiering, control design, testing, monitoring and governance reporting.

However, NIST is not a legal compliance framework. It does not by itself satisfy the EU AI Act's conformity assessment, registration, FRIA or provider/deployer obligations. Its value lies in operationalising trustworthy AI in a structured and flexible way.

3.3 ISO/IEC 42001 – AI management system

ISO/IEC 42001 is the first international management-system standard for AI. It provides a certifiable framework for establishing, implementing, maintaining and continually improving an AI management system. It is designed to help organisations manage AI responsibly through policies, roles, risk management, operational controls, performance evaluation and continual improvement.

Compared with NIST, ISO/IEC 42001 is more formalised as a management-system standard. It aligns with the structure familiar from other ISO standards, such as ISO 27001 for information security and ISO 9001 for quality management. This makes it attractive for firms that want enterprise-wide assurance and potentially external certification.





For insurers, ISO/IEC 42001 can serve as a backbone for a global AI governance framework. It can help define:

- AI policy and governance structure;
- Roles and responsibilities;
- AI system inventory and lifecycle controls;
- Risk assessment and treatment;
- Supplier and third-party controls;
- Monitoring and continuous improvement;
- Internal audit and management review.

ISO/IEC 42001 can support EU AI Act readiness, but it does not replace legal compliance. Certification may provide useful evidence of governance maturity, but firms will still need to meet the specific obligations of the AI Act, Solvency II, IDD, GDPR, DORA and local supervisory expectations.

3.4 UK FCA/PRA – principles-based and outcomes-focused approach

The UK has not adopted an EU-style horizontal AI Act. Instead, the FCA and PRA have pursued a principles-based and outcomes-focused approach, relying largely on existing regulatory frameworks. The FCA has stated that it does not currently plan to introduce additional regulations for AI and will rely on existing frameworks that mitigate many AI-related risks.

The Bank of England, PRA and FCA published Discussion Paper DP5/22 on Artificial Intelligence and machine learning in 2022, followed by Feedback Statement FS2/23 in 2023. Respondents generally favoured principles-based or risk-based approaches rather than a rigid regulatory definition of AI and emphasised the importance of regulatory coordination and alignment.

The FCA's current approach links AI oversight to existing obligations such as:

- Consumer Duty;
- Senior Managers and Certification Regime;
- Systems and controls;
- Governance and accountability;
- Fair customer outcomes;
- Operational resilience;
- Model and data governance.

For insurers operating in the UK, this means AI will be assessed through the lens of outcomes, governance and accountability rather than through a standalone AI compliance checklist. The central question will be whether the insurer can demonstrate that AI is being used safely, responsibly and in a way that delivers fair customer and market outcomes.

3.5 Other International Frameworks

While the EU AI Act, NIST AI RMF, ISO/IEC 42001 and the UK approach are among the most widely referenced frameworks, insurers operating globally may also need to consider other more mature jurisdiction-specific supervisory expectations.

Singapore – MAS FEAT Principles and Veritas

The Monetary Authority of Singapore (MAS) has adopted a sector-specific approach through its FEAT principles, which focus on Fairness, Ethics, Accountability and Transparency in the use of AI and data analytics within financial services. To support implementation, MAS and industry participants developed the Veritas framework and toolkit, which provides practical methodologies for assessing fairness, explainability, transparency and governance of AI systems. The Singapore approach is largely principles-based and supervisory in nature, focusing on responsible outcomes, customer trust and effective governance rather than prescriptive legal requirements. Nevertheless, it is widely regarded as one of the most mature operational frameworks for AI governance in financial services.

Hong Kong – HKIA Supervisory Approach

Hong Kong has no dedicated AI regulation for insurance. The Insurance Authority (IA) has indicated, in its May 2023 Conduct in Focus commentary, that AI deployment is governed by the existing framework on a technology-neutral basis, with accountability resting with the insurer or intermediary, not the technology.

Several Guidelines operate in combination. The Guideline on Enterprise Risk Management (GL21) requires insurers to identify and manage all material risks, which the IA has indicated extends to AI: expectations that AI use cases be risk-assessed in context, tested before deployment, accompanied by disclosure of system limitations, and subject to ongoing monitoring and contingency planning. Under the Guideline on Cybersecurity (GL20), AI/ML use is an inherent-risk criterion in the Cyber Resilience Assessment Framework. The Guideline on Outsourcing (GL14) applies to externally hosted AI, and the Guideline on the Use of Internet for Insurance Activities (GL8) to AI-enabled online distribution. Board accountability follows from the Guideline on Corporate Governance (GL10), and the Insurance Ordinance conduct requirements apply fully to AI-mediated customer interactions. These AI-specific expectations reflect the IA's published supervisory thinking rather than express provisions of the Guidelines.

Dedicated IA supervisory guidance has been announced but not yet published.

3.6 Common denominators across frameworks

For multinational insurers, these frameworks demonstrate a broader global regulatory trend. While differing in legal form and supervisory intensity, they increasingly converge on common expectations around governance and risk management. The result is a growing international baseline for responsible AI that extends beyond compliance with any single jurisdiction.

First, they all require clear governance and accountability. Whether through the EU AI Act's provider/deployer obligations, NIST's Govern function, ISO/IEC 42001's management-system requirements or the UK's SM&CR-based approach, firms must be able to demonstrate ownership, oversight and escalation.

Second, they require risk-based classification and proportionality. High-risk AI use cases require stronger controls than low-risk applications. For insurers, this means underwriting, pricing, claims denial, fraud scoring and customer eligibility tools should receive more rigorous oversight than low-impact internal productivity tools.

Third, all frameworks emphasise data quality, fairness and bias control. This is especially important in insurance, where historical data can contain structural biases and where proxy variables may indirectly affect protected or vulnerable groups. EIOPA and the FCA/PRA all highlight fairness, data governance and customer outcomes as central supervisory concerns.

Fourth, they require transparency, explainability and traceability. This does not always mean full technical explainability of every model parameter, but it does mean that organisations must understand the model's purpose, limitations, outputs and decision impact. For material customer decisions, firms should be able to explain the rationale in a way that is meaningful to internal control functions, supervisors and, where appropriate, customers.

Lastly, all frameworks expect monitoring and lifecycle management. AI governance does not end at model approval. Firms need continued performance monitoring, drift detection, incident management, periodic validation, change control and retirement procedures.

4. AI governance framework – lifecycle

Good AI governance should be embedded through sound risk management principles rather than being focused primarily on specific technologies or regulations. As AI capabilities continue to evolve rapidly, principles around accountability, proportionality, transparency, human oversight, validation and ongoing monitoring provide a more durable foundation for governance than technology-specific controls.

AI governance must operate across the full lifecycle of an AI system. Regulatory expectations require ongoing oversight from design through to decommissioning.

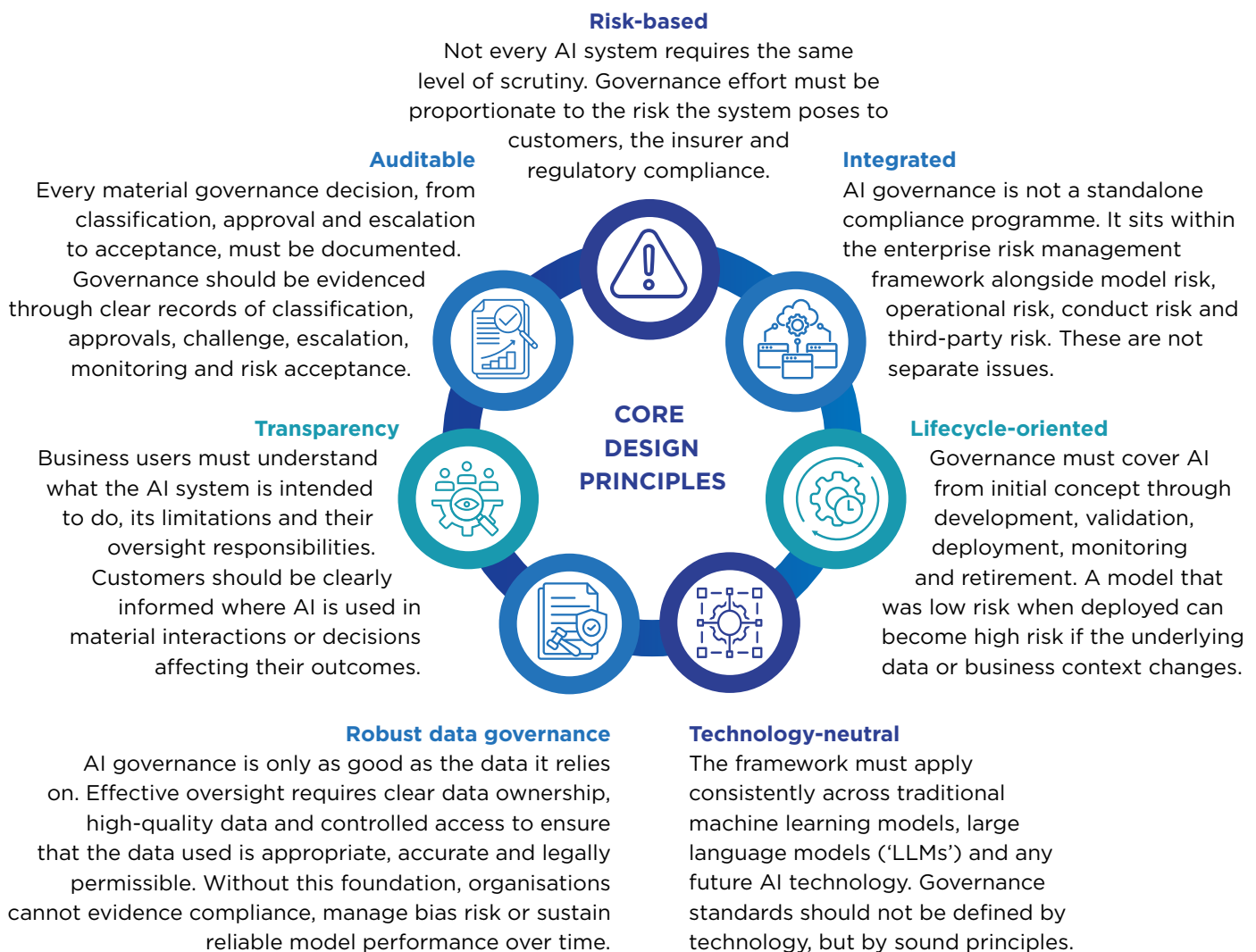
This reflects the dynamic nature of AI risk. Systems can change over time due to data, performance or use, altering their risk profile and regulatory obligations.

An effective framework must therefore ensure periodic risk assessment, control application, monitoring and auditable decision making across each stage of the lifecycle, aligned to the firm's enterprise risk framework.

4.1 Core design principles

An effective AI risk and compliance management governance framework should be built on seven principles (see figure 2).

Figure 2: Seven core design principles



4.2 Core framework components

AI inventory

The Inventory is the foundation and requires operational ownership, oversight and audit review. It should serve as your 'single source of truth'.

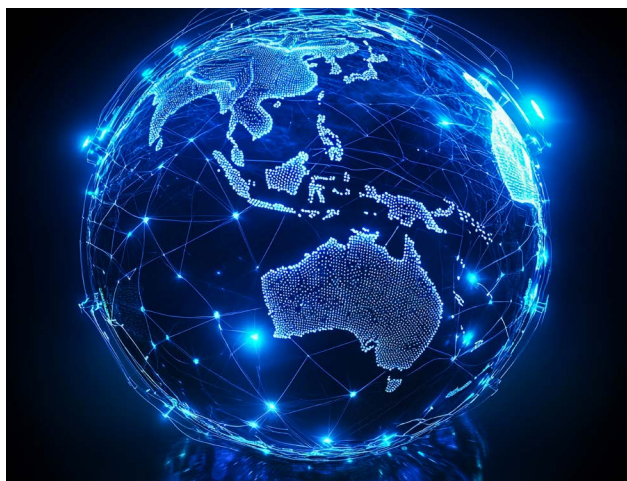
In practice, building the inventory is the hardest first step. AI systems are often not centrally registered, and ownership may be fragmented across business units. A pragmatic approach is to start broad – capturing all automated decision-support tools – and then apply (internal) classification criteria to determine which require deeper governance and additional assessments. A tool like a centralised metadata repository or AI system inventory platform can significantly accelerate this process and provide the audit trail regulators will expect.

Classification decisions – high risk, limited risk, minimal risk – must be documented and justified with reference to EU AI Act criteria and internal risk and compliance assessments.

Risk Assessment Methodology

Determining the level of governance and oversight on an AI use case requires an internal risk classification.

The EU AI Act's risk tiers were designed to protect fundamental rights and safety across all sectors – they were not calibrated to the specific risk profile of a regulated financial institution. An AI system classified as minimal or limited risk under the EU AI Act may nevertheless be material from a prudential, conduct or operational risk perspective. A customer-facing pricing tool that falls below the high-risk threshold may still carry significant conduct risk. An internal fraud scoring model may be limited or even minimal risk under the EU AI Act but operationally critical and subject to various risk, compliance and operational resilience expectations.



An LLM used for claims correspondence may only trigger transparency obligations but represent far greater reputational and customer outcome risk than its regulatory classification suggests.

Insurers therefore need their own classification layer – an internal risk assessment that determines the intensity of risk and compliance involvement based on criteria that reflect the firm's own risk profile. Factors to consider may include:

- Customer impact;
- Financial materiality;
- Operational criticality;
- Data sensitivity and quality;
- Model complexity;
- Reversibility of decisions.

This internal tiering drives the practical governance response – the depth of validation required, the frequency of monitoring, the level of human oversight, the escalation thresholds and the degree of second-line involvement. In practice, regulatory classification determines legal obligations, while internal risk classification determines governance intensity. Both are necessary and complementary.

Each AI system in the inventory requires a structured assessment. For Regulatory and risk management purposes, this has two components:

1. **Regulatory classification** – is the system high-risk as defined (for example by Annex III of the EU AI Act)? Potential follow-up: Fundamental Rights Impact Assessment ('FRIA') where required.
2. **Internal risk and impact assessment** – what is the potential harm to customers and the firm if the system e.g. fails or produces biased output as mentioned above?

The EU AI Act classification is not a substitute for an insurer's own risk assessment. The own risk assessment should drive the risk management response and governance intensity, while regulatory classification determines applicable legal obligations.

AI model validation & testing

AI model validation and testing should be completed before deployment and repeated periodically or following material model, data or business changes. For AI systems that learn, adapt or operate autonomously, validation should evolve from a point-in-time exercise to a process of continuous monitoring and periodic reassessment, supported wherever possible by automated controls and testing. The depth of validation should be

proportionate to the risk classification and potential customer, financial, operational and regulatory impact.

For material or high-risk systems, validation should be independent of AI model development and should assess performance, robustness, explainability, bias, data suitability, limitations and operational use. Validation standards should also recognise that LLMs and agentic AI may produce non-deterministic outputs and require specific testing techniques.

This implies that AI model validation capabilities are required across the 3 lines model including the risk management function for independent AI model validation for material AI systems or end-to-end process reviews.

Control framework

Controls must map to the risk classification.

For (internal) high-risk AI systems, mandatory controls include:

- Pre-deployment conformity assessment;
- Technical documentation, including system purpose, capabilities, limitations, data sources, outputs and intended use;
- Testing and validation evidence covering performance, robustness, bias, fairness, explainability and, for LLMs, hallucination risk, with documented results, remediation and acceptance;
- Human oversight mechanisms ('Human in the Loop' or 'HITL') with documented override capability;
- Logging with defined retention periods;
- FRIA (where required);
- Pre-deployment conformity assessment and approval to deploy; post deployment monitoring expectation with defined Key Risk Indicators ('KRIs') and escalation thresholds.

For lower risk systems a lighter control set should apply but documentation and review are still required.

Governance processes

The framework must define:

- Approvals for new AI deployments
- Change management for material model changes
- Escalation routes for incidents and when KRIs breach thresholds
- Periodic review cycles for existing systems (including risk classification)

These processes must have clear ownership, documented decision-making rights and audit trails. These processes must also align with existing governance forums and escalation routes within the enterprise risk framework.

Third-party AI vendor management

A significant portion of AI systems used by insurers are built by external vendors. Pricing platforms, fraud detection systems and claims triage products are routinely procured from specialist technology firms. The EU AI Act does not reduce the deployer's obligations because the systems were built externally and the insurer retains responsibility for what the system does.

There are two key AI governance aspects to address with vendors, namely structured Due Diligence and contractual obligations.

Before deploying any third-party AI system with a high-risk classification, an insurer must undertake structured vendor Due Diligence to establish:

- Is the vendor acting as a Provider under the EU AI Act? If so, have they conducted a conformity assessment, produced technical documentation and registered the system.
- Can the vendor provide the technical documentation the firm needs to fulfil its deployer obligations? This includes data governance records, validation reports and logging specifications.
- Does the Vendor's model meet the firm's data and governance fairness standards? What training data was used? Has bias testing been conducted and can documented evidence be shared?
- What are the vendor's incident notification obligations? If the model produces an error, how quickly will the firm be notified and what remediation support will the vendor provide?
- What is the exit strategy? If the vendor relationship ends, can the firm maintain continuity of the AI supported business process?

There are key contractual provisions which are required to reflect EU AI Act obligations. These should include access to technical documentation; audit rights; incident notification timelines; obligations to maintain regulatory compliance as regulations evolve; data governance and protection standards; and termination and exit support provisions. Insurers must also decide on their risk appetite regarding model training especially when considering customer data.

5. The Three Lines Model

The Three Lines Model remains a robust foundation for AI governance but must evolve in response to increasing automation and the use of advanced AI systems. This includes stronger embedding of controls within processes, clearer accountability across the AI lifecycle, and enhanced collaboration between the three lines. Particular attention is required to maintain first-line ownership, avoid blurred responsibilities, and ensure that governance keeps pace with the increasing complexity of AI-driven end-to-end processes.

5.1 Responsibilities

The Three Lines Model operates under the ultimate accountability of the Board and executive management, which are responsible for ensuring that an effective system of governance is in place, including clear allocation of responsibilities across the three lines.

The Three Lines model applies to AI risk and compliance management in the same way that it applies to any other risk type. The delineation of responsibilities must be explicit but the approach to framework ownership may be insurer dependent.

Framework ownership may be centralised within the second line or embedded within the first-line operating model. Where governance is owned by the first line, insurers should evidence clear second-line independence through approval, challenge and ongoing oversight activities to avoid conflicts between governance, advice, control design and execution.

Effective AI governance requires coordination with wider control and support functions. Compliance, Legal, Procurement and Data Protection functions play a critical role, particularly in regulatory interpretation, third-party oversight and data governance. Risk Management plays a central role in the governance of material AI use cases by supporting risk classification, defining risk assessment methodologies, providing independent challenge, overseeing validation activities and monitoring risk exposures across the AI lifecycle. It also supports management and the Board in understanding the aggregate AI risk profile and the effectiveness of associated controls.

These functions typically operate within the first and second line, but their specific responsibilities should be clearly defined within the respective company's framework.

The appropriate focus of framework ownership evolves with organisational maturity. In earlier stages of AI governance development, the second line could develop the set-up of the framework—establishing the inventory, driving classification, and building the control framework where first-line capability does not yet exist. As maturity increases, framework ownership migrates toward the first line – and in particular toward IT, data engineering or dedicated AI functions that have the technical depth to govern AI systems close to their source. At embedded maturity levels, the “1.5 line” often emerges as an effective structure: embedded governance specialists within the AI or technology function provide day-to-day governance leadership, while the second line focuses on independent oversight, challenge and escalation.

Line	Role in AI governance	Key activities
First Line (IT and Business, e.g. actuarial function, marketing, claims department)	Owns AI Systems and bears primary accountability for compliance and risk management	Build and maintain AI inventory; conduct risk classification; execute FRIA; maintain technical documentation; ensure human oversight in operations; manage AI vendors; produce and monitor KRIs; implement automated controls
Second Line (Risk, including information security, and Compliance)	Provides independent challenge and oversight; supports the Board	Challenge first line risk classification and risk assessments; conduct independent oversight of high-risk systems; monitor KRIs; escalate risks to Board
Third Line (Internal Audit)	Provides independent assurance over the AI Governance Framework	Audit inventory completeness; Test HITL mechanisms; review FRIA quality; assess third-party AI vendor governance; report to Audit Committee

The increasing use of AI, including large language models and agentic systems, requires an evolution of the Three Lines model. As decision-making becomes more automated and distributed across systems, risk management must move from point-in-time control to continuous, data-driven oversight embedded within processes including the adoption of AI itself. This increases the importance of clear accountability, strong first-line ownership and effective collaboration between the three lines, since compliance requires a combination of policy and directives, guidance and advice, technical controls and cultural awareness.

The increasing use of AI also creates opportunities for the second line. AI can support risk identification, monitoring, control testing, compliance reviews, emerging-risk detection and management reporting, allowing risk management functions to expand coverage, increase efficiency and focus more effort on areas requiring expert judgement and challenge.

5.2 Monitoring, KRIs and escalation

AI governance does not end at deployment. High-risk systems (under the AI Act or internal classification) must be monitored on an ongoing basis and throughout their lifecycle. KRIs should, where possible, be defined and monitored in real time to trigger review or escalation.

At minimum these metrics should cover:

Potential KRI categories could include (amongst others):

- Model performance and drift;
- Fairness and customer outcome indicators;
- Operational incidents and near misses;
- Human override rates;
- Complaints and challenge rates;
- Vendor performance;
- Data quality;
- Unauthorised or shadow AI use;
- Control exceptions;
- Foundation-model concentration index;
- Adversarial-resilience test results;
- AI-asset-inventory completeness;
- Time-to-detect model drift.

Model performance

How to detect degradation. Is the system's accuracy or predictive power declining? Is model drift present, whereby the relationship between the input data and the real-world changes? For example, a claims triage model may begin to wrongly classify a greater proportion of complex claims over time, leading to increased manual intervention and delayed settlement as underlying claim patterns change, and ultimately also to financial impact through gradual drift in the claims ratio.

Fairness

Is the system producing differential outcomes across protected characteristics? For example, higher premiums being produced for certain ethnic groups derived by proxy through postcodes or occupations even when ethnicity is not explicitly used as an input.

Incidents and near misses

As with other incidents, any AI incident — erroneous outputs, customer complaints due to an AI decision, or a data quality failure — must be logged, investigated and included in the governance process.

Third-party Performance

Where AI systems are provisioned by external vendors, are those vendors meeting agreed service and governance standards? Any changes in their ownership, financial health or regulatory status may affect the risk profile of the activity they support.

6. Board oversight and AI risk appetite

6.1 Board role

Boards and senior management are accountable for ensuring that AI risk is governed within the firm's overall system of governance. Boards do not need to understand every technical aspect of an AI model, but they do need to be able to approve and define the boundaries of the firm's AI risk appetite, receive regular, meaningful reporting on AI risk and oversight of AI related incidents and associated responses.

Given the EU AI Act's phased rollout, with transparency and broad compliance obligations applying from August 2026, the Board should also be receiving compliance readiness updates.

Board reporting should give directors a clear view of the firm's AI risk profile rather than a technical inventory. Useful MI should include the number and risk tier of AI systems, high-risk use cases, open control gaps, validation status, KRI breaches, incidents and near misses, vendor dependencies, fairness testing outcomes and progress against regulatory readiness plans.

6.2 Definition of appetite

Many insurers have yet to define a formal AI risk appetite and this presents a gap to remedy. A clearly articulated risk appetite is not only a control mechanism but also an enabler of innovation. It provides clarity on which AI use cases can be adopted at speed, which require additional governance and which fall outside the tolerance. Without a sufficiently enabling governance framework and risk appetite, firms may struggle to realise the benefits of AI while competitors move more aggressively.

Acceptable use cases

Boards should consider what AI applications are permissible, require additional approval or should be prohibited. Examples of unacceptable uses could include fully automated underwriting or marketing decisions with no meaningful human oversight; or black box models or third-party AI without demonstrable explainability or transparency.

Automation limits

Human oversight establishes a minimum standard, but the risk appetite should go further. Parameters, for example monetary or sensitivity, for decisions in which a human must always be the final decision maker should be explicit and monitored.

Third-party reliance

Boards should set explicit limits on concentration risk arising from reliance on third-party AI providers. Risk appetite should also define resilience expectations for AI-supported critical or important business processes, including continuity, substitutability, exit planning and vendor failure scenarios. Vendor due diligence should reflect these standards.

Bias, fairness and customer outcomes

Risk appetite should define tolerance for unfair or discriminatory customer outcomes, including use of proxy variables, vulnerable customer impacts and differential outcomes in pricing, underwriting or claims. Boards should expect pre-deployment and ongoing fairness testing, with clear remediation, escalation and documented risk acceptance where issues arise.



7. Challenges

Several practical challenges are likely to determine whether AI governance is effective in practice. These challenges range from strategic ownership and speed of adoption through to control design, validation and senior-level understanding. They are not reasons to delay adoption, but they do require clear accountability, proportionate controls and senior management attention.



Strategic ownership of AI rollout

Insurers are taking different approaches to AI ownership, with some placing responsibility with the COO and others with the CTO, CIO or data function. Each model brings different strengths and risks, and the chosen approach should make accountability clear across strategy, delivery, governance and customer outcomes.



Speed of innovation

AI use may expand faster than governance frameworks, risk taxonomies and control functions can adapt. Insurers need intake and triage processes that identify material use cases early, before models become embedded in critical processes.



Shadow AI

The accessibility of generative AI increases the risk that staff use tools outside approved channels. This creates data, confidentiality, record-keeping and conduct risks unless firms have clear policies, technical controls and monitoring. The emergence of agentic AI adds a further dimension to this challenge. AI agents may be created and deployed by end users, embedded within software platforms, or operate under existing user identities and permissions. This creates a challenge comparable to traditional End-User Computing (EUC), where tools initially developed for individual productivity can gradually become embedded within business processes and influence operational or customer outcomes. As AI agents become easier to create, share and reuse across teams, identification, inventory management, accountability and monitoring become significantly more difficult. Insurers should therefore ensure that AI governance frameworks are capable of identifying, tracking and monitoring AI agents throughout their lifecycle, including those that evolve from personal productivity tools into components of broader business processes.



Meaningful human oversight

The presence of a human reviewer does not, by itself, provide control. Staff need enough understanding, authority and time to challenge AI outputs, override where appropriate and escalate uncertain or adverse outcomes. This requires good understanding of the underlying AI systems.



Data quality, ownership and accountability

AI systems are only as reliable as the data on which they are trained and operated. AI failures originate not from the model itself, but from incomplete, inaccurate, outdated or poorly understood data.

For insurers, data often originates from multiple internal and external sources, making ownership and accountability challenging. Poor data governance can lead to biased outcomes, inaccurate decisions, model drift and difficulties demonstrating regulatory compliance. As AI adoption increases, effective data governance becomes a foundational prerequisite for effective AI governance rather than a separate discipline.

**Validation of complex AI**

Traditional model validation techniques may not be sufficient for LLMs and agentic AI, where outputs may be non-deterministic and behaviour may vary by context. Insurers need enhanced testing for robustness, explainability, fairness, prompt sensitivity, data leakage and unintended use.

**Model, vendor concentration and sovereignty dependencies**

The foundation-model market is concentrated among a small number of providers, and most insurers will access AI capability through a handful of these models and the cloud platforms that host them. This creates a third-party and concentration risk that is qualitatively different from traditional outsourcing: a common dependency on the same underlying models means that a model failure, degradation after an update, safety-driven withdrawal of a capability, or sudden change in pricing or availability could affect many insurers simultaneously. A model version change can silently alter behaviour in production, making outputs non-stationary in a way that classical model risk management did not have to address. This also sits squarely within the DORA framework for ICT third-party and concentration risk. Insurers should extend concentration-risk analysis explicitly to AI model and platform dependencies, seek contractual protections around change management, versioning, explainability and exit, and consider the operational-resilience question of what the business does if a critical model becomes unavailable or unacceptable to use.

Beyond concentration risk, insurers should also consider AI sovereignty risk. Dependence on a small number of non-domestic cloud, foundation-model and AI-service providers creates potential exposure to geopolitical developments, export restrictions, regulatory interventions, changes in market access or provider-specific policy decisions that may affect the availability, functionality or permitted use of AI capabilities. As AI becomes embedded in critical business processes, insurers should assess the resilience, portability and substitutability of key AI services and consider whether strategic dependencies align with their broader risk appetite and operational resilience objectives.

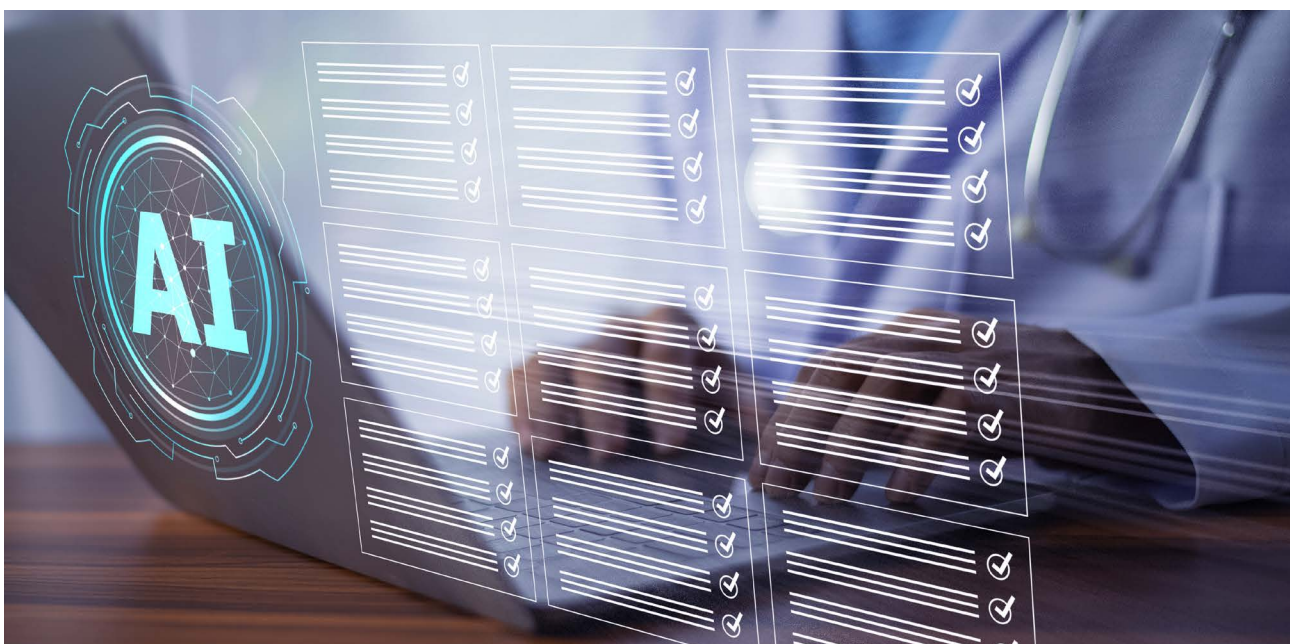
**Board and regulatory understanding**

Boards and supervisors do not need technical detail on every model, but they do need sufficient understanding to assess whether the firm's governance, risk appetite, controls and customer outcome monitoring are effective.

8. Focus areas for CROs

The following focus areas are prioritised by urgency. The first three are immediate priorities given the speed of AI adoption and the maturity uplift required. The remainder are important for building a sustainable AI governance capability.

1. **Complete and maintain an AI inventory**
Identify every AI system in production and categorise each against regulatory and internal classification. Document the classification rationale.
2. **Conduct a gap analysis against EU AI Act obligations (and other relevant regulation)**
Each high-risk system should be assessed against the mandatory requirements. Identify what is missing and build a clear remediation plan with owners and deadlines.
3. **Establish a formal approval process for high-risk use cases**
Ensure that material AI systems are identified, assessed and approved before deployment, with clear requirements for validation, documentation, human oversight, governance review and ongoing monitoring. Where third-party AI is used, obtain sufficient information to understand, oversee and challenge the operation of the system.
4. **Define your risk appetite**
Consider your firm's respective appetites for acceptable use cases, automation limits, human oversight expectations and third-party reliance thresholds.
5. **Integrate AI risk into your firm's enterprise risk management framework**
AI is not a standalone compliance matter and should be embedded within existing risk taxonomies and risk assessments.
6. **Strengthen AI vendor governance**
Review all AI vendor contracts and update Due Diligence questionnaires and processes, ensuring that technical documentation requirements are met to ensure compliance.
7. **Establish AI Model Validation**
Build up knowledge and have independent AI model validation teams assessing AI systems for processes that are critical or can have a material financial or reputational impact.
8. **Leverage AI within the Risk Management Function**
Explore how AI can support risk identification, monitoring, control testing, horizon scanning, incident analysis, regulatory change management and reporting. Risk functions should not only govern AI but also consider how AI can enhance the effectiveness and efficiency of risk management itself, while ensuring appropriate oversight and validation of these capabilities.



Appendix A

List of abbreviations

Abbreviation	Description
AI	Artificial Intelligence
AI Act / EU AI Act	European Union Artificial Intelligence Act (Regulation (EU) 2024/1689)
BCM	Business Continuity Management
CRO	Chief Risk Officer
DORA	Digital Operational Resilience Act
ERM	Enterprise Risk Management
EUC	End-User Computing; technology solutions developed by end-users outside traditional IT development processes
EIOPA	European Insurance and Occupational Pensions Authority
FCA	Financial Conduct Authority (United Kingdom)
FEAT	Fairness, Ethics, Accountability and Transparency principles issued by the Monetary Authority of Singapore
FRIA	Fundamental Rights Impact Assessment under the EU AI Act
GDPR	General Data Protection Regulation
HK IA	Hong Kong Insurance Authority
HITL	Human-in-the-Loop; a governance mechanism requiring meaningful human review or intervention
ICT	Information and Communication Technology
IDD	Insurance Distribution Directive
ISO/IEC 42001	International standard for Artificial Intelligence Management Systems (AIMS)
KRI	Key Risk Indicator
LLM	Large Language Model; a generative AI model trained on large volumes of text data
MAS	Monetary Authority of Singapore
NIST AI RMF	NIST Artificial Intelligence Risk Management Framework
PRA	Prudential Regulation Authority (United Kingdom)
RAS	Risk Appetite Statement
Solvency II	EU prudential regulatory framework for insurers
Three Lines Model	Governance model consisting of business ownership, independent oversight and independent assurance
Veritas	MAS-led framework and toolkit supporting implementation of FEAT principles

Definition Agentic AI

AI systems that pursue defined objectives with varying levels of autonomy by planning, coordinating and executing tasks or processes, often interacting with multiple systems, tools, data sources or software environments, including the generation and execution of code, with limited human intervention.



Disclaimer

Dutch law is applicable to the use of this publication. Any dispute arising out of such use will be brought before the court of Amsterdam, the Netherlands. The material and conclusions contained in this publication are for information purposes only and the editor and author(s) offer(s) no guarantee for the accuracy and completeness of its contents. All liability for the accuracy and completeness or for any damages resulting from the use of the information herein is expressly excluded. Under no circumstances shall the CRO Forum or any of its member organisations be liable for any financial or consequential loss relating to this publication. The contents of this publication are protected by copyright law. The further publication of such contents is only allowed after prior written approval of CRO Forum.

© 2026 CRO Forum

The CRO Forum is supported by a Secretariat that is run by KPMG Advisory N.V.
Laan van Langerhuize 1, 1186 DS Amstelveen, or
PO Box 74500, 1070 DB Amsterdam
The Netherlands

www.thecroforum.org

