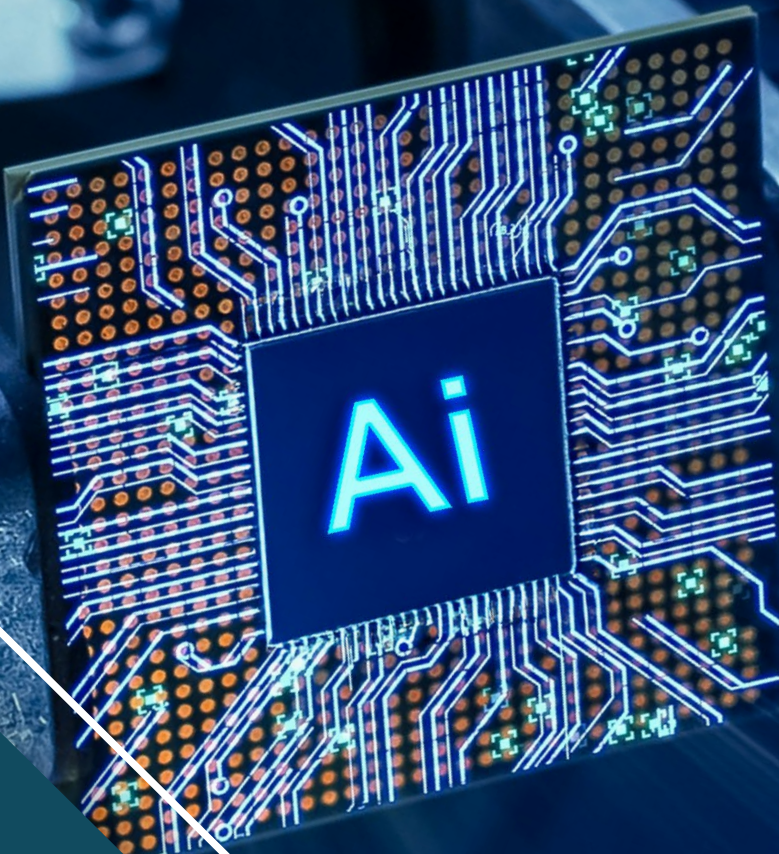


2026

AI Impact
Report Series
UK & Ireland edition

A glowing blue 'Ai' logo is centered on a square circuit board. The board is held by a mechanical gripper, and the background is a blurred industrial setting. The 'Ai' is in a stylized, sans-serif font.

How AI agents are amplifying conduct risk in insurance

Why governance doesn't stop at go-live in
the UK and Irish markets

Contents

About this report	3
Foreword	4
How AI is changing compliance in insurance: The age of AI agents	5
How AI agents are reshaping conduct risk in insurance	6
AI accountability in insurance: Regulatory expectations explained	8
Fairness, bias, and ethical risk in AI systems	9
Why insurers should be practicing continuous AI governance	11
How insurers can implement AI compliance at scale	13
Davies' AI compliance readiness checklist for insurers	15
The future of AI compliance in UK and Irish insurance	18

About this report

This report is part of Davies' AI Impact Series – a collection of research exploring the evolving role of artificial intelligence across the global insurance markets. Each report in the series will examine a critical dimension of AI adoption. As the opening report, this edition sets the foundation, focusing on AI compliance for insurers in the UK and Ireland deploying AI agents and the evolving expectations of regulators. Outlining the current regulatory landscape and highlighting key challenges and opportunities insurers face as they integrate AI into their business, it also includes a practical checklist for responsible and compliant AI deployment.

Who should read this report

Chief Risk Officers



Compliance leaders



AI and data governance teams



Insurance executives responsible for digital transformation



Understanding AI regulation in insurance

Is AI already regulated in UK and Irish insurance?

Yes. AI systems used in insurance are already subject to existing regulatory frameworks, including the FCA Consumer Duty, UK and EU GDPR, and the Central Bank of Ireland's Consumer Protection Code 2025. There is no regulatory gap – firms are expected to meet the same standards for fairness, transparency and accountability as they would for any other decision-making process.

What's changing with AI agents?

AI agents introduce continuous, adaptive decision-making. Unlike traditional models, they evolve over time, meaning compliance cannot be applied at a single point in the development lifecycle – it must be continuous.

Who's accountable for AI decisions?

Firms remain fully accountable for the outcomes of AI systems. Automation doesn't dilute responsibility, and regulatory frameworks such as Senior Managers & Certification Regime (SM&CR) reinforce clear ownership of AI-driven decisions.

Does the EU AI Act matter for UK insurers?

Yes. While not directly binding in the UK, the EU AI Act is already shaping market expectations, vendor standards, and cross-border operating models.



Foreword

Artificial intelligence (AI) is entering a defining phase for the insurance industry. What began as the use of largely static models to support discrete, well-scoped tasks is now evolving into more advanced systems capable of continuous decision-making, adaptive learning, and meaningful operational transformation. For insurers, this shift brings clear benefits: improved efficiency, faster responses, and the ability to focus human expertise where it adds the greatest value. But along with the benefits these systems bring, they also introduce a new and more complex set of compliance considerations.

As AI systems become more dynamic and autonomous, traditional approaches to governance, auditability, and accountability are being tested. Reduced human involvement in routine or administrative processes, combined with models that evolve over time, challenge established regulatory frameworks – particularly in highly regulated environments such as insurance. Ensuring appropriate oversight, transparency, and control is no longer a theoretical concern – it's a practical necessity.

“

Ensuring appropriate oversight, transparency, and control is no longer a theoretical concern – it's a practical necessity.

”

In the UK and Ireland, these challenges are unfolding against distinct but interconnected regulatory landscapes. The UK continues to pursue a principles-based, sector-led approach to AI regulation, while Ireland, as an EU Member State, prepares for the phased implementation of the EU AI Act. Despite these differences, insurers in both jurisdictions remain subject to robust existing regulatory requirements that apply fully to the use of AI today.

This report, the first in our AI Impact Series, examines what these developments mean in practice for insurers operating across the UK and Irish markets. It explores emerging challenges around responsibility, compliance assurance, and consistency as AI agents increasingly blur traditional operational and organisational boundaries. Our aim is to provide clarity, insight, and a practical foundation for organisations seeking to build AI capabilities that are not only innovative, but also compliant, accountable and sustainable.

We will cover:

How compliance is changing in the age of AI agents

The implications for conduct risk and customer outcomes

Regulatory expectations around accountability and governance

The growing importance of fairness, bias management and explainability

Why continuous AI governance is becoming essential

How insurers can operationalise AI compliance at scale

Paul O'Brien
Chief AI Officer
Davies



How AI is changing compliance in insurance: The age of AI agents

The insurance sector is undergoing a significant transformation in its use of AI. Where organisations once relied on static, rule-based or narrowly trained machine learning models, they are now increasingly deploying autonomous AI agents capable of continuous learning, adaptation, and decision-making. This marks a fundamental shift – from an incremental technology upgrade to a redefinition of the underlying risk profile.

Unlike traditional AI systems, AI agents operate with a higher degree of autonomy, often interacting dynamically with data, systems, and even customers in real time. Their ability to evolve behaviour based on new inputs introduces a level of unpredictability that

challenges conventional governance approaches. As a result, the compliance function must grapple with decision-making processes that are no longer fixed or easily auditable at a single point in time, but instead, continuously evolving.



The evolution of such technology brings into focus a critical reality: while the technology is new, the regulatory expectations are not – and there is no regulatory gap for AI. Existing, well-established frameworks – such as the FCA's Consumer Duty in the UK, the Central Bank of Ireland's requirements, and established data protection regimes including the UK GDPR and EU GDPR – remain fully applicable. These compliance frameworks already impose clear obligations around fairness, transparency, accountability, and consumer outcomes – all of which extend directly to AI-driven processes.

While regulatory frameworks are already in place, it's important to note the regulatory landscape is advancing to address the specific characteristics of AI systems. In particular, the EU AI Act – which is the first-ever legal framework on AI and is influencing countries even outside of the EU, including the UK – is setting a benchmark for AI governance, which includes these autonomous agents.

The EU AI Act also introduces a risk-based approach, placing heightened expectation on high-impact systems, and reinforcing the need for robust oversight, documentation, and human accountability. For insurers with cross-border operations, this creates an added layer of complexity in aligning multiple regulatory expectations.

It's becoming clearer to insurers who have previously treated compliance as a one-and-done, tick box exercise, that it can no longer be treated as such, nor be applied at a single point in the development cycle. Instead, it must evolve into a continuous, embedded capability, integrated into the design, deployment, and ongoing operation of AI agents. This means embedding controls that monitor behaviour in real time, ensuring traceability of decisions and enabling timely

intervention where risks emerge. It also requires closer collaboration between compliance, risk, technology, and business teams to ensure regulatory considerations are built into systems by design, rather than upgraded after deployment.



Ultimately, the role of compliance is shifting from gatekeeper to enabler, supporting innovation while ensuring AI is used safely, ethically, and in alignment with regulatory expectations. As insurers continue to adopt AI agents at scale, the ability to maintain trust – both with regulators and customers – will depend on how effectively compliance adapts to this new paradigm.



The role of compliance is shifting from gatekeeper to enabler



The following sections will explore in greater detail how compliance functions can respond to the unique challenges posed by AI agents and what practical steps organisations can take to operationalise responsible AI governance.

How AI agents are reshaping conduct risk in insurance

The adoption of AI agents in insurance introduces a new layer of complexity to conduct risk management. As these systems move beyond static outputs to continuous, autonomous decision-making, they increasingly intersect with core conduct obligations that sit at the heart of regulatory expectations. Ensuring AI-driven processes uphold principles such as fairness, transparency, and accountability is therefore critical – not only from a compliance perspective, but also to maintain customer trust.



Ensuring AI-driven processes uphold principles such as fairness, transparency, and accountability is critical to maintain customer trust”



Core conduct requirements

Even prior to the AI age, firms were being held accountable to act honestly, assess customer needs accurately, and prevent mis-selling – and this is emphasised even more so now that AI agents are becoming increasingly more autonomous.

Fair treatment of customers

At its core, conduct regulation requires insurers to deliver fair outcomes for customers. AI agents, particularly those involved in insurance underwriting, claims handling, or servicing customers, directly influence these outcomes. The dynamic nature of agentic systems means decisions may evolve over time, raising questions about consistency and fairness. Ensuring AI agents operate within clearly defined parameters – and these parameters align with customer interests – is essential to meeting obligations under frameworks like the FCA's Consumer Duty and the Central Bank of Ireland's Consumer Protection Code 2025.

Transparency of decisions

Transparency remains a fundamental expectation, particularly where AI agents inform or make decisions on pricing, claims acceptance or rejection, and eligibility. However, as models become more complex, explainability becomes more challenging and complexity can hinder AI responses. Insurers must therefore ensure AI-driven outcomes can be translated into clear, intelligible explanations that are meaningful to customers and regulators alike.

Avoidance of unfair discrimination

AI agents pose significant risks in relation to discriminatory outcomes, particularly where they rely on large datasets that may embed historical bias. Regulatory and legal frameworks – including the Equality Act 2010 (GB) and the Equal Status Acts (2000-2018) in Ireland – set strict expectations to prevent unfair discrimination on protected grounds. Under this legislation, insurers must ensure AI systems don't directly or indirectly lead to biased outcomes, particularly in areas such as pricing, policy eligibility, and claims assessment.

Contestability and redress

A key principle in consumer protection is the ability for customers to challenge decisions and seek redress. With AI agents increasingly making or influencing decisions, insurers must check customers have avenues to contest outcomes effectively. This includes providing clear explanations of how decisions were reached and maintaining accessible pathways for human review. Without these mechanisms, there's risk of undermining procedural fairness and breaching regulatory expectations.



3 key risks insurers are underestimating with AI agents

While the above conduct requirements aren't new, AI agents introduce distinct risk characteristics that amplify traditional conduct concerns.



Autonomous decision loops

AI agents can operate in continuous feedback cycles, refining decisions based on new data inputs. Without proper controls, this can lead to unintended outcomes, such as overly aggressive pricing strategies or systematic claims denials that deviate from intended policy.



Opacity and reduced explainability

The increased complexity of agentic systems can obscure how decisions are made, particularly when multiple models or data sources interact. This reduces transparency and complicates both internal oversight and external explanation when not managed.



Reinforcement of bias over time

Unlike static systems, AI agents may continuously learn from their own outputs. If bias is present – either in training data or decision logic – it can be amplified over time, leading to progressively unfair outcomes if not actively monitored and corrected.

While unique in nature, these risks must be understood in the context of existing regulatory frameworks, which already provide clear principles applicable to AI-driven decision-making.

For insurers based in, and operating in, the UK, this means considering the risks according to GDPR principles and the FCA's Consumer Duty – both of which already introduce additional guardrails including principles of fairness, transparency, purpose limitation and accountability. GDPR principles, specifically, include provisions related to automated decision-making, granting individuals the right to meaningful information about the logic involved, and in certain cases, the right not to be solely subject to automated decisions. These provisions are particularly relevant in the context of AI agents, reinforcing the need for explainability, human oversight, and robust data governance.

For Irish insurers, and those operating in Ireland, the risks will need to be assessed in line with the Central Bank of Ireland's Consumer Protection Code 2025. The code reinforces obligations around acting in customers' best interests, transparency of information and handling of complaints. AI systems must align with these requirements, particularly when ensuring decisions are clearly communicated and customers can engage effectively with firms to challenge outcomes.



What insurers need to do to manage AI-driven conduct risk

The convergence of agentic AI capabilities and existing conduct requirements underscores a critical point: conduct risk in the age of AI agents is not fundamentally new but is materially amplified. The autonomous and adaptive nature of these systems require insurers to move from periodic oversight to continuous monitoring, ensuring risks are identified and managed in real time.

To meet regulatory expectations, firms must embed conduct considerations into the full lifecycle of AI agents – from design and training through to deployment and ongoing operation. This includes implementing controls to detect bias, ensuring explainability of decisions, maintaining clear audit trails, and enabling effective customer challenge and redress mechanisms.

As AI agents become more deeply embedded in insurance operations, the ability to manage conduct risk effectively will be a key differentiator – not only for achieving compliance, but in maintaining trust and delivering consistently fair outcomes for customers.



The ability to manage conduct risk effectively will be a key differentiator



AI accountability in insurance: Regulatory expectations explained

Industry regulators are being abundantly clear in that the adoption of AI doesn't dilute accountability. Firms remain fully responsible for outcomes, regardless of the level of automation embedded within decision-making processes. In the UK financial services market, this principle is most explicitly enforced through the Senior Managers and Certification Regime (SM&CR), which serves as the primary framework for assigning and evidencing individual accountability for AI-related decisions.

“

The ability to manage conduct risk effectively will be a key differentiator

”

Under SM&CR, firms must demonstrate responsibility for AI systems is clearly owned at an appropriate seniority level – with this expectation extending across several critical dimensions:

Clear ownership of AI agent decisions:

Every AI-enabled process or decision must have a designated accountable individual. There must not be any ambiguity regarding who is responsible for the development, deployment, and outcomes of AI systems.

Audit trails and decision traceability:

Firms must ensure AI decisions are explainable and reconstructable. This requires comprehensive audit trails capturing inputs, outputs, model logic (where feasible), and decision pathways to support regulatory review and internal challenge.

Governance over third-party and vendor AI agents:

The use of external AI solutions doesn't transfer accountability. Firms must implement robust vendor governance frameworks, ensuring third-party technologies meet internal standards for risk, transparency, and compliance.

Meaningful human oversight:

Regulators are increasingly scrutinising the effectiveness of human involvement in AI-driven processes. Oversight must go beyond passive approval and rubber-stamping – it must involve active review, challenge, and intervention where necessary.



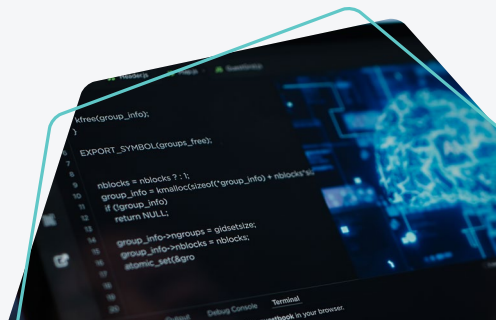
While these expectations are already well established in the UK context, firms must also consider the broader regulatory landscape. The EU AI Act, despite not being directly applicable to the UK, is emerging as a de facto benchmark for responsible AI governance and is shaping supervisory expectations globally.



“

The EU AI Act is emerging as a de facto benchmark for responsible AI governance

”



Key influential elements from the EU AI Act include:

Risk classification frameworks:

Categorisation of AI applications based on risk, with use cases such as pricing, underwriting, and claims decisions typically falling into high-risk categories which require heightened governance and control requirements.

Mandatory compliance obligations:

High-risk AI systems are subject to stringent requirements, including:

- ✓ Comprehensive documentation and record-keeping.
- ✓ Transparency obligations: These ensure stakeholders understand how AI systems operate and impact outcomes.
- ✓ Enforced human oversight mechanisms: These are designed to prevent unchecked automated decision-making.

Considered together, these regulatory frameworks reinforce a consistent message: firms must embed accountability, transparency, and control at the core of their AI strategies. Automation may enhance efficiency and scalability, but it doesn't replace the need for clear governance, demonstrable oversight, and ultimately, human responsibility.



Fairness, bias, and ethical risk in AI systems

As AI agents become more autonomous, adaptive, and deeply embedded within financial services decision-making, fairness and ethical risk can no longer be treated as abstract principles – they become measurable, auditable, and enforceable components of regulatory compliance.

Unlike traditional rules-based automation, AI agents often rely on continuous learning, complex, and multi-source data inputs, and hyper-personalised decisioning at scale. While these capabilities can enhance efficiency and accuracy, they also create new pathways through which bias can be introduced, amplified, or obscured.

Bias risk may arise from historical data reflecting past inequalities, proxy variables that unintentionally correlate with protected characteristics, or feedback loops where model outputs influence future data. In autonomous or semi-autonomous systems, these effects can compound over time, making bias harder to detect and remediate without deliberate governance and monitoring.



These risks translate directly into regulatory and compliance exposure across several dimensions:



Discrimination risk

AI systems used in pricing, creditworthiness assessments, or underwriting may inadvertently disadvantage certain customer groups if models correlate risk with socio-economic, geographic, or behavioural proxies. For example, an underwriting model that optimises for loss ratios could systematically exclude specific postcodes or employment types, creating indirect discrimination even in the absence of explicit protected attributes.



Customer detriment

When using AI in insurance claims handling, it can lead to unfair outcomes if models are overly conservative, insufficiently explainable, or trained on biased historical settlement data. This may result in higher rejection rates, lower settlement values, or inconsistent treatment of similar claims, exposing firms to conduct risk and customer and Financial Conduct Authority (FCA) complaints.



Reputational and conduct risk

In areas like fraud detection, AI systems that disproportionately flag certain customer segments for investigation can damage trust and brand reputation. False positives affecting vulnerable or minority groups may attract regulatory scrutiny, reputational damage, and loss of customer confidence, even where intent is absent.

Given these risks, regulators increasingly expect firms to demonstrate fairness, and ethical considerations are actively managed throughout the AI lifecycle. This includes clear evidence that bias is not only considered at design stage but monitored and addressed on an ongoing basis.



Key regulatory expectations include:

Regular testing for bias and disparate outcomes

Particularly for high-impact use cases such as underwriting, pricing, claims settlement, and credit decisions

Evidence of fairness outcomes

Including metrics, thresholds, and management actions taken where imbalance or unintended harm is identified

Explainability and transparency

Ensuring AI-driven decisions can be understood, challenged, and justified by both internal stakeholders and regulators

In practice, this means firms must be able to explain why an AI system made a particular decision, demonstrate that outcomes are consistent and proportionate, and show that appropriate controls exist to prevent unfair or unethical impacts.



Why insurers should be practicing continuous AI governance

Traditional compliance models – based on pre-deployment approvals, static documentation, and periodic reviews – were developed for systems whose behaviour remained largely stable once live. AI agents fundamentally disrupt this assumption.

When deployed into environments, AI agents can operate autonomously, interact with customers and intermediaries, adapt to new data, and influence decisions across underwriting, pricing, claims handling, fraud detection, and customer engagement. As a result, compliance risk doesn't halt at deployment; it evolves continuously throughout the agent's operational life.

For insurers, this creates a clear challenge: compliance can no longer be treated as a prior activity AI agents pass through before production. Instead, it must operate as a continuous governance capability that remains active for as long as the agent is in use. Therefore, the core compliance question has shifted from "Was the AI agent compliant when deployed?"

to "Is the AI agent continuing to operate compliantly within approved risk boundaries today?"

This shift is particularly important in insurance, where AI agents may affect regulated outcomes such as pricing, fairness, claims decisions, customer treatment and fraud interventions. In these cases, inappropriate agent behaviour can quickly translate into conduct risk, prudential risk, or breaches of regulatory expectations around fairness, transparency, and accountability.



The core compliance question has shifted from "Was the AI agent compliant when deployed?" to "Is the AI agent continuing to operate compliantly within approved risk boundaries today?"



In practice: governing AI agents in claims

In claims handling – one of the most conduct-sensitive areas of insurance – AI agents are increasingly used to support triage, workflow orchestration and decision support. Davies' ClaimPilot product suite illustrates how insurers can deploy these agentic capabilities while maintaining continuous governance.

By combining AI-driven automation with defined escalation thresholds, decision traceability and human oversight for higher-impact outcomes, ClaimPilot demonstrates how conduct risk controls can remain active throughout the lifecycle of an AI agent, rather than ending at deployment.



How to have continuous oversight of AI agents

Once an AI agent is deployed into an insurance workflow, its behaviour must be continuously supervised. Unlike traditional rules-based systems, AI agents may change their behaviour over time due to data drift, evolving claims patterns, seasonal effects, or retraining. From a compliance perspective, this introduces ongoing exposure to risks such as indirect discrimination, inconsistent claims outcomes, unexplained decision-making, or deviations from approved underwriting or pricing strategies.

Continuous monitoring therefore becomes a core compliance control for insurers. This includes:

- ✓ Real-time oversight of agent decisions and actions
- ✓ Monitoring of input data for drift or bias
- ✓ Ongoing assessment of performance against approved conduct and risk thresholds

Importantly, this monitoring supports both operational risk management and regulatory assurance by generating auditable evidence that the AI agent remains within its approved scope of use.

Monitoring alone, however, is insufficient. Deployed AI agents must also be subject to real-time risk controls that can actively enforce compliance boundaries. In the insurance industry, this may include automated thresholds that trigger alerts or escalation when claims outcomes diverge from expected norms, human-in-the-loop review for high-impact

decisions (such as claims rejections or premium adjustments), or technical controls that restrict agent autonomy in sensitive or vulnerable customer interactions. These mechanisms ensure compliance is enforced dynamically, at the point of decision-making, rather than retrospectively through complaints or regulatory findings.

The role of lifecycle-based compliance

Effective compliance for agentic AI in insurance must span the full lifecycle of the agent – from initial design through deployment, live operation, retraining, and eventual retirement. Each phase introduces distinct compliance considerations, including:

Risk classification



Documentation



Explainability requirements



Change management



Accountability for outcomes



This lifecycle-based approach aligns closely with established model risk management practices already familiar to insurers – for example, Solvency II internal-model governance, the actuarial TAS standards, as well as the FCA and PRA guidance.



ISO/IEC 42001 further strengthens this approach by framing AI governance as a management system embedded into organisational operations. For insurers, this supports the integration of compliance into day-to-day AI agent deployment processes such as machine learning operations (MLOps) pipelines, incident management, model updates, and change control. The standard's emphasis on risk-based controls, accountability, and continuous improvement mirrors regulatory expectations that insurers demonstrate not just initial compliance, but sustained control over AI-enabled decision-making.

Together, these frameworks reflect a clear trajectory for insurance regulation – AI agents must be governed as living systems rather than static models, with compliance mechanisms that operate continuously and adapt as risks evolve.



“

AI agents must be governed as living systems, not static models

”



How insurers can implement AI compliance at scale

As insurers move from isolated AI use cases to more widespread, agent-based models embedded across the value chain, the challenge is no longer whether AI can be governed, but how that governance is operationalised at scale.

In practice, insurers are increasingly managing multiple AI agents operating across underwriting, claims, customer service, pricing, fraud detection, and internal operations. These systems often span both legacy platforms and modern cloud-based architectures, making it difficult to apply consistent controls, monitoring and documentation across the enterprise. Fragmented technology estates can result in uneven risk management, duplicated assurance activity and limited end-to-end visibility over how AI is being used.

Operationalising AI compliance requires effective coordination between compliance, risk, IT, data science, and frontline operations. In many organisations, accountability for AI risk is spread across many different people, divisions, and departments, with unclear ownership across the model lifecycle – from design and training through to deployment, monitoring and change management. For UK and Irish insurers operating under increasing supervisory scrutiny, this lack of clarity presents a material governance risk.

The emergence of agentic AI systems further intensifies these challenges by initiating actions, interacting with other

systems, and adapting its behaviour over time, increasing exposure to AI-specific security risks like prompt injection, data poisoning and adversarial attacks. These risks sit alongside more established concerns around data protection, explainability, bias, and model drift, and require targeted controls that many existing risk frameworks were not designed to address.

Third-party risk is another critical dimension to consider. Insurers are increasingly reliant on vendor-provided AI models, embedded tools and cloud-based AI services, often integrated deep within core processes. Regulators in both the UK and Ireland are clear that outsourcing AI does not outsource accountability. Firms are expected to ensure third-party AI agents meet the same governance, security and ethical standards as internally developed systems. This drives an increased need for robust due diligence, ongoing oversight and contractual safeguards, particularly where vendors are using opaque or rapidly evolving models.

Operational resilience considerations are also crucially important for insurers – particularly because greater reliance on a small number of critical AI and cloud service providers heightens concentration risk and increases regulatory focus on resilience, exit planning and service continuity. For insurers, AI governance must therefore align closely with existing operational resilience frameworks, ensuring critical business services remain tolerable in the event of AI system failures, data corruption or provider outages.

To address these challenges, insurers are increasingly turning to a set of key enablers:

AI governance platforms & tooling:

These help centralise oversight, automate controls and provide real-time visibility across diverse AI estates.

Standardised risk frameworks:

Aligned to regulatory expectations and enterprise risk management, these enable more consistent assessment and decision-making across functions and use cases.

Employee training:

Ensuring employees undertake AI compliance training enables them to use AI responsibly and understand its limitations, failure modes, and associated risks.



From a practical governance perspective, Data Protection Impact Assessments (DPIAs) and AI Impact Assessments (AIAs) are becoming essential mechanisms for assessing and mitigating AI risks prior to deployment. When embedded into product

development and change processes, these assessments provide a structured way to detect data protection, ethical, operational and security risks early, and to evidence compliance to regulators and auditors.

Insurers also need to be aware of the growing expectation for them to maintain a centralised AI system inventory or register. This acts as the backbone of effective AI governance, supporting auditability, regulatory engagement and ongoing risk management. A well-maintained inventory enables firms to answer fundamental supervisory questions including

Where is AI being used?



What purpose is AI being used for?



What data is AI using?



Who is accountable for AI use?



What controls are in place for AI usage?



In an environment of accelerating AI adoption and evolving regulatory expectations, the ability to operationalise AI compliance at scale is rapidly becoming a core capability for insurers in the UK and Ireland. Now more than ever, it's crucial for insurers not to treat it as a compliance exercise, but a foundational element of sustainable and resilient AI adoption.



Davies' AI compliance readiness checklist for insurers

For insurers operating in the UK and Ireland, the transition to agentic AI requires more than high-level governance principles – it demands clear, evidence-based controls that can withstand regulatory scrutiny. Supervisory expectations from the FCA, PRA, and Central Bank of Ireland continue to emphasise accountability, transparency and operational resilience, particularly where AI systems influence customer outcomes.

The following checklist outlines the core elements insurers should be able to evidence when deploying and scaling AI agents across the enterprise.

01 Clear ownership and accountability

Each AI agent should have defined ownership across its lifecycle, with named individuals responsible for:

- ☐ Design
- ☐ Deployment
- ☐ Oversight and;
- ☐ Risk acceptance

This must align with existing Senior Management accountability frameworks where relevant.

02 Decision transparency, auditability & logging

Insurers must maintain robust audit trails, capturing how AI agents make decisions including:

- ☐ Inputs
- ☐ Outputs and;
- ☐ Any human interventions

Logging should be sufficiently detailed to support:

- ☐ Internal audit
- ☐ Model validation
- ☐ Regulatory enquiries
- ☐ Customer complaints and;
- ☐ Dispute resolution

03 Risk classification and materiality assessment

AI use cases should be formally classified according to risk, drawing on concepts aligned to the EU AI Act (e.g., risk-based categorisation) where appropriate, even for UK-based firms. This enables insurers to apply:

- ☐ Proportionate controls
- ☐ Prioritise oversight on high-impact use cases and;
- ☐ Ensures consistency across jurisdictions - especially for cross-border operations

The 5 pillars of AI compliance in insurance:

Accountability

Transparency

Fairness

Continuous monitoring

Operational resilience



04 Human oversight and intervention mechanisms

Insurers are expected to retain meaningful human oversight, particularly for decisions that impact customers, pricing, claims outcomes, or access to products. This includes clearly defined:

- ☐ Escalation pathways
 - ☐ Override mechanisms and;
 - ☐ Defined thresholds for human review
-

05 Bias, fairness and outcomes testing

Regular bias and fairness testing should be embedded into the model lifecycle, with a focus on customer outcomes and protected characteristics. Insurers should be able to demonstrate:

- ☐ Ongoing monitoring for unintended bias
 - ☐ Periodic review of model performance and;
 - ☐ Documented remediation actions
-

06 Data governance, lineage and GDPR compliance

The data-intensive nature of AI means firms must evidence full data lineage and strong AI data protection compliance frameworks and controls including:

- ☐ Traceability of data used for training, validation and live operation
 - ☐ Lawful basis for processing under GDPR
 - ☐ Data minimisation
 - ☐ Purpose limitation and;
 - ☐ Controls to prevent unauthorised data use or leakage
-

07 Third-party and vendor risk management

Insurers must ensure equivalent governance standards apply to third-party systems. This requires:

- ☐ Robust pre-contract due diligence
- ☐ Transparency over model design
- ☐ Training data and limitations
- ☐ Contractual provisions on audit rights
- ☐ Performance standards
- ☐ Incident reporting and regulatory cooperation, as well as;
- ☐ Ongoing monitoring of vendor performance and risk

08 Continuous monitoring and model lifecycle management

Insurers must establish ongoing monitoring and performance review processes covering:

- ☐ Model drift and degradation
 - ☐ Changes in data inputs and operating environment and;
 - ☐ Emerging risks linked to agentic AI capabilities and evolving model behaviour
-

09 AI literacy and employee training

As AI becomes embedded across business functions, insurers must invest in AI compliance training at all levels of the organisation. This training should cover:

- ☐ Appropriate use of AI tools and outputs
 - ☐ Understanding limitations
 - ☐ Risks and failure modes, as well as;
 - ☐ Role-specific responsibilities for managing AI risk
-

10 Pre-deployment risk assessments

Data Protection Impact Assessments (DPIAs) and AI Impact Assessments (AIAs) should be embedded as standard controls within development and change processes to help firms:

- ☐ Identify privacy, ethical and operational risks early
- ☐ Define mitigation strategies prior to deployment and;
- ☐ Evidence AI and GDPR compliance and broader governance expectations

Importantly, DPIAs and AIAs should be living documents, updated as systems evolve.

11 Incident management and response frameworks

Firms should implement AI-specific incident response processes, integrated into broader operational risk frameworks. These should address:

- ☐ Model failures and unexpected behaviour
- ☐ Data breaches or integrity issues and;
- ☐ Security incidents such as prompt injection or adversarial attacks

Clear reporting, escalation and remediation protocols are essential in these AI compliance frameworks.

12 Customer outcomes, complaints and redress

AI governance must ultimately support fair customer outcomes. Insurers should ensure:

- ☐ Transparent communication where AI is being used in decision-making
- ☐ Accessible complaint mechanisms and;
- ☐ Clear redress processes where AI-driven decisions are challenged or found incorrect

This aligns closely with Consumer Duty expectations in the UK and similar conduct standards in Ireland.

13 Operational resilience and third-party dependencies

AI deployment must be considered through the lens of operational resilience. Insurers should assess:

- ☐ Reliance on critical AI and cloud providers
- ☐ Potential single points of failure
- ☐ Exit and substitution strategies and;
- ☐ The impact of AI disruption on important business services

These considerations should be integrated into existing resilience frameworks and scenario testing.

Individually, these controls are familiar. The challenge for insurers lies in implementing them consistently across increasingly complex, distributed and dynamic AI ecosystems.



A robust compliance framework will depend on:

Centralised oversight (e.g. AI inventories and governance platforms)



Standardisation of frameworks and processes



Strong cross functional collaboration



For UK and Ireland insurers, demonstrating this level of control is not only key to AI regulatory compliance – it is central to building trust, resilience and long term value from AI adoption.



The future of AI compliance in UK and Irish insurance



The widespread deployment of AI agents marks a pivotal shift for insurers in the UK and Ireland – one that brings significant opportunity, but equally heightened regulatory, ethical, and operational responsibility. As agentic systems become more autonomous, adaptive, and embedded within core insurance processes, the challenge is no longer simply adopting the tech but implementing AI governance and compliance in a way that is demonstrably fair and resilient.



Crucially, while the technology is evolving rapidly, the underlying regulatory expectations remain consistent. Existing frameworks, spanning conduct risk, data protection, accountability, and consumer outcomes, apply fully to AI-driven decision-making. At the same time, emerging standards such as the EU AI Act are raising the bar for what effective AI governance looks like, particularly for insurers operating across borders. Together, these forces are reinforcing a clear message: firms must be able to evidence control, transparency, and accountability across the full lifecycle of AI agents.

This requires a fundamental shift in how compliance is approached. Static, point-in-time assessments are no longer sufficient in the context of systems that learn and evolve. Instead, insurers must look at embedding continuous AI governance into the design and operation of AI agents at scale – supported by real-time monitoring, robust controls, and clear ownership structures.

Firms that succeed will be those that move beyond viewing compliance as a constraint and instead treat it as an enabler of sustainable innovation. By operationalising AI governance effectively – through strong frameworks, cross-functional collaboration, and investment in tools, skills, and oversight – insurers can not only meet regulatory expectations but also build trust with customers and regulators alike.

Ultimately, the safe and responsible deployment of AI agents will become a defining capability for the industry. Those who can demonstrate consistent, fair, and accountable AI-driven outcomes will be best positioned to unlock the full value of these technologies, while maintaining the confidence of both regulators and the customers they serve.



Davies Group Limited

Registered Company No. 06479822. Registered in England and Wales
Registered Office: 5th Floor, 20 Gracechurch Street, London EC3V 0BG

Disclaimer and Copyright Notice

This document is intended for general informational purposes only, does not take into account the reader's specific circumstances and is not a substitute for professional or legal advice. Readers are responsible for obtaining such advice from licensed professionals. The information included in this document has been obtained from sources we believe to be reliable and accurate at the time of issue. Davies Group Limited accepts no liability for errors or omissions in this document. This document and the information contained within it are provided for personal use only. No part may be reproduced, stored in a retrieval system or transmitted in any form or by any means electronic, mechanical photocopying, microfilming, recording, scanning or otherwise for commercial purposes without the written permission of the copyright holder.

© Davies Group 2026. All rights reserved. Telephone calls may be recorded.