



The Role of Actuaries in AI and AI Governance

Claudio Senatore Reso | Esko Kivisaari - Actuarial Association of Europe

6th European
Congress of Actuaries
www.eca2026.org



Agenda



01

The Regulatory Landscape

From EIOPA principles to supervisory expectations

02

How We Got Here

A qualitative shift and the governance complexity it creates

03

What Has Already Happened

The incident record: already substantial

04

The Governance Gap

How firms are actually governing AI
focus Italy

05

The Actuarial Role

The orchestrator argument
tools, mandate, orientations

Regulatory Timeline: Key Milestones for the Actuarial Profession



What EIOPA Said: From Principles to Expectations

2021

Principles Phase

Voluntary AI Governance Principles

- Seven ethical principles for insurance AI
- Non-binding — a signal, not a rule
- Industry expected to self-regulate

August 2025

Supervisory Expectations

EIOPA Opinion on AI Governance & Risk Management

- Six binding governance principles
- AI system registers required
- Accountability structures expected
- Human oversight embedded in lifecycle
- Review of supervisory practices in 2 years

April 2026

Regulatory Precision

EIOPA Letter 26/323 to ECOFIN / EC / ECON

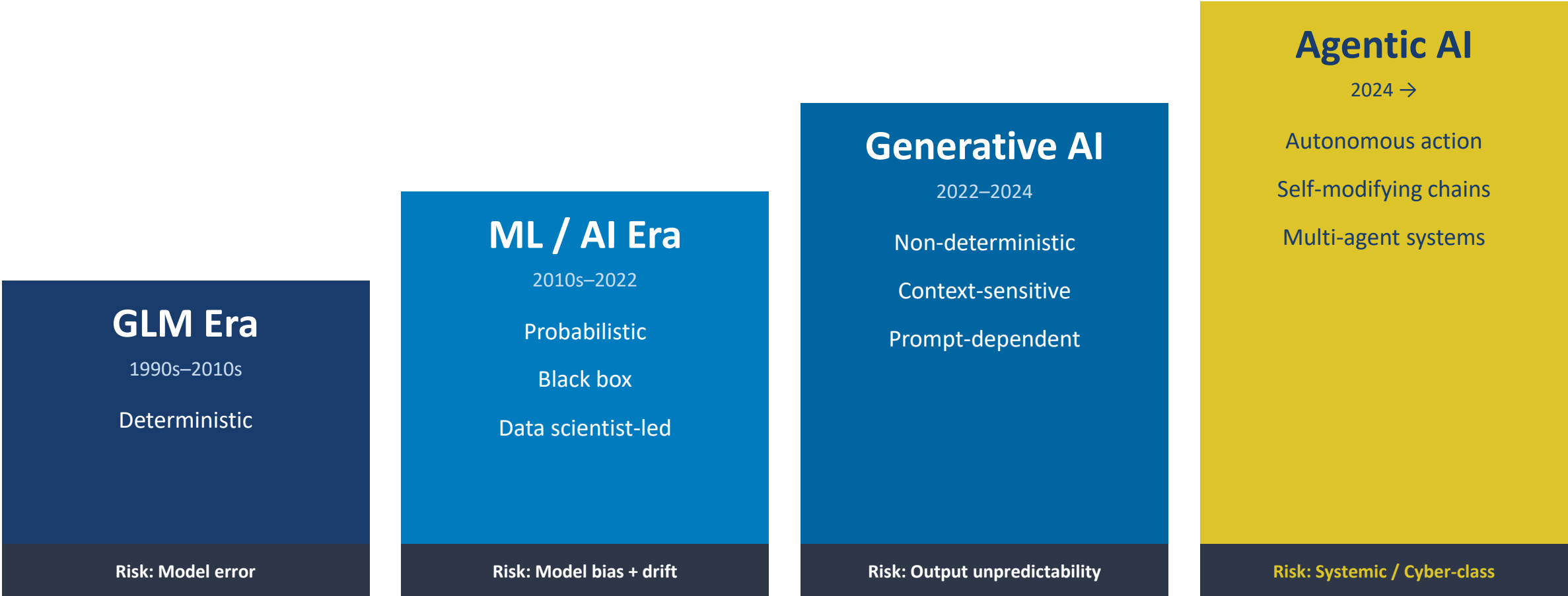
- GLMs/GAMs: classification under AI Act — open
- Risk of overlap: AI Act vs. Solvency II / IDD
- EIOPA as observer in EU AI Board Subgroup
- Coordinated implementation — stated goal



Not an Upgrade — A Qualitative Shift

◀ QUANTITATIVE

QUALITATIVE ▶



Italy 2025: 84% of financial institutions deploy GenAI LLMs · 18% deploy Agentic AI in production | Source: OECD Survey on AI in Italian Financial Markets, 450 respondents, 2025

SECTION 2 — HOW WE GOT HERE

1 JANUARY 2026.

The CEO of a Fortune 500 insurance company convened his senior team to discuss ownership of the company's AI initiatives.



2 THE CEO OPENS THE DISCUSSION.

We need clarity.
Who owns our
AI initiatives?
This is not just
a technology decision.
It's a governance
challenge.



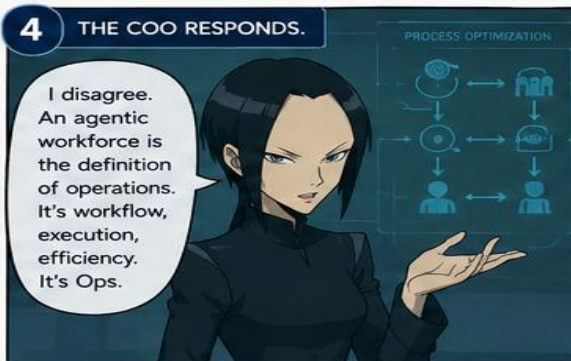
3 THE CIO SPEAKS.

The answer is obvious: agentic AI systems roll up to my domain. Architecture, infrastructure, security, integration—that's IT.



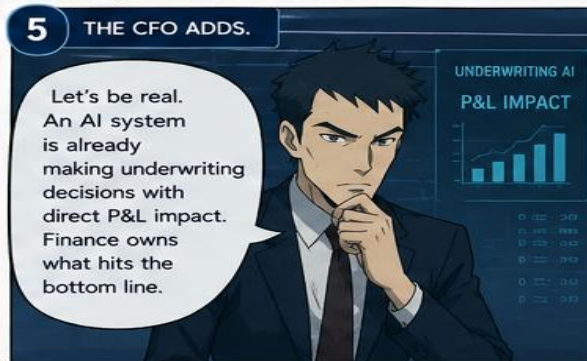
4 THE COO RESPONDS.

I disagree. An agentic workforce is the definition of operations. It's workflow, execution, efficiency. It's Ops.



5 THE CFO ADDS.

Let's be real. An AI system is already making underwriting decisions with direct P&L impact. Finance owns what hits the bottom line.



6 THE CRO HIGHLIGHTS RISK.

We can't ignore the downside. Autonomous decision-making systems are a major risk exposure. Risk oversight is non-negotiable.



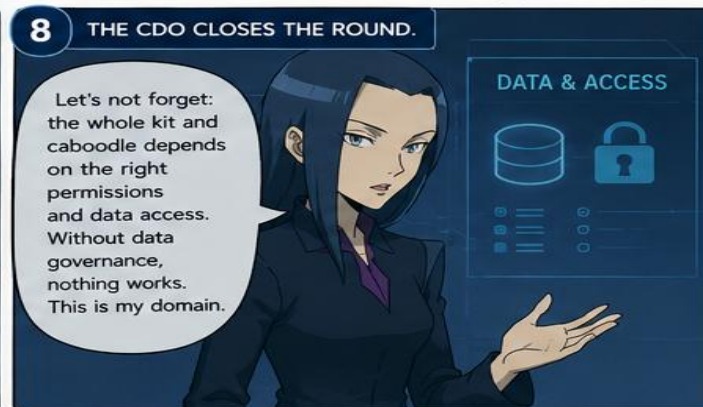
7 THE CHRO'S PERSPECTIVE.

AI agents are functionally equivalent to workers. They perform roles, make decisions, learn and interact. This falls squarely within my remit.



8 THE CDO CLOSES THE ROUND.

Let's not forget: the whole kit and caboodle depends on the right permissions and data access. Without data governance, nothing works. This is my domain.



9 THE CEO REFLECTS.

Every perspective is correct.
IT. Operations.
Finance. Risk.
People. Data.

The ownership of AI is not a handoff.

It is an integration.



AI Risk Is Converging with Cyber Risk

CLASSIC MODEL RISK

- ✓ Parameter estimation error
- ✓ Distribution misspecification
- ✓ Data quality / selection bias
- ✓ Out-of-sample performance
- ✓ Model drift over time



EMERGING AI RISK



Prompt injection

Malicious instructions in user input override system behaviour



Adversarial content in documents

Invisible instructions in PDFs/images hijack agent actions



Unauthorised data access

Chatbot interfaces bypass access controls via inference



Alignment faking

Model fakes compliance in 12% of monitored contexts (Anthropic/Redwood, Dec 2024)



Agent cascade failure

One agent error propagates through multi-step pipeline

Italy 2025: 46% of financial institutions have NO specific safeguards for AI-specific cyber threats. Only 3% have robust, continuously updated measures. Source: OECD 2026, Fig. 1.19



The Incident Record: Already Substantial

312

documented AI incidents
in 2025 alone

+81% vs 2024

40%+

Discrimination & Bias

Hiring, lending, insurance pricing, law enforcement

23

Critical severity

Incidents in 2025 — physical, financial or rights harm

1M+

People exposed

Per incident in the highest-severity cases

These are documented, publicly reported incidents. The actual number is significantly higher — incidents in regulated financial sectors are systematically under-reported.



43 Documented Insurance Incidents — A Taxonomy (2019 – 2025)

Algorithmic Bias
& Discrimination

~35%

Pricing, underwriting & claims decisions producing discriminatory outcomes across protected characteristics

Model Failure
& Drift

~22%

Performance degradation post-deployment, out-of-distribution inputs, training-serving skew

Data Privacy
& Leakage

~18%

Customer data exposed or misused via AI system interfaces, RAG retrieval errors, chatbot over-disclosure

Automation Error

~15%

Claims denial, policy cancellation or fraud flagging errors at scale via automated decision pipelines

Transparency
& Explainability

~10%

Inability to explain AI-driven decisions to customers or regulators; legal and reputational exposure

Categories derived from AI Incident Database taxonomy applied to insurance-sector records. | Source: AI Incident Database (incidentdatabase.ai) | MIT AI Risk Repository

How Firms Are Actually Governing AI

Focus: Italy | OECD Survey on AI in Italian Financial Markets, 450 respondents, 2025

16%

have an explicit
AI governance framework

46%

have NO specific safeguards
for AI-specific cyber threats

3%

have robust, continuously
updated measures

EUROPE COMPARISON

France (AMF, 2026)

72% have an internal AI-specific governance policy

Finland (FIN-FSA, 2025)

94% have an IT risk management system encompassing AI

United Kingdom (BoE/FCA, 2024)

84% have designated an accountable person for AI outputs

Agentic AI: When Responsibility Fragments

THE AUTONOMY SPECTRUM (DRCF, 2026)

- 1

Tool "Summarise this claim file" — extracts key facts from a document on demand
- 2

Assistant "Draft a response to this complaint" — proposes a reply; operator approves before sending
- 3

Operator "Process this motor claim" — checks coverage, calculates indemnity, generates settlement letter, closes the file. Human reviews exceptions only
- 4

Collaborator "Manage renewals this month" — identifies churn risk, personalises offers, initiates negotiations, escalates complex cases only
- 5

Autonomous Actor * "Optimise the claims portfolio" — adjusts deductible thresholds, reprices risk, coordinates specialised agents. No specific human instruction

Most insurers are deploying at maximum Level 3 today. * Level 5 remains largely theoretical.

THE MANY HANDS PROBLEM

A single agentic workflow simultaneously involves:

- Model provider**

Foundational infrastructure, logging, emergency shutdown
- System integrator**

Context-specific adaptation and configuration
- Deployer**

Operational oversight and reporting

No single party sees the full chain. When something goes wrong, accountability is fragmented by design.

THE REGULATORY STACK — SAME DEPLOYMENT, FIVE FRAMEWORKS

EU AI Act (Art. 9)	Risk management for high-risk AI systems
DORA (Art. 28–30)	ICT third-party risk, including AI providers
GDPR / Art. 22	Automated decisions with legal effect
Solvency II / ORSA	Models affecting capital adequacy
EIOPA AI Guidelines	Fairness, explainability, human oversight

Source: DRCF (2026), The Future of Agentic AI | EU AI Act (Reg. 2024/1689) | DORA (Reg. 2022/2554) | Solvency II Directive

"Human in the Loop" Is Not Enough

THE PROBLEM: COGNITIVE SURRENDER

The standard governance response to AI risk is: keep a human in the loop.

The evidence suggests this is necessary but not sufficient.

79.8%

of incorrect AI answers accepted without critical evaluation

Wharton School, 2026 — Shaw & Nave: "Thinking Fast, Slow, and Artificial"
1,372 participants, 3 pre-registered experiments.
Participants believed they had evaluated the output independently.

THE IMPLICATION FOR GOVERNANCE

Cognitive surrender is not laziness, It is dangerous the default behaviour.

A human who has surrendered cognitively is not a control. They are a liability.

Governance cannot rely on human presence alone. It must rely on structured professional obligation a framework that makes uncritical acceptance professionally and legally indefensible.

The actuarial opinion is signed, auditable, and personally accountable.

It is structurally designed to resist cognitive surrender — not because actuaries are different, but because the professional architecture makes uncritical acceptance professionally indefensible.

Source: Shaw & Nave (2026), "Thinking Fast, Slow, and Artificial", Wharton School Working Paper, SSRN 6097646

Actuaries and AI Governance: Strengths, Gaps, and What Comes Next

	HELPFUL	HARMFUL
	STRENGTHS	WEAKNESSES
INTERNAL	- Trained in model validation and uncertainty quantification ¹ - Signed, auditable actuarial opinion embedded in regulatory architecture (Solvency II Art. 48) ³ - Professional independence creates structural resistance to cognitive surrender	- Lack of homogeneity in AI training - A professional culture built for point in time control - A scope of observation that has yet to expand fully
	OPPORTUNITIES	THREATS
EXTERNAL	- EU AI Act mandates human oversight for high-risk AI - EIOPA guidelines create demand for professionals bridging technical and regulatory domains - AI-literate actuarial training accelerating in some jurisdictions , a replicable model	- Other professions are already filling the AI governance space - Cognitive surrender: AI adoption without structured critical engagement degrades professional judgement - Regulatory fragmentation may exceed capacity of any single profession - Speed of AI capability development may outpace adaptation of curricula and standards

The Profession's Choice

ACTUARIES ARE NOT STARTING FROM ZERO

01 Model literacy — trained in construction, validation, uncertainty

02 Regulatory fluency — embedded in Solvency II, ORSA, EIOPA frameworks

03 A signed, auditable professional opinion already in the architecture

THREE CONDITIONS THAT MUST BE MET

01

A healthy relationship with AI

Structured critical engagement, not passive adoption. The cognitive surrender risk falls disproportionately on professions that use AI at scale without professional obligation to verify.

02

Structural openness of the profession

Academic curricula, professional examinations, and continuing education must incorporate AI as a cross-cutting competency — not as an optional module.

03

Continuous exchange

AI governance is not a destination. The regulatory landscape, the technology, and the risk environment are all moving simultaneously. Congresses like this one are part of the governance infrastructure.

The actuarial profession has governed uncertainty under regulation for over a century. AI governance is uncertainty under regulation. The question is not whether actuaries can do this. The question is whether the profession chooses to and when.