



SIFM AI Roundtable

Elaine Lehnert, Ankura Consulting Group

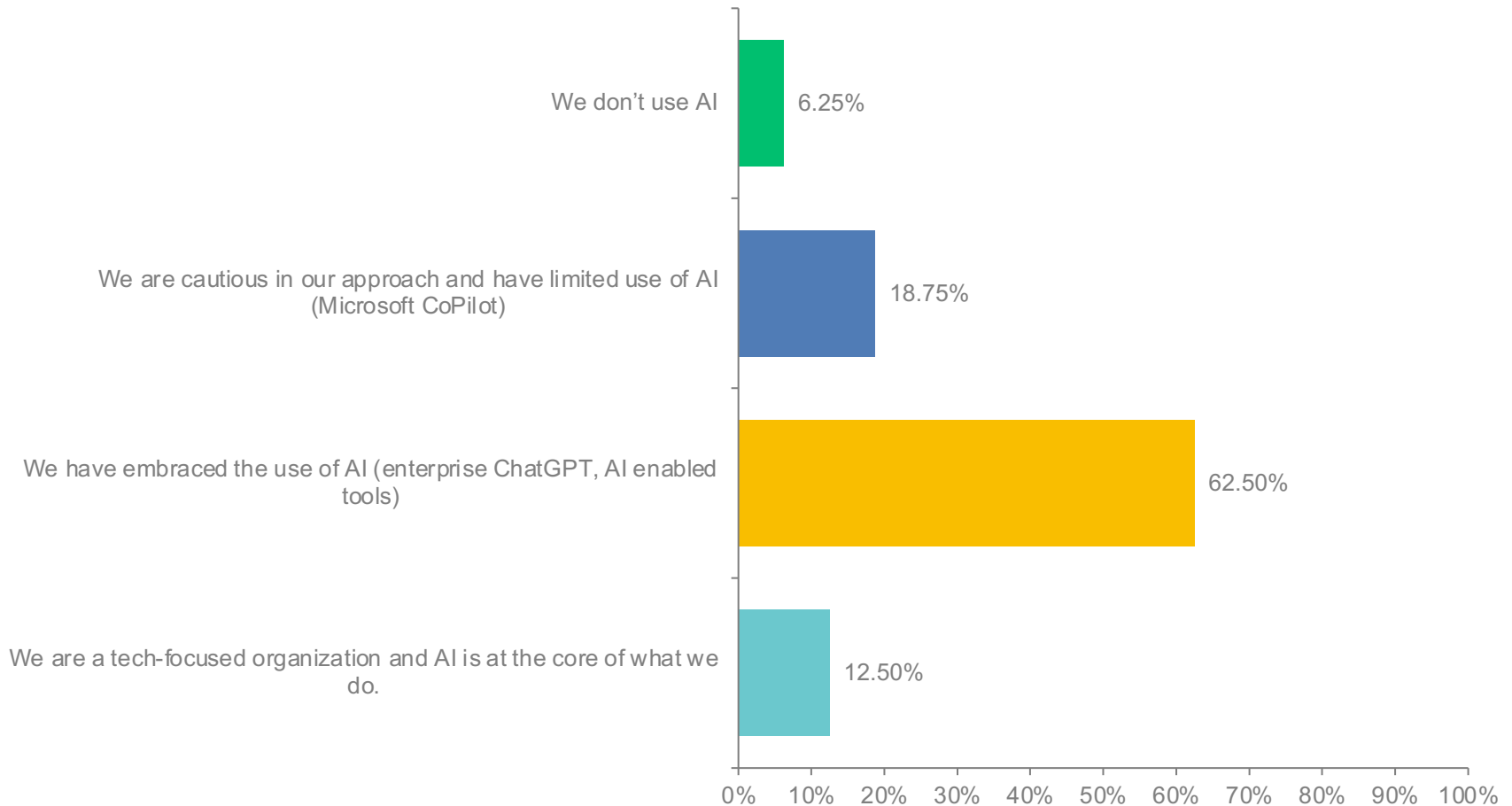
Julia Paromchik, QBE Insurance NA

Russell Sommers, Baker Tilly Advisory Group

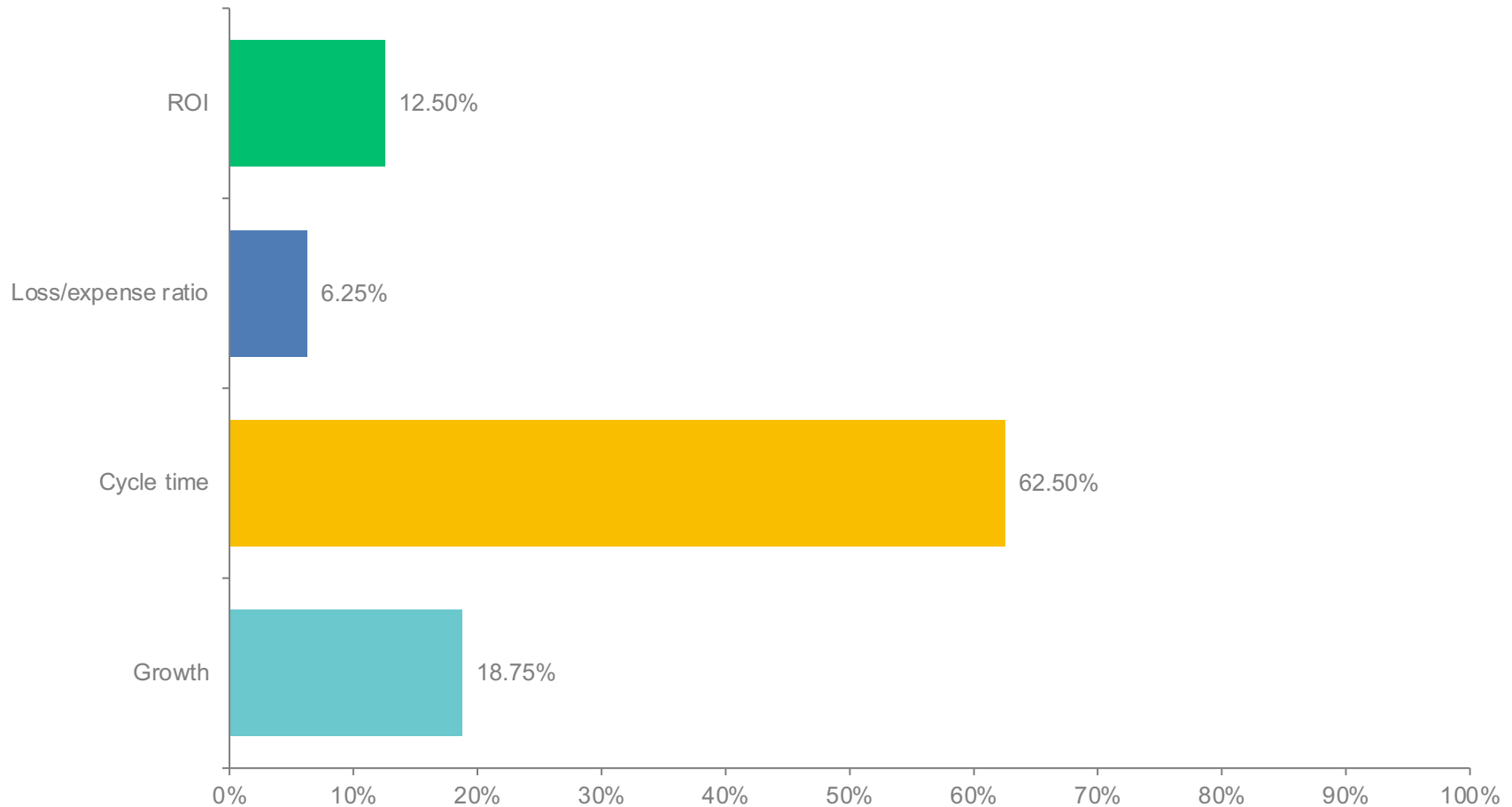
June 24, 2026

SIFM Quarterly Conference

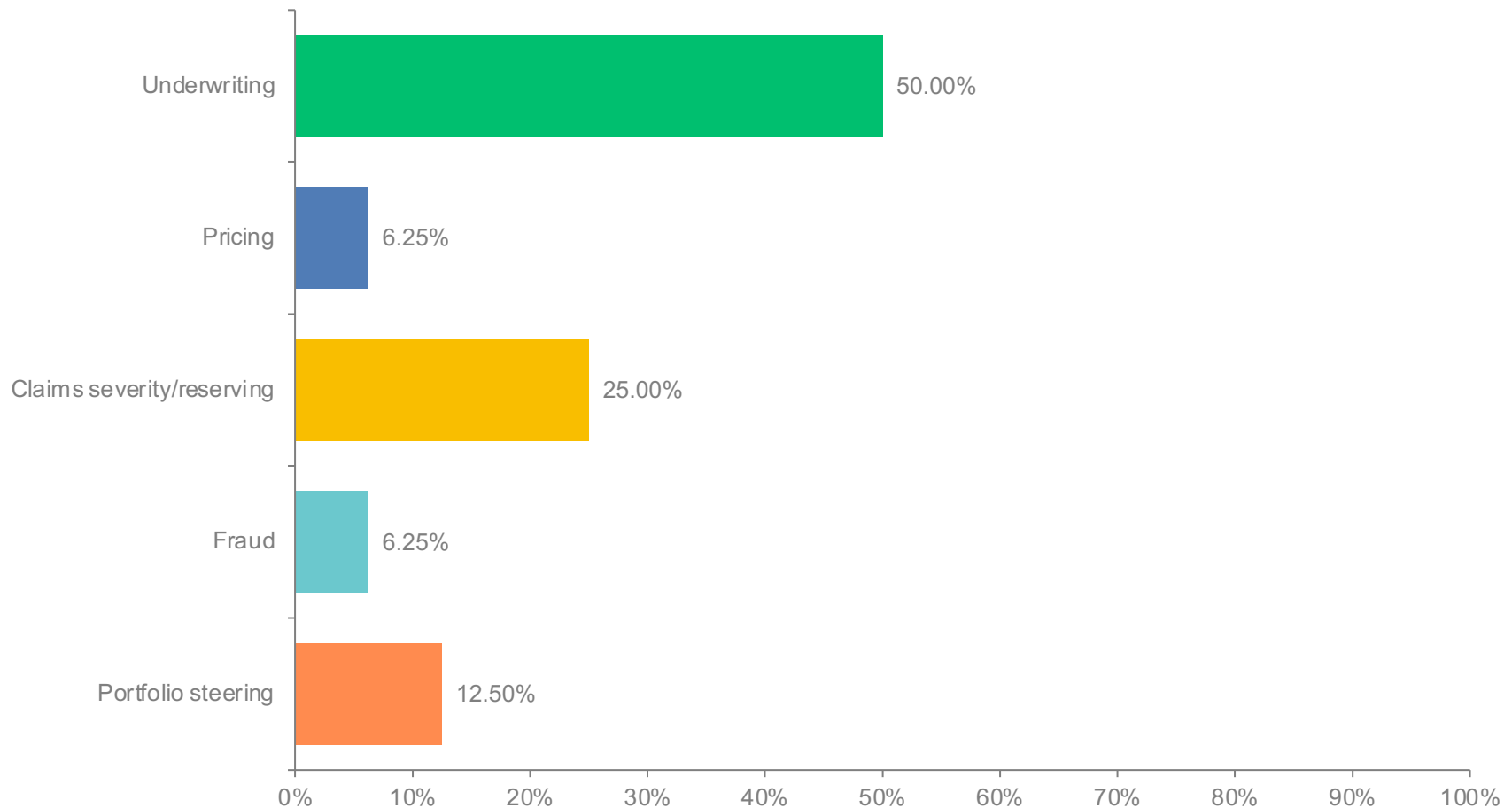
Which best describes your organization's current AI maturity?



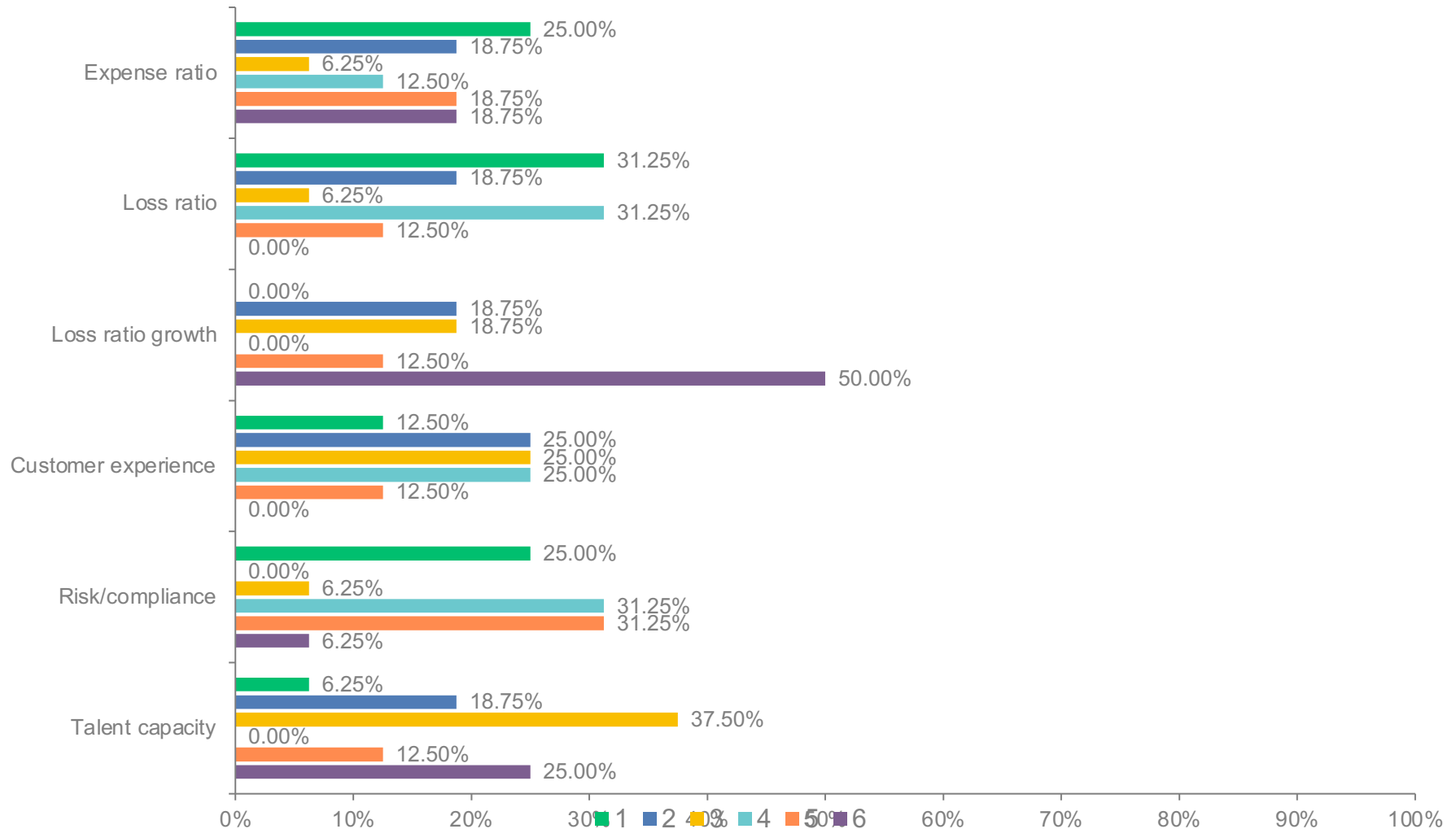
Which AI use cases have the most tangible business value?



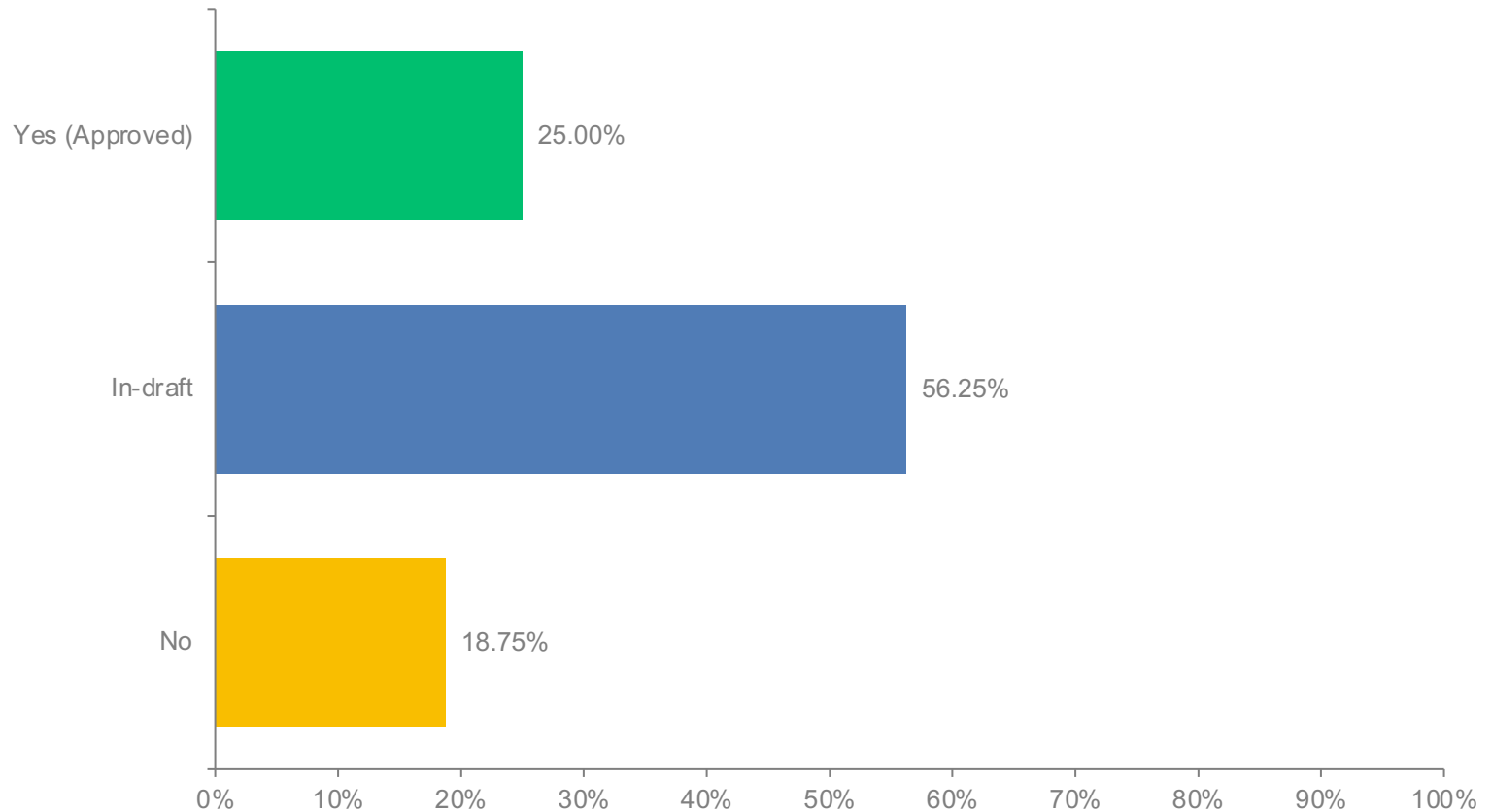
In which department has AI improved the quality of decisions, not just efficiency?

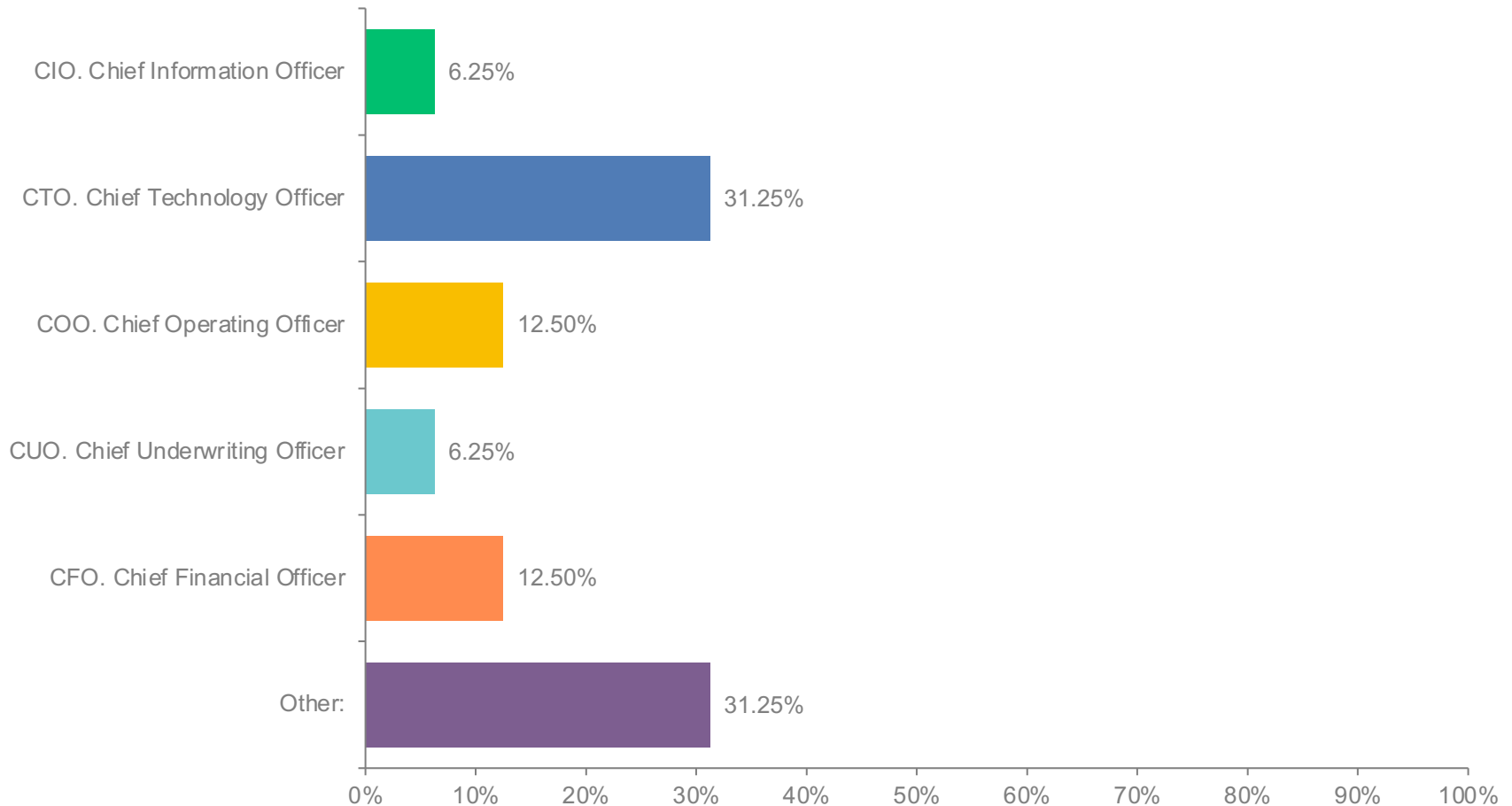


At leadership level, what outcomes really matter most when judging AI success?

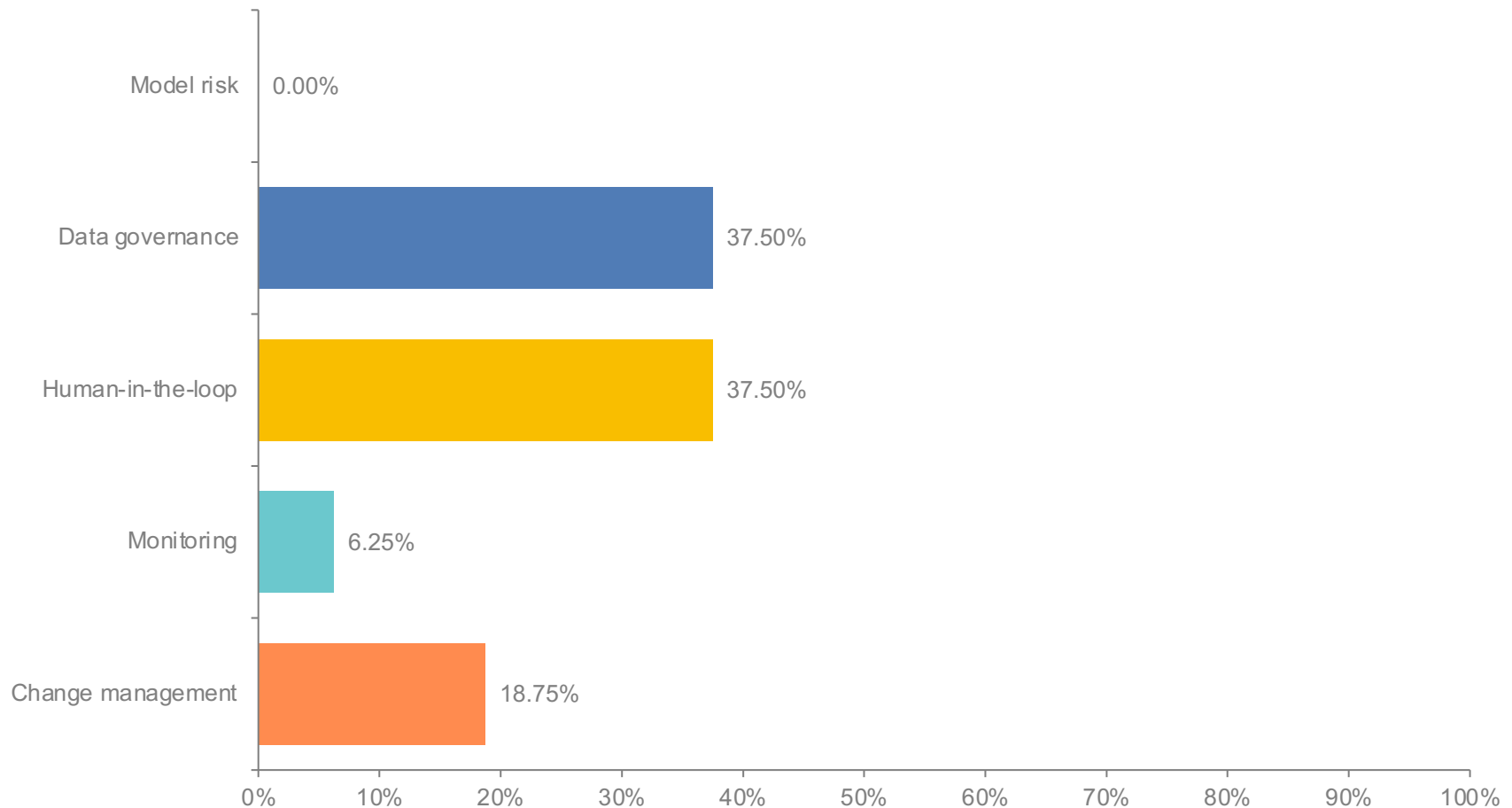


Do you have a formal enterprise AI strategy/roadmap?

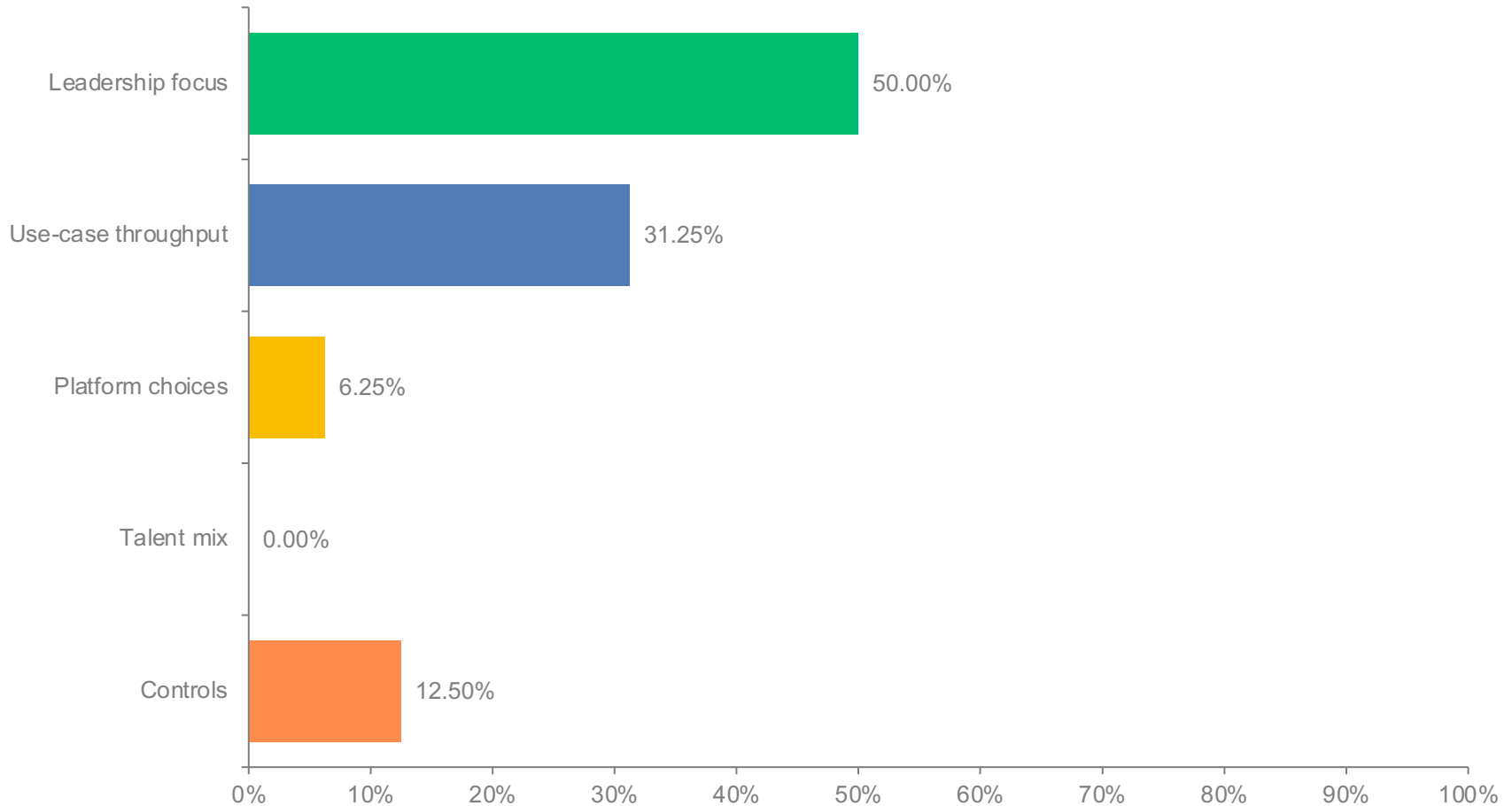


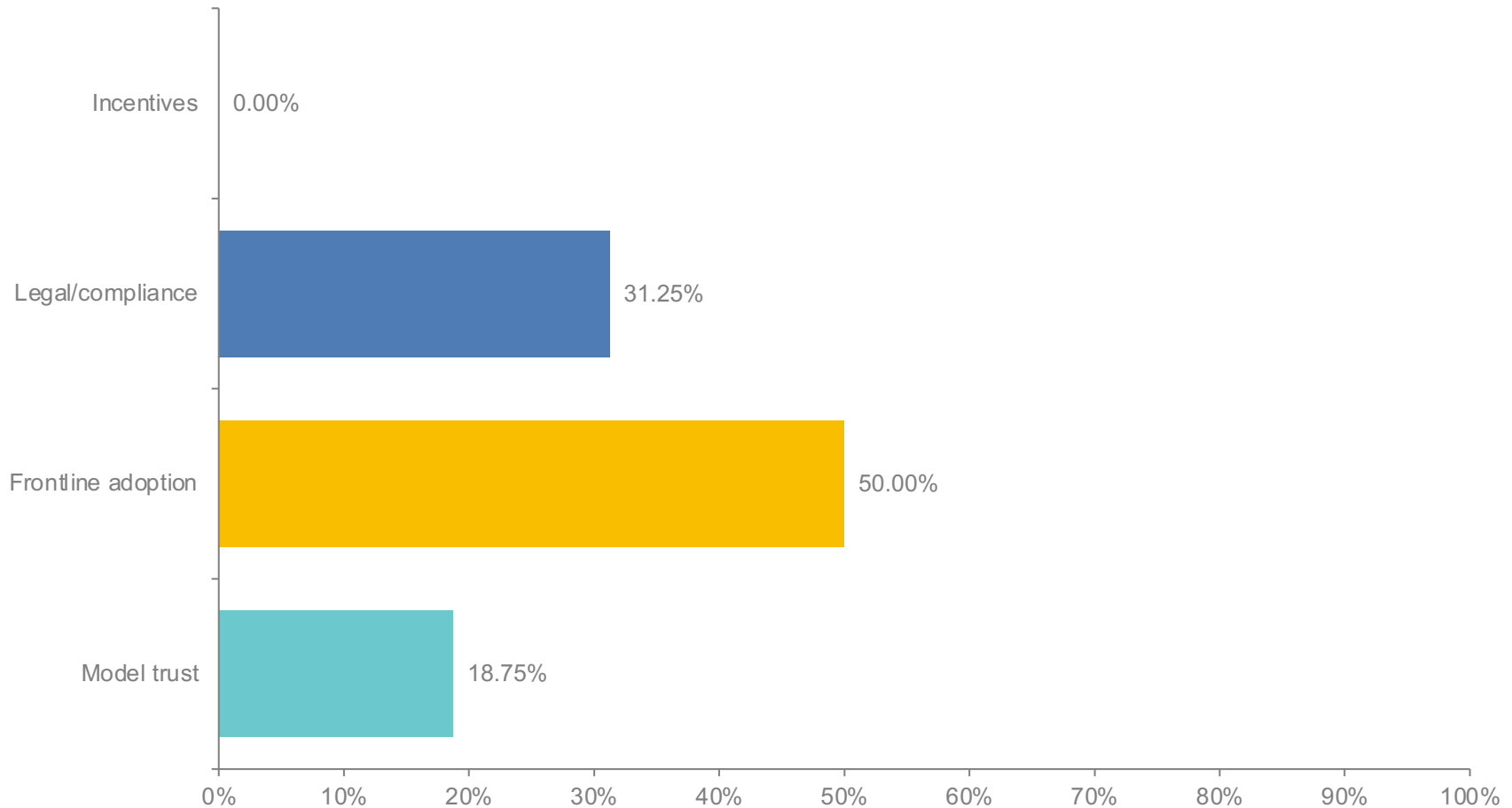


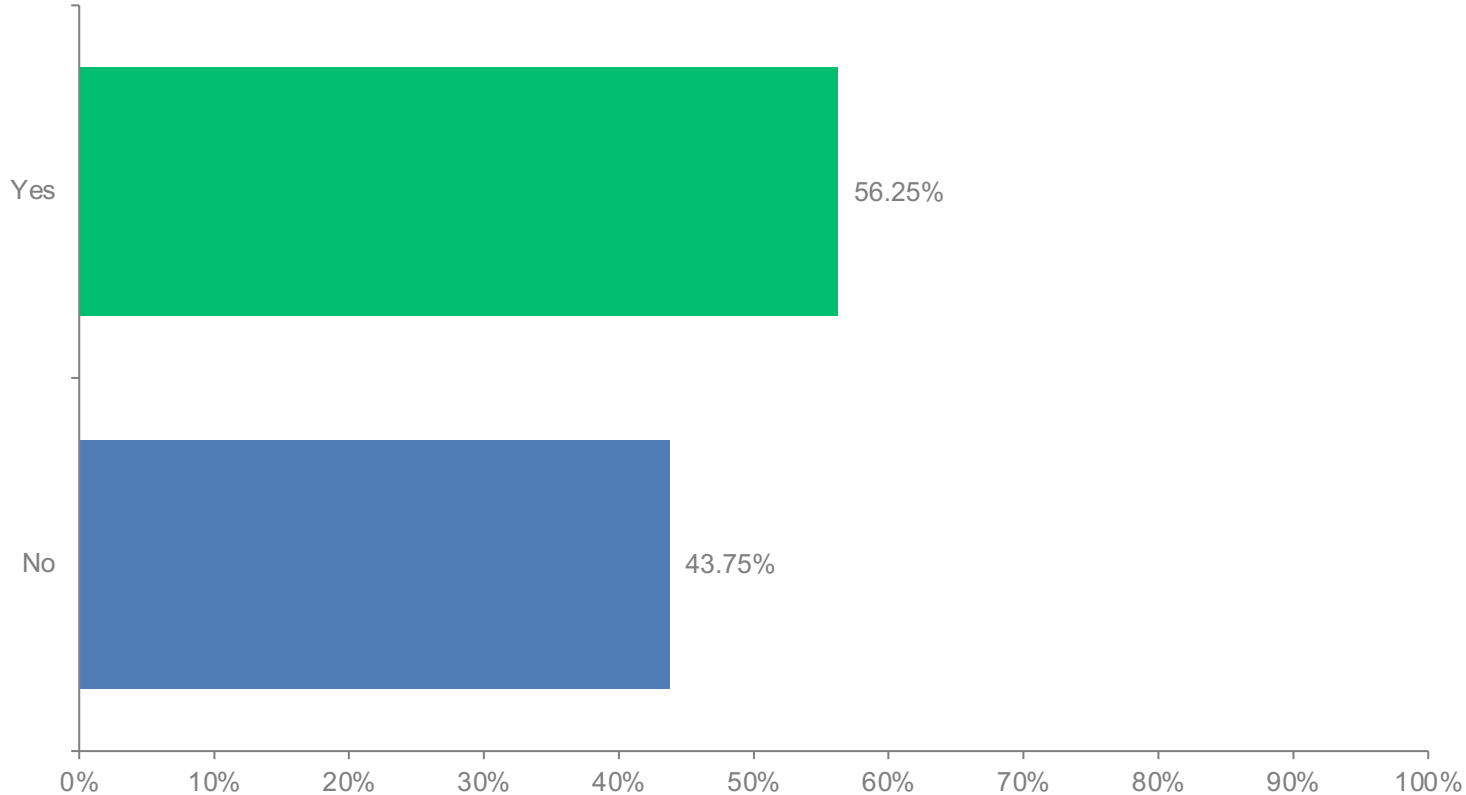
What governance approaches have enabled innovation without slowing progress?



Looking across the industry, what are the early signs that an insurer is falling behind on AI maturity?







1. Defining AI Boundaries

- Emergence of “Shadow AI”
- Need clear definitions of approved AI use vs tools

2. AI Inventory & Classification

- General Productivity AI
(*writing assistants, summarization tools*)
- Decision-support AI
(*claims, UW, financial forecast*)
- Autonomous / High Impact AI
(*auto decision making*)
- Level of oversight will vary

3. Accuracy & Reliability Expectations

- Exploratory / productivity use → 90–97% may be acceptable
- Financial reporting / regulatory outputs → near 100% required

4. Change Management Evolution

- AI introduces continuous change vs static releases
- Shift to continuous monitoring, re-testing and drift detection

5. Explainability & Model Risk

- Persistent “black box” limitations
- Outputs may be non-deterministic and difficult to explain

6. Monitoring & Accountability

- “Who watches the AI?”
- Emerging concept of AI monitoring AI
- Authority may be delegated but accountability remains with humans

7. Human Transformation

- Employees gain capacity but must shift to more strategic work
- Tasks vs judgement-based responsibilities need to be rebalanced
- Upskilling vs de-skilling workforce

8. ICFR Impact

- Challenges validating AI driven outputs
- Increased need for controls, auditability, and validation evidence

NAIC Model Bulletin on Use of AI Systems by Insurers

- Establishes board accountability, AI inventory requirements, third-party oversight, and governance program expectations. Adopted by 20+ states.

NY DFS · CL7 (2024) Circular Letter 7 — AI in Underwriting & Pricing

- Requires written policies, comprehensive AI documentation inventory, bias testing (proxy assessment, adverse impact ratios), third-party contractual audit rights, and adverse action notice requirements.

AI Risk Management Framework (RMF) (2023)

- Voluntary, non-authoritative but widely adopted across industries. Four core functions: GOVERN, MAP, MEASURE, MANAGE. Provides a technology-neutral, lifecycle approach to AI risk.

COSO Generative AI Risk Framework

- Applies COSO internal control principles (control environment, risk assessment, control activities, information & communication, monitoring) specifically to generative AI. Bridges AI risk to familiar ERM and internal audit frameworks already used in examinations.

**Principle 1: Accountability**

Insurers and all AI actors are responsible and accountable for decisions made with AI. Human oversight of AI decision-making must be maintained — examiners assess whether governance structures enforce this.

**Principle 2: Compliance**

All applicable insurance laws and regulations must be followed when using AI. An AI system does not excuse non-compliance with underwriting, rating, or consumer protection requirements.

**Principle 3: Fairness & Non-Discrimination**

AI systems must not engage in unfair discrimination. Bias testing, proxy assessments, and adverse impact analysis are required controls. NY CL7 operationalizes this into specific quantitative testing methodologies.

**Principle 4: Transparency**

- Insurers should be transparent with regulators and consumers about their use of AI and the basis for AI-driven decisions. Consumers must be able to understand and contest adverse decisions.

**Principle 5: Security, Safety & Robustness**

AI systems must be secure, safe, and resilient. This is the direct linkage to Model 668 cybersecurity requirements and Exhibit C IT controls — AI systems are information systems requiring full security program coverage.

- **AI/model risk management policy and procedures**
- **Comprehensive AI system inventory**
- **Board and committee meeting minutes**
- **Organizational chart showing AI governance**
- **AI-related training program**
- **Internal audit reports and findings**
- **ORSA summary — AI/technology risk AI governance committee**

- **Model validation reports**
- **Bias/fairness testing reports**
- **Model monitoring reports**
- **Data governance documentation**
- **Vendor AI due diligence**
- **Cybersecurity risk assessment**
- **Incident response plan**
- **Adverse action notice samples (NY CL7)**

Take 7-10 minutes within your tables and discuss one of the following questions:

Biggest Risk

- What is the greatest AI risk your organization is dealing with today?

What's Working

- Where is AI delivering clear, tangible value right now?

What's Not Working

- Where has AI fallen short or caused issues?

Coollest Use Case

- What is the most innovative or impactful AI use you've seen?

Governance vs. Agility

- How are you enabling AI without compromising controls?

Decision Integrity

- How do you ensure AI is informing—not distorting—decisions?