

Insights into AI's Role in Financial Services

[View the X9 AI Study Group](#)

X9 AI Study Group Newsletter - Issue No. 12 (2026)

Tacit Knowledge Is Your Next Competitive Moat

Organizations are increasingly recognizing that one of the biggest challenges in scaling AI is capturing tacit knowledge, the judgment, experience, and informal know-how that employees develop over years of work. Much of this expertise is never documented in procedures or manuals, yet it often drives critical decisions and operational effectiveness. As organizations adopt generative and agentic AI, attention is shifting toward identifying, preserving, and operationalizing employee expertise before it is lost through retirement or workforce turnover. The article argues that organizations that successfully capture and apply this institutional knowledge may gain a durable competitive advantage that cannot be replicated through technology investments alone.

[Continue Reading](#)



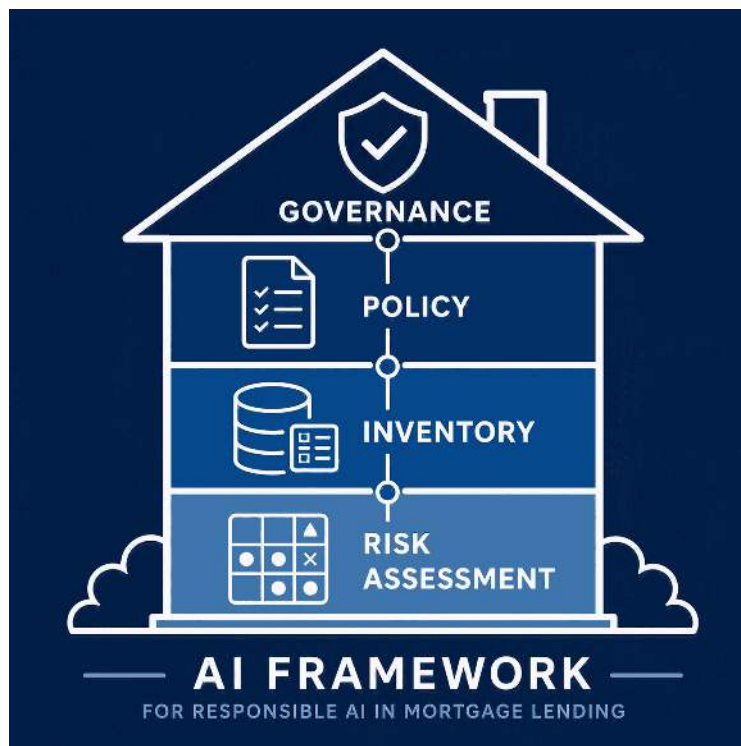
Mortgage Industry Introduces AI Governance Toolkit

MISMO has released the Framework for Responsible AI in the Mortgage Ecosystem (FRAME), a governance toolkit designed to help mortgage lenders, servicers, and technology providers manage AI risks while supporting responsible innovation.

Developed by MISMO's AI Community of Practice at the request of the Mortgage Bankers Association's Residential Board of Governors, FRAME provides mortgage-specific guidance aligned with broader AI governance standards. It includes a governance policy template, AI system inventory, AI risk assessment, implementation guidance, and a getting-started guide. The framework is intended to help organizations identify where AI is being used, assess associated risks, document decision-making, and establish repeatable governance processes.

MISMO emphasizes that FRAME is a practical risk management tool, not a new regulatory requirement, and is designed to be scalable for organizations at any stage of AI adoption.

[Continue Reading](#)



G7 Leaders Debate AI Sovereignty and Access to Frontier Models

Artificial intelligence emerged as a major topic at the G7 summit, where government leaders and AI executives discussed how advanced AI systems should be governed. French President Emmanuel Macron and OpenAI CEO Sam Altman called for greater international

cooperation and the creation of global mechanisms to develop AI safety standards and oversight.

The discussion was fueled by a recent U.S. decision restricting foreign access to Anthropic's newest AI models. European leaders expressed concern that heavy reliance on a small number of U.S. AI providers could leave countries vulnerable if access to critical AI capabilities is limited or withdrawn.

The debate highlights a growing global tension between national security, responsible AI governance, and technological sovereignty as nations seek both access to advanced AI and greater control over their digital futures.

[Continue Reading](#)



Fed Chair Warsh Launches Five Task Forces to Reshape Central Bank

Federal Reserve Chair Kevin Warsh has begun what observers describe as a quiet but potentially far-reaching effort to reshape the central bank. Following his first policy meeting as chair, Warsh announced five task forces that will review key aspects of monetary policymaking, including Fed communications, economic data, inflation frameworks, the impact of technologies such as AI, and management of the Fed's \$6.7 trillion balance sheet.

[cnbc.com], [cnbc.com]

Warsh said the reviews will start from first principles and consider alternative approaches to existing practices. The initiative signals a willingness to rethink how the Fed measures the economy, communicates policy, and evaluates long-term economic trends. While Warsh previously criticized the institution, he is pursuing change through a collaborative process that many analysts view as a significant, though measured, shift in the Fed's direction

[Continue Reading](#)



A blueprint for Democratic Governance of Frontier AI

OpenAI's Blueprint for Democratic Governance of Frontier AI argues that increasingly capable AI systems require formal government oversight rather than relying solely on voluntary industry commitments. The proposal centers on three pillars: a national framework for frontier AI safety, a strengthened Center for AI Standards and Innovation (CAISI) as the federal government's primary evaluation body, and a broader resilience strategy to address emerging national security and public safety risks.

The paper argues that democratic governments are best positioned to balance innovation with accountability, transparency, and public oversight. It calls for ongoing capability monitoring, independent evaluations, and greater government access to technical expertise so policymakers can respond as AI capabilities evolve. The central message is that decisions about the most powerful AI systems should ultimately be shaped through democratic institutions rather than by AI developers alone.

[Continue Reading](#)



AI-Driven Fraud Losses Surge Globally

AI ranks as the top fraud threat facing financial institutions in 2026. Fraudsters use generative tools to produce convincing content and automate attacks at scale. You face faster, cheaper, and more credible scams than a year ago.

The data shows the scale. Consumers lost more than \$12.5 billion to fraud in 2024. Voice cloning and vishing attacks now exceed 1,000 AI scam calls per day at major retailers. Deepfakes account for a rising share of global fraud activity.

The methods keep shifting. Emotionally intelligent bots run romance and relative-in-need scams with no human operator. Cloned websites harvest credentials and feed identity theft. Your defenses need AI-based detection to match the speed of these attacks.

[Continue Reading](#)



From Breakthrough to Regulation in One News Cycle:

How Anthropic's newest models helped accelerate the debate over AI governance

Breakthrough: Anthropic Brings Frontier AI to the Mainstream

Anthropic launched Claude Fable 5 and Claude Mythos 5, its most capable AI models to date. Fable 5 is a public-facing version of the more powerful Mythos platform, delivering major advances in software engineering, research, vision, and sustained multi-step reasoning. To reduce misuse risks, Anthropic implemented safeguards that redirect certain cybersecurity and biology-related requests to a less capable model. The launch marked one of the first large-scale deployments of a frontier AI system designed to balance cutting-edge capability with safety controls, signaling a new phase in the commercial availability of advanced AI.

[Continue Reading](#)



Regulation: Government Halts and Then Restores Access

Just three days after Anthropic launched Claude Fable 5 and Mythos 5, the U.S. government imposed export controls that forced the company to suspend access to both models worldwide. The action reflected growing concern that frontier AI systems with advanced cybersecurity capabilities could pose national security risks if misused. After several weeks of review, the controls were lifted and access was restored with additional safeguards and closer government coordination. The episode marked one of the first major cases of direct government intervention in the deployment of a frontier AI model.

[Continue Reading](#)

Stay Ahead in Financial AI

Subscribe to our newsletter for the latest insights and updates on AI in financial services.

[Subscribe Now](#)