



FINANCIAL
STABILITY
BOARD

Sound Practices for Responsible Adoption of Artificial Intelligence (AI): Consultation report

Response to Consultation

OATHOR Ltd

1. Do you agree with the benefits and risks of AI adoption by financial institutions described in this report? Are there any substantive benefits or risks not covered?

Yes. OATHOR LTD broadly agrees with the benefits and risks described in the report. AI adoption can improve efficiency, risk detection, customer service, fraud monitoring, compliance workflows, cyber defence, data analysis and decision support across financial institutions.

One risk that could be made more explicit is execution-boundary risk.

Execution-boundary risk arises when an AI-enabled output, instruction, recommendation, workflow or system action becomes operational, financial, legal, customer, market or institutional consequence.

This risk becomes more important as AI systems move from analysis and decision support toward workflow initiation, tool use, autonomous coordination, transaction-related processes, customer-affecting actions, access changes and operational decisioning.

Model governance, data quality, explainability, cybersecurity, monitoring and audit remain necessary. They may not be sufficient where consequential digital action can occur before an independent control is applied.

The report could be strengthened by recognising execution-boundary risk as a distinct concern within responsible AI adoption, especially for high-consequence, autonomous, semi-autonomous and agentic AI environments.

2. Are the sound practices sufficiently comprehensive and clear to enable financial institutions' responsible AI adoption?

The sound practices are broadly comprehensive and clear. They address governance, accountability, senior management oversight, lifecycle management, risk assessment, data governance, explainability, performance management, cyber and ICT risk, third-party risk and human oversight.

One area where additional clarity may be helpful is the distinction between different control functions in the AI governance stack.

Policy defines obligations. Cybersecurity protects access. Monitoring observes behaviour. Audit records occurrence. Where AI-enabled systems can influence or initiate consequential actions, responsible adoption may also require control at the point before execution becomes consequence.

The sound practices would be strengthened by clarifying that accountability after an event is necessary, but not always sufficient. In high-consequence settings, institutions should assess whether control exists where execution is still preventable.

This would not require a prescriptive technical standard. It would clarify a principle. The closer an AI system is to consequential execution, the stronger the need for independent pre-execution control.

3. Do the sound practices strike an appropriate balance between managing risks relating to all forms of AI, and addressing some of the risks relating to emerging and new complex forms of AI, such as GenAI and agentic AI?

The sound practices strike a generally appropriate balance. A technology-neutral and principles-based approach is important because financial institutions use many forms of AI, from traditional machine learning to GenAI and agentic AI.

GenAI and agentic AI make execution proximity more important. The risk is not only that an AI model may produce an inaccurate, biased or poorly explained output. The risk is also that an AI-enabled system may initiate, route, approve, modify or coordinate actions across tools and workflows before appropriate control is applied.

The report could more clearly distinguish between AI that supports analysis, AI that recommends action, AI that initiates action, AI that coordinates actions across systems and AI that can produce consequential operational effects.

This distinction would help financial institutions apply proportionate controls. AI systems with low execution proximity may require conventional governance, monitoring and review. AI systems closer to consequential execution may require stronger independent controls before action proceeds.

4. Are the sound practices sufficiently flexible to accommodate and address newer types of AI and responsible adoption over time?

Yes. The sound practices are sufficiently flexible because they are principles-based, proportionate and not tied to one AI architecture.

This flexibility could be strengthened by adding an execution-proximity lens.

A practical governance question would be:

How close is the AI system to causing, initiating, approving, modifying, preventing or coordinating a consequential operational, financial, legal, customer, market or institutional action?

This lens remains adaptable over time. Whether a system is traditional AI, GenAI, agentic AI, multi-agent AI, embedded AI or a future autonomous digital system, the governance concern remains materially similar. Does the system influence consequence, and is appropriate control present before consequence begins?

Execution proximity would also help institutions avoid assessing AI risk only by model complexity or explainability. A simpler system close to consequential execution may present greater institutional risk than a more complex system used only for internal analysis.

5. Do the case studies in this report sufficiently highlight how different types of financial institutions can benefit from responsible AI adoption? Are there additional case studies for inclusion in the report? If so, please provide such case studies, particularly for nonbanks.

The case studies are useful and help illustrate responsible AI adoption across different financial institution contexts. An additional case study could focus on a nonbank financial institution or financial technology provider using AI-enabled workflow automation in a customer or transaction-affecting environment.

For example, a nonbank financial institution may use AI-enabled systems to support onboarding, fraud review, transaction monitoring, customer access decisions, account actions, escalation workflows or operational decisioning. The AI system may not merely produce analysis. It may influence workflow routing, escalation, approval, suspension, review or downstream processing.

In this scenario, responsible AI adoption should not rely only on model validation, access control, monitoring and downstream audit. The institution should identify which AI-enabled outputs can become consequential actions, define which actions require pre-execution control, and ensure that authority, policy alignment, risk conditions, human oversight requirements and exception-handling rules are verified before consequence occurs.

This case study would be useful for nonbanks because many nonbank environments rely on technology-enabled workflows, third-party systems and operational automation where AI may influence action before traditional review mechanisms are triggered.

6. Do the case studies in this report provide actionable insights for financial institutions in their responsible AI adoption?

The case studies provide useful actionable insights, particularly in relation to governance, proportionality, oversight, testing, explainability, risk management and responsible deployment.

Future case studies could become more actionable by focusing more directly on the transition from AI output to institutional action.

In practice, many risks do not arise only when an AI model produces an output. They arise when that output is accepted, routed, escalated, automated or executed within an operational environment.

Additional guidance could ask financial institutions to identify which AI outputs can influence consequential actions, which actions require independent control before execution, where human oversight must occur before execution rather than after, which actions should never be executed automatically, how exceptions are handled when an AI-enabled workflow exceeds authorised boundaries, and how accountability is preserved when AI systems interact with third-party tools, internal infrastructure or customer-facing systems.

This would help institutions translate responsible AI principles into operational controls.

7. Are the definitions in the glossary clear and aligned with industry sound practices, including recent developments in AI?

The glossary definitions are generally clear and aligned with current industry practice. Additional clarity could be useful around execution-related concepts, particularly as agentic AI and autonomous workflows become more relevant to financial institutions.

Execution-boundary risk could be defined as the risk that an AI-enabled output, instruction, recommendation, workflow or system action becomes operational, financial, legal, customer, market or institutional consequence without appropriate control being applied before execution.

Execution proximity could be defined as the degree to which an AI system can influence, initiate, approve, modify, prevent, coordinate or execute a consequential action.

Pre-execution control could be defined as control applied before a consequential digital action proceeds, particularly where downstream audit or post-event accountability may be insufficient to prevent harm or institutional exposure.

The definition of agentic AI could also clarify that risk increases where AI systems can plan, sequence, initiate or coordinate actions across tools, systems or workflows, especially where those actions can create institutional consequence.

These additions would help institutions distinguish between AI systems that inform decisions and AI systems that influence or cause consequential execution.

8. Are there any comments you would like to make on this report? Please fill in.

OATHOR LTD welcomes the FSB's consultation on sound practices for responsible adoption of artificial intelligence by financial institutions.

OATHOR LTD is an ADGM-incorporated technology infrastructure company based in Abu Dhabi, United Arab Emirates, focused on Execution Control Infrastructure for high-consequence digital and autonomous environments (oathor.com).

This submission is provided as a non-confidential, category-level contribution. It does not disclose proprietary mechanisms, seek product recognition or request endorsement. Its

purpose is to highlight a structural issue relevant to responsible AI adoption: execution-boundary risk.

As AI-enabled systems become more capable of influencing instructions, approvals, transactions, workflows, access changes, operational decisions and autonomous agent activity, responsible AI governance must extend beyond upstream policy, perimeter cybersecurity, model governance, monitoring and downstream audit.

Those layers remain necessary. They are not always sufficient.

The core issue is that digital systems can act at institutional speed. Where consequence is automated or semi-automated, control must exist where execution is still preventable.

OATHOR LTD recommends that the FSB consider explicitly recognising execution-boundary risk within the sound practices. Financial institutions should be encouraged to assess AI systems not only by model type, complexity, explainability or use case, but also by proximity to consequential execution.

The relevant question is:

How close is the AI system to causing, initiating, approving, modifying, preventing, coordinating or executing an operational, financial, legal, customer, market or institutional consequence?

The closer an AI system is to consequential execution, the stronger the case for independent pre-execution control.

A relevant control category for addressing this risk is independent execution control. OATHOR refers to this category as Execution Control Infrastructure: an independent control layer at the point where AI-enabled digital action becomes operational, financial, legal, customer, market or institutional consequence.

The distinction is structural.

Policy defines obligation.

Cybersecurity protects access.

Audit records occurrence.

Execution control governs consequence.

This distinction may be especially important for GenAI, agentic AI, third-party AI tools, AI-enabled workflows and systems that interact with operational infrastructure. In such environments, the issue is not only whether an AI model is accurate, explainable or monitored. The issue is whether consequential execution is independently governed before consequence begins.

Recognising execution-boundary risk would strengthen the sound practices while remaining proportionate, technology-neutral and adaptable to future AI development.