



FINANCIAL
STABILITY
BOARD

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

Response to Consultation

Duška Ž (independent response)

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.

The report provides a balanced overview of both the opportunities and risks associated with AI adoption in financial services.

One area that could be further emphasized is the relationship between AI autonomy and decision materiality. The primary governance challenge is often not the model itself, but the degree of autonomy granted to the model relative to the potential impact of failure.

As decision materiality increases, acceptable AI autonomy should decrease and independent validation mechanisms should become more prominent.

In addition, cumulative risk deserves explicit consideration. Large volumes of individually low-impact AI decisions may collectively create material financial, operational or conduct risks if a common error, bias or control failure exists.

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

Yes.

The proposed practices provide a strong foundation for responsible AI adoption.

Several practices could be further strengthened through additional guidance on:

- proportional calibration of AI autonomy based on decision materiality;
- aggregate risk monitoring across large populations of low-impact decisions;
- decision traceability and explainability at the individual decision level;
- governance of AI orchestration architectures involving multiple models, agents and control layers;

- the portability and independent auditability of core business logic when deploying proprietary third-party AI solutions to mitigate concentration and lock-in risks.

Increasingly, financial institutions are deploying AI systems as interconnected decision ecosystems rather than standalone models. Governance expectations for orchestration and control architecture may therefore warrant additional attention.

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?

Yes.

The principles appear sufficiently technology-neutral and therefore should remain relevant as AI capabilities evolve.

The proposed framework appropriately focuses on governance, accountability, risk management and oversight rather than specific technologies.

For emerging agentic AI systems, additional emphasis could be placed on:

- clearly defined decision boundaries;
- escalation mechanisms;
- role separation between execution, validation and monitoring functions;
- controls governing interactions between multiple AI agents.

These considerations become increasingly important as institutions move from model-centric deployments toward orchestrated AI systems capable of initiating and executing actions.

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

The case studies are useful and illustrate several practical applications.

Additional examples could include:

- AI-enabled credit decisioning with independent validation controls;
- AI-supported fraud detection operating under predefined escalation rules;
- multi-agent operational workflows where decision execution, validation, monitoring and escalation are distributed across separate control functions.

Such examples would help demonstrate how responsible AI adoption can improve operational efficiency while maintaining strong governance and control standards.

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.

Yes.

However, additional examples illustrating control architecture design would further increase practical value.

Many implementation failures originate not from model performance deficiencies but from weaknesses in governance architecture, escalation design, accountability structures or decision controls.

Examples showing how institutions operationalize fail-safe mechanisms, escalation pathways and control responses when predefined materiality thresholds are breached would be particularly beneficial.

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

The case studies provide useful illustrations of AI-related risks and demonstrate the importance of governance, oversight and risk management.

However, their practical value could be strengthened by providing greater visibility into the control architectures surrounding the AI systems involved.

Many significant AI failures originate not from model deficiencies alone, but from weaknesses in system design, decision governance, escalation mechanisms, control frameworks or responsibility allocation.

Additional value could be provided by including examples that demonstrate:

- how decision authority was allocated between humans and AI systems;
- how escalation mechanisms were designed and triggered;
- how materiality assessments influenced the level of autonomy granted to the system;
- how institutions implemented fail-safe controls when predefined thresholds were breached;
- how organizations recalibrated models, controls and decision frameworks following incidents;
- how multiple AI models, agents or control layers interacted within a broader decision architecture.

As AI adoption increasingly evolves from individual models toward interconnected decision ecosystems, practical examples of orchestration, control design and governance architecture would provide financial institutions with more actionable implementation guidance than model-level examples alone.

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

Generally yes.

As AI adoption evolves, consideration could be given to expanding definitions relating to:

- AI autonomy;
- AI orchestration;
- agentic systems;
- human oversight;
- decision explainability;
- control architecture.

Clear terminology in these areas may improve consistency of interpretation across institutions and jurisdictions.

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

The report appropriately emphasizes governance, risk management and oversight.

One observation is that responsible AI adoption increasingly depends on the design of the overall decision architecture rather than on the performance of any individual model.

In practice, many operational and governance failures emerge when a single AI system is expected to observe, reason, decide, execute and validate its own actions.

As AI adoption matures, institutions may benefit from architectures that separate these responsibilities across independent control functions, whether implemented through AI systems, deterministic controls or human oversight.

The effectiveness of AI governance is therefore not determined solely by model quality, but also by how decision authority, validation, escalation and accountability are distributed throughout the system.