



FINANCIAL
STABILITY
BOARD

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

Response to Consultation

ÆQUOM - Fair Ops & AI Governance

- 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?**

Broadly agree with the benefits and risks as described. The one risk dimension I would weight more heavily is the verifiability of oversight claims at scale: the gap between oversight asserted and oversight evidenced, which my responses to questions 2-4 address.

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

I respond as a compliance and GRC practitioner with 15+ years across Financial Services, Tech and Telecom, including compliance operations supporting certification work under a U.S. Department of Justice Deferred Prosecution Agreement (DPA). My work is the operating layer these practices address: compliance operating models, control libraries and testing scripts, evidence standards and sampling rules, monitoring and self-testing, and remediation through RCSA and the Three Lines of Defence, in recent years applied to AI-enabled processes, aligned to the EU AI Act where applicable.

The practices are comprehensive on what institutions should do. Where they could be one degree more explicit is the evidence layer. In operation, the recurring failure is not the absence of oversight, but the absence of evidence that oversight occurred. 'Reviewed' is asserted; when tested, it is undefined: no named reviewer, no criteria the review was performed against, no record that would satisfy an internal auditor, let alone a supervisor. I would encourage the final report to make explicit that institutions should define, in advance, what constitutes evidence of oversight for each AI use case: who reviews, against what criteria, recorded where, and how that record is sampled and tested. This is standard controls discipline, evidence standards and sampling rules, applied to a new risk class.

- 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?**

As autonomy increases toward agentic systems, the volume and speed of actions make oversight claims progressively harder to verify after the fact. An evidence standard defined per use case scales across the whole oversight spectrum, because it attaches to the

oversight design rather than to any one technology. That is what allows the practices to manage risk across all forms of AI while still addressing what is genuinely new in GenAI and agentic systems: the verification burden moves with the design, not the model.

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

The oversight model for each use case should be a documented design decision. The report helpfully maps oversight along a spectrum. One further step would strengthen it: institutions should record, per use case, which point on that spectrum was chosen, why, by whom, and under what conditions the choice is revisited, for instance, when the model, the data, or the materiality of the use case changes. In practice, oversight levels are rarely chosen; they accrete. A process launches with full human review, volume grows, review quietly thins, and no one can later show that the resulting oversight level was a decision rather than a drift. A recorded design decision with a named owner converts that drift into something governable, and the revisit conditions are what keep the practices flexible as newer types of AI arrive, without prescribing any particular model.

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.

I do not offer additional case studies. On how the existing ones could serve different institution types, particularly non-banks, more effectively, see my response to question 6.

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

The case studies would gain actionability by showing, for at least some examples, the evidence artefacts behind the practice: what the oversight record, the monitoring output, or the escalation log actually looked like, rather than the process description alone. That is where smaller institutions and non-banks most need a concrete anchor.

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

I would suggest defining 'human oversight' by the decision rights it carries: the authority to approve, to bound, to monitor, and to stop, rather than by the presence of a human in the process. Without that, the same term covers genuine command and rubber-stamping, and this definition will be doing real supervisory work within a few years of the final report.

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

My response is deliberately narrow: two recommendations on the human-oversight and monitoring dimension, drawn from what I have seen hold, and fail, when oversight claims are tested. It asks the final report to be one degree more explicit about evidence and decision records: the two artefacts that make every oversight model on the spectrum testable. I support the direction of the report and the decision to frame the practices as a menu rather than a rulebook. I am happy to elaborate on any point.