



FINANCIAL
STABILITY
BOARD

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

Response to Consultation

Ekrem N (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?

The taxonomy is comprehensive and the agentic section is thorough. There is one risk I would add, and I think it belongs in the report's own list rather than as a matter of framing: the risk that institutions misread the scope of a supervisory framework as a statement of expectation.

The clearest current example is the revised US interagency guidance on model risk management issued on 17 April 2026 (Federal Reserve SR 26-2), which the report cites. It supersedes SR 11-7 and keeps traditional AI and machine learning models fully in scope, while placing generative and agentic AI outside it on the ground that they are novel and rapidly evolving. I think that scoping decision is defensible. Writing prescriptive validation requirements for a technology that changes faster than the drafting cycle would produce guidance that is obsolete on issue.

The effect on the ground is what concerns me. The AI whose behavior is best understood is subject to the most prescriptive coverage, and the AI that nobody can fully explain is subject to the least, at the same moment it is being deployed fastest and closest to customers. Nothing in that follows from a judgement that the risk is lower. It follows from the fact that the rules were easier to write for the first category.

The failure mode is a firm concluding that out of scope means unregulated. It does not. Safety and soundness expectations, consumer protection and fair lending law, conduct rules, privacy obligations, and board accountability all continue to apply, and supervisors are applying them to AI during ordinary examinations. What has actually changed is not the level of expectation but its form: the institution has moved from following a rule to defending a judgement, and a judgement is the harder of the two to defend under examination. A firm that reads the carve-out as relief has misread it in the direction that will hurt it.

This matters for an international report because the pattern is not confined to one jurisdiction, and because the FSB is well placed to name it in a way that no single national authority comfortably can. The report's statement that adopting the sound practices does not absolve an institution of its local legal and regulatory obligations addresses this exactly,

and I would give it more prominence than a closing caveat in the introduction. Narrowed formal scope raises the burden on the institution's own judgement. It does not lower it.

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

Yes, broadly. The three-part split, organization-wide governance, lifecycle, and then cyber and third-party risk, matches how risk functions are actually organized, which matters more than it might sound. Frameworks that cut against the existing org chart get implemented on paper and nowhere else.

The part I would single out for support is sound practice 1 on risk appetite. Most guidance in this space addresses materiality and tiering. Those answer the question of how carefully to govern a given system. Almost nothing addresses the question that comes before it, which is what level of AI risk the institution is prepared to accept in the first place. A firm that has never stated an AI risk appetite has not thereby avoided having one. It has one, set implicitly, at the level of whatever its least-controlled deployment happens to be doing. Naming that explicitly is the most useful thing in the report and I would keep it prominent.

One thing I would add. A tolerance with no defined consequence is not a control. In practice, appetite statements tend to arrive as qualitative sentences that name no indicator, no trigger point, and nobody empowered to act when the line is approached. They read well in a board pack and change no behavior. It would be worth saying explicitly that each stated tolerance should carry a measurable indicator, a trigger, and a named person who acts when it breaches.

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 architecture is right: general practices, with flags where autonomy changes the picture. I would also note that the report's candor on human oversight of agentic systems, specifically the acknowledgement that real-time human monitoring becomes impractical as agent use scales and may require augmentation by another AI, is more forthright than most supervisory material I have read on the subject. It is worth keeping that sentence exactly as it is. There will be pressure to soften it.

My substantive comment concerns the supply chain beneath a vendor.

The report is strong on the agent as an actor inside the institution and strong on data governance. It is thinner on what sits underneath a third-party product. Sound practice 12 covers third-party risk including supply chain and concentration, but there is a practical problem it does not quite reach: a vendor register records the vendor. It does not record the foundation model the vendor's product is built on.

So an institution buys four AI-enabled tools from four suppliers, books four independent relationships, and believes it has diversified. If all four sit on the same underlying model, it has not. A defect in that model, a bias, a vulnerability, a bad update, a working attack,

surfaces in all four at once, and nothing in the institution's own third-party inventory would have warned it. The register is measuring the wrong layer.

Suggested change. Sound practice 12, or the documentation attributes under sound practice 3, could recommend that institutions record the underlying foundation model or models a third-party product depends on, where that is ascertainable, and assess concentration at that layer and not only at the vendor layer. The report already asks for documentation of dependencies on third-party AI providers, so this is a small extension rather than a new obligation. And where a vendor declines to say, that refusal is itself a finding worth recording.

Who is able to see correlation?

There is a mismatch in the report that I think is worth naming plainly.

The risk analysis identifies that reliance on common models, datasets, and infrastructure produces correlated behaviour across institutions, and that this can amplify herding and procyclicality and worsen market stress. I agree, and it is properly identified as a financial-stability risk.

The sound practices then assign the management of concentration risk to the individual institution. But an individual institution cannot see system-wide correlation. It can map its own dependencies with effort. It has no way of knowing what its competitors have built on, and no single firm can determine whether the sector as a whole has converged on three providers or thirty. The risk is systemic and the mitigant is institutional, and a firm that follows every sound practice diligently will still be unable to answer the question the report itself raises.

The sound practices are addressed to institutions, so this is not something they can solve alone. Closing it needs aggregation, and aggregation needs an authority. I would encourage the FSB to consider, here or in later work, whether supervisors are better placed to collect foundation-model dependency data at sector level, in the way critical third-party dependency is now being approached in several jurisdictions.

There is a useful complementarity here. A sound practice asking institutions to record their foundation-model dependencies in a consistent, comparable form would generate exactly the raw material such aggregation requires. The institutional practice is a precondition for the systemic remedy, which is an argument for adding it even setting aside its value to the individual firm.

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

Yes. Outcomes rather than prescriptions is the right call for something moving this fast, and sound practice 4 on adaptability is the mechanism that keeps the rest current.

One refinement. The report says institutions reassess materiality and risk periodically or on the occurrence of certain events. Periodic review is the weaker of the two and, being easier to administer, tends to win. For AI the events that matter are frequently external and invisible to a calendar: a vendor changing the model under a product without telling you, a step-

change in a model class's capability, a new attack technique becoming cheap. I would be more directive that reassessment triggers must include external capability and supply-chain events, not only internal incidents and performance degradation. Sound practice 4 gestures at this; connecting it explicitly to the reassessment duty in sound practice 5 would strengthen both.

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.**
6. **Do the case studies in this report provide actionable insights for financial institutions in their responsible AI adoption?**
7. **Are the definitions in the glossary clear and aligned with industry sound practices, including recent developments in AI?**

Clear and consistent with how the terms are used in practice. Two things worth preserving: the distinction drawn between explainability and transparency, and the framing of explainability as a continuum rather than a binary. The observation that explainability is not an end in itself but a means to accountability is the right formulation and I would not lose it in editing.

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