

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

Response to Consultation

Better Markets

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?

We broadly agree with the benefits and risks catalogued in Section 1, and we commend the balance the report strikes, including its important observation that not adopting AI can itself create risk (for example, in fraud and financial-crime monitoring). The taxonomy of risks (governance, data, explainability, performance, human oversight, cyber, third-party, consumer conduct, and agentic AI) is comprehensive at the level of the individual institution. However, several substantive risks are either absent or materially understated, and they share a common feature: each is a risk the adopting institution may have limited private incentive to manage because the harm falls largely on others.

1. Systemic risk from correlated, homogeneous AI behavior is understated

The report acknowledges herding, procyclicality, and correlated outcomes under third-party concentration, but treats them primarily as a byproduct of shared providers. We urge the FSB to elevate model and strategy correlation to a first-order financial-stability risk in its own right. When many institutions rely on similar models, trained on similar data, responding to similar signals, they can buy and sell in lockstep, amplifying shocks rather than absorbing them. This is not hypothetical: algorithmic trading already accounts for an estimated 60–70% of trades, and in October 2025 a crypto flash crash triggered roughly \$19.3 billion in forced liquidations and wiped out approximately 1.6 million accounts within an hour after simultaneous algorithmic selloffs, with no single entity accountable for the cascade. Oxford researchers have separately documented agentic AI in financial simulations exhibiting price collusion and manipulation that no one programmed and that traditional oversight is not designed to detect. The sound practices, which are addressed to individual institutions, are structurally ill-suited to a risk that emerges only in the aggregate. We recommend the FSB state explicitly that firm-level practices are necessary but not sufficient, and that supervisory, sector-level monitoring of model and strategy concentration is required.

2. Consumer harm and algorithmic discrimination deserve equal billing with prudential risk

Section 1.3 treats consumer protection and market conduct as one risk category among many. In our view, for the decisions that AI will increasingly adjudicate for households and small businesses, it is the central one. AI systems already influence whether a person gets

a card, a loan, or an insurance payout. The evidence that these systems can encode and scale discrimination is not speculative. In July 2025, a student-loan lender settled with the Massachusetts Attorney General in the United States for \$2.5 million after its AI underwriting model produced discriminatory outcomes for Black and Hispanic applicants by using college cohort default rates as an input, penalizing graduates of Historically Black Colleges and Universities even with strong credit. A multi-year independent fair-lending monitorship of a major AI lending platform found approval-rate disparities for Black applicants that were 'statistically and practically significant,' and the firm declined to adopt an available less-discriminatory alternative model. Supervisory examinations of credit-card and auto lenders have found disparities in AI-scored outcomes that the institutions lacked the systems to even detect. We recommend the report treat testing for unacceptable bias and disparate impact as a named, non-optional performance dimension for consumer-facing use cases, not a subsidiary consideration.

3. The illusion of inclusion: new predation at algorithmic scale

The report rightly notes AI's promise to broaden access through alternative-data and cash-flow underwriting. We share that hope, but the FSB should name the mirror-image risk: AI-powered instant approval can make it frictionlessly easy to extend unaffordable credit to the borrowers least able to absorb it. Products such as 'earned wage access' and 'buy now, pay later,' underwritten by AI and approved in seconds, have proliferated fastest among lower- and middle-income, Black, Hispanic, and women borrowers in the United States. Expanded access is only a benefit if it is access to non-predatory products. We recommend the benefits discussion be paired with an explicit caution that frictionless approval can facilitate consumer harm, and that affordability and suitability assessment is itself a dimension of responsible adoption.

4. Competitive concentration and the erosion of community banking

The report addresses third-party and provider concentration but not the concentration AI drives among financial institutions themselves. The largest institutions can build data and compute advantages that community banks and credit unions cannot replicate, even though those smaller institutions make a disproportionate share of small-business and relationship-based loans in the U.S. If responsible-adoption expectations are calibrated only to the sophisticated, well-resourced adopter, they may inadvertently accelerate consolidation and reduce the diversity of the financial system, which is itself a source of resilience. We encourage the FSB to recognize competitive concentration and reduced institutional diversity as a financial-stability-relevant consequence of AI adoption, and to ensure proportionality does not become a two-tier system in which only large firms can afford to participate.

5. Macro-financial and buildout externalities

Finally, two adjacent risks fall outside a strict enterprise-risk frame but bear on financial stability. First, concentration risk in the AI supply chain and in markets: investors and retirees across income levels are increasingly exposed to a small number of AI firms whose valuations now represent a significant share of major equity indices. Second, the real-economy capital allocation effects of the AI buildout, including the possibility that concentrated AI capital expenditure crowds out lending and investment in the broader

productive economy, and that data-center energy costs are passed through to ratepayers. We do not suggest the sound practices can resolve these, but the FSB's benefits-and-risks overview would be more complete for acknowledging them as part of the macro-financial context in which adoption is occurring.

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

The twelve sound practices are comprehensive in coverage and generally clear in exposition. The lifecycle structure is logical, and the linkage from governance through use-case management to cross-cutting cyber and third-party risk is coherent. We offer four observations aimed at making the practices operationally effective rather than aspirational.

Clarity is high; enforceability and specificity are where the gap lies

Our principal concern is not clarity of language but clarity of expectation. The report is explicit that the practices are a non-binding 'menu,' are not a standard, and do not impose a prescriptive approach. We understand the FSB's mandate constraints. But a menu from which institutions may select, applied by the same firms that have a commercial incentive to move fast, risks becoming a document that validates whatever an institution was already doing. History teaches that voluntary, principles-only frameworks tend to be honored in proportion to the adopter's pre-existing willingness. We encourage the FSB to be candid that these sound practices are a floor for supervisory dialogue and national implementation, not a ceiling, and to invite member authorities to translate them into supervisory expectations with teeth.

The accountability gap needs a sharper answer

The practices repeatedly and correctly state that a financial institution 'remains accountable for AI outcomes even when the AI was not built in-house.' This is the single most important principle in the document. But the practices stop short of the harder question the report's own risk analysis raises: when an autonomous system produces a harmful outcome that even its creators cannot fully explain, which human or entity is answerable, to whom, and through what mechanism? We recommend SP 2 and SP 10 be strengthened to require that, for every material or high-risk use case, a named, senior, human owner be accountable for outcomes, that this accountability be non-delegable to a vendor, and that it survive the institution's inability to explain a specific decision. Accountability that evaporates precisely when a decision becomes unexplainable is not accountability.

'Shadow AI' and the practical limits of firm visibility

We welcome the report's treatment of shadow AI. We would go further: as general-purpose AI tools become embedded in ordinary software, the assumption that an institution can maintain a complete inventory of its own AI use will come under strain. The practices should acknowledge this candidly and treat comprehensive discovery (not just documentation of known systems) as an ongoing control, since an institution cannot govern what it cannot see.

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 balance is largely appropriate, and the report deserves credit for engaging seriously with agentic AI rather than deferring it. The technology-neutral core, supplemented by agentic- and genAI-specific overlays is a sensible architecture. Two points of emphasis follow.

Agentic AI is arriving faster than the guardrails, and the report should say so plainly

Many financial institutions are already testing AI agents and these systems have moved well beyond trading into consumer lending, insurance underwriting, claims, debt collection, and wealth management, often at nonbanks with minimal regulatory coverage. Meanwhile, when U.S. banking regulators issued revised model risk management guidance in April 2026 (SR 26-2), they explicitly excluded generative and agentic AI from scope as 'novel and rapidly evolving.' There is therefore a live and widening gap between deployment and supervision. This strengthens, rather than weakens, the case for the FSB's work, and we encourage the report to name the gap directly: the most consequential forms of AI are precisely the ones least covered by existing model-risk and third-party frameworks, and that is where sound practices are most urgently needed.

The speed of harm is the distinctive agentic risk

The report's agentic risk list is strong. We would foreground one attribute above the others: agentic harms can materialize and scale at machine speed, faster than human oversight or remediation can operate. This is why we strongly endorse the report's emphasis on pre-deployment testing in controlled environments, human approval and dual authorization for financial transactions (especially those involving customer funds), least-privilege access, and hard limits on an agent's ability to act on the external environment until it is demonstrably safe. We would make one addition explicit: for agents acting on customer funds or making legally consequential consumer decisions, there should be a strong presumption against full autonomy absent a compelling, documented justification and commensurate controls. The burden should rest on the institution to justify autonomy, not on the customer to discover its absence.

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

Yes. The outcomes-based, technology-neutral drafting is the correct choice for a field moving this quickly, and the report's explicit statement that the practices are non-exhaustive and will evolve is appropriate. Flexibility, however, carries a cost that the FSB should manage deliberately.

The risk of a flexible, principles-based framework is that 'flexibility' becomes a reason for inaction, and 'proportionality' becomes a justification for doing less. We offer three suggestions to keep flexibility from becoming a loophole:

- Pair flexibility with a review cadence. Commit to revisiting the sound practices on a defined schedule (for example, no less than every 18–24 months) and after significant capability jumps or AI-related incidents, so the framework adapts as fast as the technology it governs.
- Guard against proportionality as a race to the bottom. Proportionality should scale controls to risk, not excuse high-risk deployments by low-resource firms. The report should be explicit that a use case's materiality and risk, not the institution's size or appetite for control, sets the floor for oversight.
- Make capability-triggered escalation explicit. Because newer AI can acquire capabilities its deployers did not anticipate (a point the report itself makes), the framework should require re-assessment whenever an AI system demonstrates behavior outside its documented scope, treating capability surprise as an automatic trigger for re-tiering.

With those safeguards, the framework's flexibility is a strength. Without them, flexibility and adaptiveness can quietly become permission to defer.

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 concrete, and the fraud-detection and operational-efficiency examples in particular illustrate real benefits. But the set has two systematic biases the FSB should correct.

The case studies are almost uniformly success stories

Nearly every vignette recounts a benefit realized: fraud losses down 20%, preparation time halved, 400,000 manual processes eliminated. Only the 'developmental testing' case (the bank that declined to deploy an underperforming ML model) shows responsible adoption resulting in restraint. A report on responsible adoption is more credible, and more instructive, when it shows what responsible adoption looks like when things go wrong: a use case retired after monitoring detected drift, a vendor dropped after disclosures proved inadequate, a consumer-facing model pulled after bias testing failed, an agent's autonomy rolled back after an incident. We recommend adding failure-and-remediation case studies. The lesson that a firm walked away from a benefit to avoid a harm is the most valuable lesson the report can teach.

Nonbank coverage is thin, and that is where the risk is concentrating

As the FSB anticipates in its question, nonbanks are under-represented. This matters because much of the fastest, least-supervised agentic deployment is happening at nonbanks: AI-native lending platforms, fintech BNPL and earned-wage-access providers, insurtechs, robo-advisers, and third-party AI vendors that are themselves becoming systemically important nodes. We offer several directions for additional, nonbank-focused case studies:

- AI-native lending platform seeking a bank charter. A platform that originates loans with over 90% full automation and no human intervention, whose independent fair-lending monitor nonetheless found significant approval-rate disparities. This illustrates both the inclusion promise and the accountability problem when automation is near-total, and it is directly on point for a report about responsible adoption.
- Robo-advisory and fiduciary duty. A nonbank wealth manager using agentic AI for portfolio construction and rebalancing, illustrating how suitability and fiduciary standards designed for human advisers must be operationalized for autonomous systems.
- Insurtech claims adjudication. An insurer using AI to approve or deny claims, showing human-in-the-loop design where the cost of a wrong denial falls on a policyholder rather than the firm, and the redress and contestability mechanisms that responsibility requires.
- Third-party AI provider as a critical node. A case centered on the vendor rather than the adopter, illustrating concentration risk, model-change-without-notice risk, and the business-continuity and exit-strategy planning SP 12 calls for, from the perspective of the many small institutions depending on one provider.
- BNPL / earned-wage-access underwriting. A nonbank consumer-credit provider using AI for instant approval, illustrating affordability and suitability assessment as a dimension of responsible adoption, and the difference between expanding access and scaling predation.

We would be glad to work with the FSB and its members to develop any of these into fuller case studies.

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

Partly. The most actionable studies are the ones that show a decision process and its constraints, the developmental-testing case, the third-party-risk-management case, and the documentation/inventory case, because a practitioner can extract a transferable method from them. The least actionable are the outcome-focused vignettes that report a headline result (e.g., fraud losses down over 20%) without the governance detail that would let another institution reproduce the responsibility, not just the result. To make the case studies more actionable, we recommend each one explicitly surface: (i) the specific risk that was identified; (ii) the concrete control, guardrail, or decision that addressed it; (iii) what the institution chose not to do, and why; and (iv) how the outcome was monitored after deployment. The point of a responsible-adoption case study is to transmit the judgment, not to advertise the payoff. Quantified benefits without the accompanying controls risk being read as marketing for the technology rather than guidance for governing it.

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

The glossary is clear, current, and well-aligned with prevailing usage; the definitions of agentic AI, foundation models, RAG, data and model drift, and the learning paradigms are accurate and accessible. We offer targeted suggestions.

- Explainability vs. interpretability. The report distinguishes these well in the body but the glossary defines only explainability. Adding 'interpretability' to the glossary, and stating plainly that neither guarantees a model is fair, accurate, or lawful, would prevent the common error of treating an explanation as a validation.
- 'Unacceptable bias' should be defined, not just used. The body relies heavily on the concept of 'unacceptable bias' and notes some bias may be acceptable. Given how much rides on that line, the glossary should define the term and cross-reference disparate impact and applicable fair-lending / anti-discrimination concepts, rather than leaving 'unacceptable' to the reader.
- Accountability and human oversight. The glossary defines model risk but not 'accountability,' 'human oversight,' or the oversight taxonomy (human-in-the-loop, on-the-loop, in-command). Given their centrality, brief definitions would aid consistent use across jurisdictions.
- 'Responsible AI adoption.' The report's titular concept is never defined. A short glossary entry, framed around outcomes for customers, markets, and stability rather than process alone, would anchor the whole document.
- Autonomy as a spectrum. The agentic-AI entry would benefit from noting that autonomy is a continuum (from suggestion to execution), mirroring the report's own useful treatment of explainability as a continuum rather than a binary.

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

We close with four cross-cutting points that do not fit neatly under a single question but bear on the report as a whole.

1. Keep the public interest, not just the institution's balance sheet, in frame

The report's implicit reader is a risk manager protecting the institution. That is appropriate, but it means the framework is most fully developed where the institution's interest and the public's coincide, and least developed where they diverge. The risks with the weakest private incentive to manage (discrimination against a rejected applicant, systemic correlation, cost-shifting onto ratepayers and displaced workers, consolidation of the smallest institutions) are exactly the ones a public-interest body like the FSB is uniquely positioned to foreground. We urge the FSB to make explicit that responsible adoption is measured by outcomes for customers, markets, and financial stability, and not solely by the adopting firm's risk-adjusted return.

2. Firm-level practices cannot address system-level risk alone

The gravest AI-driven financial-stability risks (correlated behavior, herding, common-provider concentration, and cross-institution agent interactions) are emergent properties of the system, invisible from inside any single firm. A document addressed entirely to individual institutions cannot manage them. Given the speed of change and learnings from rapid adoption, we recommend the FSB commit to updating its 2024 Monitoring Adoption of Artificial Intelligence and Related Vulnerabilities in the Financial Sector, focusing specifically on system-level monitoring of AI model, data, and provider concentration, and on the sector-

wide scenario exercises the report gestures toward, as the necessary complement to firm-level sound practices.

3. Avoid pendulum policymaking

The question is not whether to allow AI in finance; that is already an irreversible reality. It is whether we build appropriate transparency, accountability, and liability now, before a preventable failure provokes a public backlash and reactive over-regulation. The industry's own self-interest should favor getting this right early. The worst outcome for everyone (Main Street and the financial sector alike) is pendulum policymaking that swings from under-regulation to over-regulation and back. Sound practices adopted and implemented seriously now are the best defense against that swing.

4. Inclusive process produces better outcomes

Finally, a word on process. Consultations of this kind draw responses overwhelmingly from the regulated industry and its vendors. We encourage the FSB to actively solicit and weight the perspectives of consumer, labor, community-banking, civil-rights, and public-interest voices, whose constituents will experience these systems as subjects rather than operators. The legitimacy and durability of the final framework will depend in part on whether the people most affected by AI in finance had a seat at the table when the practices were written.

Better Markets appreciates the FSB's leadership on this issue and its evident seriousness of purpose. We would welcome the opportunity to discuss any of these comments further, to contribute the nonbank case studies described above, or to otherwise support the FSB's important work toward AI adoption that turns change into progress for the many, not just the few.