

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

Response to Consultation

Bloomberg LP

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?

Bloomberg broadly agrees with the report's framing of benefits and risks, including efficiency gains and improved risk management, alongside governance, data, and operational risks.

It is also important to recognise that the adoption of AI does not alter the underlying regulatory obligations that already apply to regulated financial institutions. Firms remain responsible for managing cyber and operational risks, treating customers fairly, maintaining appropriate governance and risk management frameworks, and ensuring accountability for model risk and outcomes. AI should therefore be viewed as another technology operating within existing regulatory frameworks, rather than requiring a fundamentally different allocation of responsibilities.

Regarding additional benefits, there are several worth drawing out further:

First, AI can significantly improve the accessibility and usability of financial data. Search, summarisation and natural-language interfaces allow users to interrogate large and complex datasets without specialist technical skills, lowering barriers to entry and enabling a broader range of market participants to extract insight from information that was previously difficult or costly to work with. This is particularly valuable in financial markets, where the volume of relevant data has long outstripped the capacity of any individual to process it manually.

Second, AI can enhance market efficiency and price discovery. By enabling faster and more sophisticated processing of data, AI helps market participants identify and act on relevant information more quickly and accurately. Over time, this supports more efficient allocation of capital and can improve the quality and timeliness of pricing across asset classes, including in less liquid or less well-covered segments of the market.

On the risk side, we would highlight the safety risks associated with retrieval-augmented generation (RAG). Where a model retrieves data from untrusted sources, this can introduce inaccuracies and make its outputs less reliable. Importantly, these risks can be managed: by ensuring that models draw on trusted, high-quality data and applying appropriate guardrails, both the general risks of RAG and the risks specific to financial services can be addressed. In this context, we would note Bloomberg's own research on this topic, which

finds that general-purpose AI guardrails often fail to detect the content risks that are specific to financial services, and that a domain-specific approach is therefore needed(*). That research also reinforces a wider point, which is that firms' existing governance, oversight and risk management frameworks are generally well able to accommodate AI-related risks. In most cases, a targeted update to those existing frameworks is needed to deal with AI-specific issues, as opposed to building entirely new governance structures.

(*) Gehrmann et al., Understanding and Mitigating Risks of Generative AI in Financial Services, Bloomberg, 2025, available at <https://arxiv.org/html/2504.20086v1>.

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

Bloomberg considers the proposed sound practices to be broadly comprehensive and sufficiently clear to support the responsible adoption of AI by financial institutions.

We welcome the FSB's principles-based, outcomes-focused approach. However, it would be helpful to clarify that the practices are non-exhaustive and should be applied proportionately, reflecting firms' business models, risk profiles, and use cases.

We also encourage the FSB to reinforce the principle of technology neutrality. Financial institutions are already subject to governance, risk management and accountability obligations, and AI should be integrated within these existing frameworks, with firms managing any additional AI-related risks accordingly. The sound practices should avoid becoming overly prescriptive. For example, under Sound Practice 6, expectations regarding the selection of AI systems or providers should preserve firms' ability to exercise risk-based judgment, rather than creating de facto mandatory requirements or prescribed selection criteria.

Overall, the sound practices provide a useful foundation for responsible AI adoption. Greater emphasis on proportionality, flexibility, and technology neutrality would support effective implementation while enabling continued innovation.

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?

Broadly speaking, Bloomberg considers that the sound practices generally strike an appropriate balance. Framing all twelve practices as generally applicable to any form of AI, while flagging where specific technologies warrant additional attention, is consistent with the principle that financial regulation is best designed to be technology-neutral so that it can endure across successive waves of innovation. This matters especially for AI, given how fast this technology is advancing.

That baseline is then reinforced for more complex forms of AI: enhanced documentation and logging, prompt versioning, and chain-of-thought capture for GenAI (Sound Practice 3); explicit treatment of autonomy, scope of action, and access to systems and sensitive data within materiality and risk assessment (Sound Practice 5); reasoning-path and intermediate-step explainability for agentic systems when feasible (Sound Practice 8); and recognition

that performance assessment becomes progressively harder to automate as one moves from traditional AI to GenAI and agentic AI (Sound Practice 9). This mirrors the approach our own research advocates: using general-purpose frameworks (such as NIST) as a starting point and then adapting them to the specific technology, use case, and domain, rather than treating a generic control set as sufficient in itself.

However, the balance is delicate. These practices should not be taken to mean that controls calibrated generally for traditional AI will transfer cleanly to GenAI and agentic systems with only incremental adjustment, because our findings suggest they often do not. Bloomberg's research on GenAI in financial services (*) found that general-purpose safety guardrails and taxonomies failed to detect most of the domain-specific content risks it evaluated. This safety gap emerges when generic tooling is assumed to be adequate for a specialised, knowledge-intensive domain.

The practical implication is that, for the newer and more complex forms of AI, the sound practices could be more explicit that: (i) risk must be assessed as a function of both likelihood and severity within the specific sociotechnical system in which the AI is deployed, and (ii) single-layer safeguards are insufficient, so that responsible AI adoption depends on layered, domain-grounded mitigation embedded in governance rather than on any one control. Making this expectation clearer would sharpen the treatment of emerging risks without disturbing the technology-neutral foundation, which we support.

(*) <https://assets.bbhub.io/company/sites/51/2025/04/arXiv-Understanding-and-Mitigating-Risks-of-Generative-AI-in-Financial-Services-FINAL-4-25-25.pdf>

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

Bloomberg strongly supports the report's flexible and adaptive approach and emphasises that the sound practices should remain technology-neutral, outcomes-focused, and capable of evolving alongside advances in AI. In this regard, it is important to avoid static definitions and overly rigid controls that may quickly become outdated, as well as requirements that do not adequately reflect technical feasibility or the current state of the art.

This approach recognises the adaptability of existing regulatory frameworks, which generally regulate use cases, outcomes and firms' conduct rather than the technologies they deploy. Firms should remain accountable for the outcomes of AI-enabled activities and compliance with their existing regulatory obligations, irrespective of the underlying technology.

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.

No comment.

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

No comment.

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

No comment.

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

No comment.