



FINANCIAL
STABILITY
BOARD

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

Response to Consultation

Elliptic

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?

Our Position

We broadly agree with the benefits and risks identified in the consultation report. The bilateral framing of AI as simultaneously creating opportunities and threats in the same domain (for example, improving and undermining cybersecurity) is accurate. We comment as a provider of AI-powered blockchain analytics and financial-crime compliance software to financial institutions, VASPs and public authorities, rather than as a financial institution ourselves; our observations are grounded in that supplier and AI-provider perspective.

Additional Benefits Warranting Recognition

The following systemic benefits are underrepresented in the report:

Regulatory compliance efficiency at scale. Across the US, EU, UK, and APAC, financial institutions are deploying AI to automate compliance workflows that previously required substantial manual effort: AML/KYC, transaction monitoring, and sanctions-screening workflows. The productivity gains documented in the FSB's own case studies are even more pronounced in compliance functions, yet the report does not give this use case commensurate prominence.

AI-powered blockchain analytics for financial crime risk management. AI-powered blockchain analytics is an underrecognised benefit absent from the report. Public blockchains generate transaction data at a scale incompatible with manual review: AI-driven graph analytics and cross-chain entity clustering are now the primary mechanism for tracing illicit fund flows and mapping sanctions evasion in digital asset markets. Elliptic's AI Copilot, an LLM-based investigative-summary feature, illustrates the operational impact: it produces plain-language summaries of blockchain risk and transaction analysis that materially reduce analyst review time, while human analysts retain the decision. The FATF Travel Rule, now implemented in over 99 jurisdictions, creates the compliance infrastructure these tools operationalise. The FSB should recognise AI-powered blockchain analytics as a sound practice benefit and note the attendant governance obligations: human review of any output

that informs an enforcement or reporting decision, and traceability of AI-generated outputs to the underlying data and its known limitations.

The report should also recognise a tension specific to financial-crime detection: full disclosure of detection logic or model features would materially aid evasion by illicit actors, so explainability in this domain is appropriately satisfied through compensating controls, such as independent validation, data lineage and outcome traceability, rather than through disclosure of the methodology itself.

Additional Risks Warranting Recognition

Regulatory fragmentation as a systemic risk. The proliferation of jurisdiction-specific AI governance frameworks now creates a compliance architecture risk that is itself material. A technology provider serving financial institutions across multiple jurisdictions must map a single product simultaneously against the EU AI Act, DORA, and divergent US and APAC expectations. Where these obligations are directionally compatible but terminologically and procedurally divergent, the compliance cost can exceed the operational benefit of the AI system. The FSB's sound practices should explicitly identify this risk and commit to fostering framework interoperability. Regulatory fragmentation is not a theoretical concern: the EU AI Act's Omnibus amendment, passed by the European Parliament on 16 June 2026, which deferred high-risk AI compliance obligations to December 2027, was in part motivated by the compliance burden imposed on institutions already managing multiple overlapping regulatory regimes.

AI governance capability asymmetry between large and small institutions. Smaller financial institutions and non-banks lack the internal resources to build comprehensive AI governance programmes. The MAS Project MindForge Consortium model, which brought together institutions of varying sizes to develop a shared AI Risk Management Toolkit (published March 2026), and the UK FCA's AI Live Testing pilot (first cohort October 2025), which explicitly sought to help smaller firms overcome 'proof of concept paralysis,' provide empirical evidence that this asymmetry is real and addressable through structured public-private collaboration. In the digital-asset sector in particular, we observe that many smaller VASPs and non-bank institutions rely on their technology vendors to supply much of the AI governance scaffolding they cannot build in-house, which places a corresponding responsibility on providers and is a dimension the report does not address.

AI-empowered systemic cyber threats. The OCC's May 2026 Semiannual Risk Perspective and the HKMA's June 2026 circular on cyber resilience both identify frontier AI models as marking a potential step change in the global cyber threat landscape. The report addresses AI cyber risks but does not adequately capture the systemic dimension: if multiple financial institutions face AI-empowered attacks simultaneously, the aggregate impact on financial stability could exceed what any individual institution's controls can absorb.

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

Our Position

The 12 sound practices are largely comprehensive. Our observations below are in the areas where we have direct experience as an AI provider and third party to financial institutions.

Cross-Framework Alignment

Third-party risk. As an ICT third-party service provider to EU financial institutions, we are directly subject, through our customers' supply chains, to the contractual, resilience-testing and register obligations DORA imposes. We would encourage the FSB to align its third-party AI guidance with DORA rather than create a parallel construct, and to recognise a category the current drafting overlooks: third-party AI that is itself a risk or compliance control (such as sanctions screening or blockchain analytics) where the appropriate assurance model differs from that for AI embedded in an institution's own operations.

Areas for Improvement

Cross-framework mapping annex. A valuable addition to the final report would be an annex mapping each of the 12 sound practices to their equivalents in the major frameworks (EU AI Act, DORA, UK FCA/PRA guidance, US model risk management guidance, NIST AI RMF, ISO 42001). Compliance teams spend significant resources on this mapping exercise independently. A FSB-published concordance would reduce duplicative effort and signal that regulatory convergence is an objective rather than an aspiration.

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?

Our Position

The graduated treatment of GenAI and agentic AI within each practice, rather than through a separate regulatory tier, is the correct structural approach. A single, technology-neutral framework that scales obligations to risk is easier for vendors and their customers to implement consistently across a product than a set of technology-specific tiers.

EU AI Act Tension with Technology-Neutral Framing

The EU AI Act introduces a categorical, use-case-based risk classification system that sits in some tension with the sound practices' technology-neutral framing. Certain financial AI applications, including credit scoring and insurance pricing, are classified as high-risk regardless of the technology used, while others, including fraud detection, are explicitly excluded from high-risk classification. This categorical approach creates practical compliance questions for institutions operating both EU and non-EU AI governance frameworks simultaneously. The FSB's sound practices should acknowledge this structural difference and provide guidance on how providers and institutions can apply the technology-neutral proportionality framework in a manner consistent with the EU AI Act's categorical risk classification.

GenAI-Specific Guidance

For GenAI used in financial-crime investigation, our own domain, the principal quality risks are unfaithful or fabricated summaries rather than adverse consumer outcomes. We

manage these through output-to-source traceability, prompt and model versioning, evaluation against known cases, and mandatory human review before any output informs a reporting or enforcement decision. The performance-management practice could usefully acknowledge this class of assistive, non-scoring GenAI, which neither scores nor profiles individuals for decision-making.

Agentic AI Guidance

We welcome the forward-looking treatment of agentic AI. From a provider's standpoint, the key control as autonomy increases is graduated human oversight tied to the consequence of the action, together with a clear allocation of responsibility between the provider and the deploying institution. We treat any move from assistive to agentic functionality as a change that must re-enter design review and re-assessment, rather than as an incremental update.

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

Our Position

The technology-neutral framing is the sound practices' most durable structural feature. The concern is not with the framing itself but with two practical risks: that the guidance becomes too abstract to provide operational value, and that the absence of an explicit interoperability mechanism allows national divergence to compound over time.

Divergence between the major regimes is a real long-term risk, particularly for suppliers of AI products across multiple jurisdictions and regulatory regimes. We would support the FSB providing operational guidance that helps both institutions and their vendors satisfy the strictest applicable requirement once, rather than maintaining parallel programmes.

Recommendations for Future-Proofing

Standing update mechanism. The FSB should commit to periodic updates and establish a standing mechanism for incorporating practitioner learnings, including from technology providers, as AI capabilities evolve.

Interoperability standard. The FSB should develop a reference architecture identifying the minimum governance elements with which national frameworks should be compatible. This would function analogously to the Basel Committee's capital adequacy minimum standards: a floor below which national frameworks cannot fall, above which they may impose additional requirements, and across which bilateral recognition agreements can be negotiated.

Foundation model governance. No current framework across the US, EU, UK, or APAC provides adequate guidance on financial institutions' governance of the foundation model providers on which most financial AI now depends.

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.

Our Position

The case studies are well-chosen and provide instructive operational detail. The consultation specifically invites nonbank case studies; we note a gap the existing studies do not cover.

Gaps in Existing Case Studies

Third-party AI as a compliance control. To our knowledge, no case study addresses AI supplied by a third-party vendor and used by an institution as a compliance control. This is now a pervasive feature of financial-crime compliance and raises assurance questions, provider transparency, data quality, human oversight, contractual allocation of responsibility, that differ from those for AI an institution builds and operates itself.

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

Our Position

The case studies are directionally useful but would significantly increase in operational value if structured more systematically and anchored more explicitly to the binding regulatory obligations that practitioners must navigate.

Recommendations

Standardise structure and add regulatory mapping. Each case study should include a section identifying the regulatory frameworks applicable to the institution described and how the practices illustrated would satisfy those frameworks' requirements. A case study on fraud detection AI, for example, should note that fraud detection is excluded from the EU AI Act high-risk category, and identify which sound practices the described controls satisfy.

Include failure and near-miss cases. The case study on the ML trading model that failed developmental testing and was not deployed is the most instructive case study in the report precisely because it illustrates governance working as intended. More failure cases would strengthen the practical value of the report.

Address the governance-innovation tension. The tension between rapid adoption and rigorous governance is not confined to the deploying institution; it is felt acutely by the providers whose AI they adopt, and is often best resolved there. When a provider builds governance in by design, the institution inherits assurance rather than rebuilding it, and adoption is faster and better governed. The policy concern that governance impedes adoption, reflected in the White House AI Action Plan (July 2025) and the UK government's January 2026 letters to regulators, is most effectively answered at this supply-chain layer. Case studies illustrating how a provider embedded governance into a deployed AI product could be more useful than those treating governance as a downstream, ex-post constraint.

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

Our Position

The glossary is substantially sound. The following observations identify specific definitional points that bear on our position as an AI provider and third party.

Specific Definitional Observations

‘High-risk AI’. The glossary does not define this term, despite the EU AI Act’s high-risk classification system being the primary mechanism through which providers and EU-regulated institutions will determine the applicable intensity of governance obligations under the sound practices. The FSB should explain how its proportionality framework relates to that categorical risk classification.

‘Materiality’. The sound practices use materiality as a key governance trigger but the glossary does not define it. The EU AI Act uses a categorical classification rather than a materiality threshold; US model risk management guidance refers to the significance and complexity of model use; UK PRA guidance uses a qualitative materiality judgment. A workable FSB definition of materiality, anchored to concepts from operational resilience frameworks (critical operations, important business services) and compatible with all three major framework families, would reduce interpretive divergence and enable more consistent international implementation.

‘Human oversight’. The glossary does not define human oversight despite it being the subject of a dedicated sound practice and a key requirement in every major framework. The taxonomy of oversight modalities in the relevant sound practice (human-in-the-loop, human-on-the-loop, human-in-command, kill switch, contestability) should be incorporated into the glossary with brief descriptions, given that different frameworks use these terms differently.

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

Concluding Observations

The FSB’s sound practices represent a timely and important contribution to global AI governance in financial services. Their greatest potential lies not merely in setting standards for individual institutions but in providing the international reference architecture that enables national frameworks to converge toward a workable interoperability standard.

The current landscape is characterised by directional compatibility and operational divergence. The EU AI Act, UK principles-based framework, and US model risk management tradition all share commitments, but differ substantially in how these commitments are operationalised: through categorical risk classification and binding conformity assessment (EU), through supervisory judgment and outcomes monitoring (UK), and through model validation documentation (US). APAC frameworks add further variation, with the HKMA’s sandbox-led learning approach and MAS’s structured consortium model providing tested mechanisms for translating principles into practice.

The FSB has a unique opportunity to address this. We would recommend the following, from our perspective as an AI provider and third party to financial institutions:

Publish a cross-framework mapping as an annex to the final report, mapping each of the 12 sound practices to their equivalents in major international frameworks. This single document would reduce duplicative compliance effort.

Recognise AI-powered blockchain analytics as a benefit in its own right, and treat third-party AI that is itself a compliance control (such as sanctions screening or blockchain analytics) as a distinct category within the third-party risk practice.

Engage explicitly with the EU AI Act's Omnibus amendment and what the deferral of high-risk AI compliance obligations to December 2027 signals about the sustainable pace of AI governance implementation. The FSB's sound practices should be calibrated to implementation timelines that institutions and their vendors can realistically meet.

Commit to addressing regulatory fragmentation as a systemic risk in its own right, with a concrete work programme to develop the interoperability standard described above.

We look forward to the FSB's continued engagement on these issues and remain available to provide further input or to participate in consultative workshops on any of the themes raised in this response.