

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

Response to Consultation

techUK

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 set out. The report's coverage of governance gaps, data quality, explainability, performance degradation, human oversight, cyber risk, third-party dependency and agentic-specific risks is thorough and reflects what our members are experiencing in live deployment. [FSB Consultation Report, Section 1.3 "Risks and implementation challenges", p.12-17].

Three gaps are worth flagging:

The role of insurance in managing AI risk is missing. The report treats AI risk primarily as something to be governed and supervised, with no discussion of insurance as a market mechanism that prices and incentivises good governance. Insurers evaluating an organisation's assurance practices when setting coverage and premiums makes safety legible in the bottom line, not just a compliance requirement. Cyber insurance provides an analogy: insurers requiring evidence of security controls before underwriting has driven measurable improvements in firms' cybersecurity practices. A similar dynamic applied to AI governance would make assurance a commercial imperative, not just a regulatory one.

Liability across multi-party AI supply chains remains functionally unresolved. The report acknowledges third-party dependency risk but understates how genuinely difficult accountability becomes once harm is separated from its cause by several steps. [FSB Consultation Report, Section 1.3 "Third-party dependencies and concentration", p.12-17]. Existing accountability frameworks were designed for decisions made by individuals, not for harm that emerges from a chain of autonomous system interactions e.g. SM&CR assigns personal accountability to named senior managers, but that framework strains when the decision that caused harm was made by a foundation model, executed by a cloud-hosted application, and deployed by an institution whose compliance team approved a system they could not fully audit. Our member company procurement teams describe liability as 'increasingly unplaceable' once foundation model providers, cloud infrastructure, application vendors and deploying institutions are all in the chain. [AI Procurement Frontline report, "Liability remains the unanswered question", p.9] The same accountability-fragmentation problem came up independently in our critical national infrastructure roundtables. [CNI Roundtable notes ("Additional Input 1"), Section 2].

The report's own frontier AI carve-out deserves more attention than a footnote. The FSB has been explicit that these practices were not developed to address frontier AI model risks, "although some sound practices would help financial institutions respond to such risks." [FSB press release, 10 June 2026, quoting Ho Hern Shin] Given the pace at which frontier capability is advancing relative to the report's own ambition to inform a US G20 deliverable by October 2026 [FSB press release, 10 June 2026, quoting Michelle Bowman], we think the FSB should say more about how it intends to close this gap.

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

Largely yes on structure. The 12 practices map sensibly onto governance and lifecycle stages, and the proportionality principle is well judged. [FSB Consultation Report, Executive summary, p.1] Two things would strengthen clarity in practice.

The distinction between nominal and substantive human oversight needs sharper, more operational treatment. Sound Practice 10 rightly distinguishes "meaningful" oversight where humans have "sufficient ability, authority, and incentive to intervene" from oversight that is "nominal or driven by 'tick-the-box' compliance," but doesn't give firms much to test against. [FSB Consultation Report, Sound Practice 10, p.41] Our own report reaches the same conclusion from the deployment side: "the risk is that human oversight becomes a procedural requirement rather than a genuine control." [Agents of Change, Chapter 3, p.16] In practical terms, this is what one AI risk practitioner described in our recent webinar as a reviewer expected to check hundreds of AI decisions a day being unable to do so meaningfully [AI Assurance Webinar Transcript, approx. 41:35-42:11]. We'd suggest the FSB add a practical marker: oversight is only meaningful where the person responsible has the skill, time, and authority to actually challenge the output not simply the formal responsibility to sign it off. [CNI Roundtable notes ("Additional Input 1"), Section 9, "Materiality and human oversight standards"].

Testing and validation should not be conflated with ongoing assurance. Sound Practice 9 is thorough on developmental testing and performance dimensions but leans throughout on point-in-time verification language. [FSB Consultation Report, Sound Practice 9, p.37-40] Our CNI roundtable participants were emphatic on this distinction: "audit is not assurance, and testing AI is not the same as assuring it." [CNI Roundtable notes ("Additional Input 1"), Section 1] We'd welcome the final report being clearer that dynamic, continuous assurance built into deployment pipelines, not scheduled as a periodic exercise, is the standard being asked for, not an optional enhancement for high-risk use cases only.

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, yes. [FSB Consultation Report, Introduction, p.3-4] We'd note one practical tension worth naming explicitly rather than leaving implicit. The report's structure treats agentic AI risk as an extension of existing sound practices with "additional considerations" layered on. [FSB Consultation Report, Sound Practice 10, p.42-43] Our members' experience deploying agentic systems in production suggests agentic governance is better understood as an

architectural decision made at design time, not a bolt-on control applied afterwards. A couple examples which illustrate this:

- DXC's payments agents distinguish advisory, assisted action and execution agents by design, with segregation of duties and audit trails built in from the outset. [Agents of Change, Chapter 2, p.12; Use Case Annex, p.5]
- Microsoft's claims triage system and Lloyds' fraud contact centre assistant follow the same pattern. [Agents of Change, Chapter 2, p.9-10, p.12-13; Use Case Annex, p.15, p.21].

Our own report states this directly: "governance of agentic AI is most effective when it is an architectural decision from the outset." [Agents of Change, Chapter 2, p.13] We'd encourage the final report to state this more directly, since it has real implications for when in the AI lifecycle firms need to be applying Sound Practice 10.

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

This is our main area of concern. The FSB has been open that this toolkit does not address frontier AI risk and was developed on an accelerated timeframe. [FSB press release, 10 June 2026] We think that's the right call for getting something useful out quickly, but it creates a structural problem: frontier models are being released and materially updated on cycles measured in weeks, while governance frameworks, including this one, are designed, consulted on, and embedded on cycles measured in years. The report identifies this speed mismatch as a risk in principle [FSB Consultation Report, Section 1.3, p.12-17] but doesn't apply the same logic to itself. Our own report names this dynamic explicitly: "Frontier AI models are being released, updated, and refined on timescales quicker than typical governance process timescales." [Agents of Change, Chapter 3, p.16]

We'd recommend the final report commit to an explicit, time-bound review mechanism just a general statement that the practices "may be refined as AI technologies evolve" [FSB Consultation Report, Executive summary, p.1] so that the toolkit's own currency doesn't quietly erode between now and whenever the next iteration is scheduled. There is precedent for treating this kind of adaptive cycle as a feature rather than a weakness: our notes on AI insurance markets suggest firms sometimes adopt safety practices to satisfy insurer requirements ahead of regulation catching up, effectively creating a faster-moving complement to formal standard-setting.

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.

We can offer several. Our Agents of Change report and use case annex include contributions from a wide range of UK institution types, several of which speak directly to gaps in the FSB's current case study set:

- Insurance: Marsh's Sentrisk platform, which maps supply chain risk from customs and logistics records to bring previously uneconomic products to market [Agents of Change,

Chapter 2, p.12; Use Case Annex, p.20]; Microsoft's agentic claims triage system for FNOL processing [Agents of Change, Chapter 2, p.12-13; Use Case Annex, p.21].

- Capital markets: Google Cloud's agentic investment research deployment across multiple London-based hedge funds [Agents of Change, Chapter 2, p.11; Use Case Annex, p.18]; LSEG's AI Ready Content, delivering licensed data via Model Context Protocol [Agents of Change, Chapter 2, p.11; Use Case Annex, p.19].
- Payments: DXC Technology's agentic payments operations [Use Case Annex, p.5] and Visa's Intelligent Commerce infrastructure for consumer-facing agent-initiated transactions [Use Case Annex, p.12].
- SME lending: Giffitid's deterministic scoring architecture, a useful counterpoint to the FSB's own explainability discussion, since it's a live example of a firm choosing an inherently auditable design over a marginally higher-performing but less explainable alternative. [Use Case Annex, p.6]
- Cross-sector infrastructure: the FCA's Supercharged Sandbox (NayaOne/NVIDIA), which moved from 132 applications for 22 places in its first cohort to a second cohort specifically focused on agentic AI. [Agents of Change, Chapter 2, p.13; Use Case Annex, p.7]

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

Partially. The FSB's case studies are well-written but anonymised (a large internationally active bank), which limits how actionable they are as a firm reading them can't tell how mature the deployment is or what scale it's operating at. [FSB Consultation Report, e.g. case studies at Sound Practice 1, p.20 and Sound Practice 9, p.39].

Our own use case annex tags every entry with a maturity marker: live at scale, live limited scale, pilot, or technical testing, precisely because that context is what makes a case study usable rather than just illustrative. [Use Case Annex, throughout, e.g. p.2-21] A reader needs to know whether they're looking at a proven, production-scale pattern or an early-stage experiment before they can extract anything actionable from it. We'd recommend the FSB adopt a similar maturity classification for the case studies in its final report.

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

Mostly, yes -- the definitions are clear and technically sound. [FSB Consultation Report, Glossary, p.56-58] Two additions would strengthen it:

"Assurance" itself isn't defined, despite being central to the report's purpose. Given the confusion our own roundtables have surfaced about how assurance differs from audit or testing, a distinction we think matters enough to be worth stating explicitly [CNI Roundtable notes ("Additional Input 1"), Section 1]. We'd suggest the glossary define it directly, rather than leaving it to be inferred from how the term is used throughout the report.

Data drift and model drift are defined, but behavioural drift is not treated as a distinct category [FSB Consultation Report, Glossary, p.56], despite arguably being the harder problem. Our own report defines it precisely: "model outputs gradually change over time as real-world data diverges from training conditions... A model can pass every health check while quietly producing different results than it did six months before." [Agents of Change, Chapter 3, p.14-15] This is a genuinely different failure mode from data and model drift as currently defined, and worth naming as such, it is, in our members' experience, significantly harder for existing monitoring tools to detect. [Agents of Change, Chapter 3, p.15].

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

techUK were happy to input into this and would be delighted with further engagement with the FSB.