



To:
Ms. Michelle W. Bowman
Chair
Financial Stability Board Standing Committee on Supervisory and Regulatory Cooperation (SRC)

Ms. Ho Hern Shin
Lead
Financial Stability Board SRC Workstream on Artificial Intelligence

CC:
Mr. John Schindler
Secretary General
Financial Stability Board

21 July 2026

By email:

Re: GFMA Comment on the Financial Stability Board’s Consultation Report on “Sound Practices for Responsible Adoption of Artificial Intelligence (AI)”

The Global Financial Markets Association (GFMA)¹ welcomes the opportunity to respond to the Financial Stability Board’s (FSB) “Sound Practices for Responsible Adoption of Artificial Intelligence (AI): Consultation report.” GFMA and its members appreciate the FSB’s leadership in this important area and its continued efforts to promote a globally consistent and coordinated approach to the responsible use of AI in financial services.

GFMA members have long utilized AI across a range of applications. Recent advances in technology, including generative AI (GenAI) and agentic AI, are further expanding these use cases and have the potential to deliver significant benefits to financial institutions, their clients, and the broader economy. Where novel or significant risks may arise from these advances, GFMA believes the sound practices are well positioned to support effective mitigation. Additionally, while many of these developments are reflected in the report’s illustrative case studies, which are primarily focused on banking activities, GFMA encourages the FSB to further consider AI adoption across non-bank financial institutions as it finalizes the report, consistent with its cross-sectoral mandate.

GFMA is generally supportive of the proposed sound practices, in particular their ability to tailor to existing processes and governance, as well as their flexible and outcomes-based approach. In GFMA’s view, it is important that the sound practices are not intended to create prescriptive supervisory expectations and remain principles-based, proportionate, technology-neutral, and aligned with internationally recognized risk management approaches. GFMA also encourages the FSB to further clarify that firms can leverage existing governance, model risk, operational risk, cybersecurity, privacy, compliance, and third-party risk frameworks where they already achieve the intended outcomes, thereby avoiding unnecessary duplication, mission creep, or parallel AI-specific governance structures.

¹The GFMA represents the common interests of the world’s leading financial and capital market participants, to provide a collective voice on matters that support global capital markets. We advocate on policies to address risks that have no borders, regional market developments that impact global capital markets, and policies that promote efficient cross-border capital flows, benefiting broader global economic growth. The [Global Financial Markets Association](#) (“GFMA”) brings together three of the world’s leading financial trade associations to address the increasingly important global regulatory agenda and to promote coordinated advocacy efforts. The [Association for Financial Markets in Europe](#) (AFME) in London, Brussels and Frankfurt, the [Asia Securities Industry & Financial Markets Association](#) (ASIFMA) in Hong Kong and Singapore, and the [Securities Industry and Financial Markets Association](#) (SIFMA) in New York and Washington are, respectively, the European, Asian and North American members of GFMA.



afme/

asifma

sifma

GFMA believes that avoiding regulatory fragmentation should remain a key priority. To that end, coordination with other international standard-setting bodies will be important, including with the International Organization of Securities Commissions (IOSCO), which recently announced it will be focusing its attention on “emerging industry practices across governance, disclosure, and recordkeeping and reporting of AI systems in capital markets.”² Greater alignment across initiatives will help promote consistency and support effective global implementation.

GFMA also encourages continued coordination between industry and the official sector, including participation from non-financial firms such as AI providers and developers. This public-private collaboration will support the collective understanding of these rapidly evolving technologies and their associated impacts, as well as overall system transparency.

GFMA’s responses to the consultation questions are set out below. We look forward to continued dialogue as this work progresses and are grateful for the opportunity to provide these comments.

Sincerely,

A handwritten signature in black ink, appearing to read "Allison Parent", with a stylized flourish at the end.

Allison Parent
Executive Director
Global Financial Markets Association (GFMA)

²See [Media Release](#) “IOSCO publishes AI Supervisory Toolkit” (May 2026)

Responses to the questions raised in the Consultation Report

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?

GFMA broadly agrees with the benefits and risks of AI adoption as described in the report. GFMA also welcomes recognition that not adopting AI itself may create risks, as this provides a more balanced perspective on AI adoption.

GFMA believes the final report should continue to preserve a principles-based, proportionate, and technology-neutral approach. In particular, sound practices and case study examples should not be transposed into prescriptive supervisory expectations. The report could also more explicitly recognize the benefits that AI can serve as an enabler of control functions and as a tool to scale organizational knowledge and decision-support capabilities across large institutions. Additional risk considerations are included below.

Fragmentation Risks

The final report should acknowledge the risk that fragmented or inconsistent supervisory expectations across jurisdictions may create implementation challenges for global financial institutions and that global firms need AI governance frameworks that can be implemented consistently across legal entities, business lines, and jurisdictions while still accommodating local legal and supervisory requirements. GFMA and its members are already witnessing significant fragmentation across regulatory elements of AI frameworks within certain regions, even with regard to, but not limited to, basic elements such as definitions. This adds unnecessary complexity and cost and gives rise to the risk of conflicting and/or duplicative requirements.

Fragmentation risk may also arise at the technical implementation layer. Industry experience indicates that major AI and foundational model providers often offer proprietary guardrails, safety controls, audit logging, and control interfaces. Where firms use multiple providers, this can require bespoke and non-portable security and monitoring implementation across vendors, increasing operational complexity and potentially exacerbating vendor lock-in and related concentration risk. The report could acknowledge this challenge and, where appropriate, encourage continued industry work towards greater interoperability of relevant control artifacts and interfaces, while recognizing that implementation approaches will evolve over time.

AI-Specific Considerations

We welcome the report's recognition that many AI-related risks are not entirely new and that existing tech-neutral frameworks, including board and senior management governance structures, are a suitable foundation for managing AI risks. Governance boards should focus on strategic direction, risk appetite, accountability, and effective challenge rather than technical model-level approval. The report should more clearly distinguish between risks that are a manifestation of risks long present within financial services, versus those unique to AI technologies. References to explainability should also not be read as general constraints on the use of less explainable models where performance, risk, and business objectives justify their use and appropriate controls are in place. Additionally, we recommend that the report clarify that utilization of AI to support or replace an existing process should not automatically create an assumption of increased risk for the process.

Incident Reporting

Discussions on cyber and ICT risk information-sharing should recognize that the cross-border sharing of AI incidents depends on applicable global and national legal and regulatory frameworks, and the need to harmonize and coordinate overlapping incident reporting requirements. As an example, the report could reference the FSB's Format for Incident

Reporting Exchange (FIRE)³ as a potential model, while recognizing the need to harmonize AI incident and cyber reporting. The report could also encourage greater clarity over time on AI incident taxonomies (for example, covering issues such as model drift, data poisoning, prompt injection, and other AI-specific failure modes), and how these map to existing cyber and ICT incident reporting obligations. This would support interoperability and consistency while avoiding duplicative reporting regimes.

Third-Party Risk

With respect to discussion of third-party risk management practices, we encourage the FSB to more explicitly recognize the challenges and proportionate expectations for AI embedded within existing third-party products and services (i.e., beyond the reference in footnote 97). The use of AI in non-AI service delivery as well as in third-party supply chains is increasingly prevalent in financial services and may present distinct challenges, including relating to transparency.

Additionally, while the report implicitly recognizes the challenges associated with third-party AI transparency, we believe this could be elevated and framed as a more explicit challenge for deployers. Limited access to information may not itself constitute a distinct AI risk, but it can constrain firms' ability to identify, assess, and manage AI-related risks effectively, and therefore act as a practical barrier to responsible AI adoption. The report should expressly note that this challenge stems from the inherent information asymmetry between deployers and third-party AI providers; effective third-party AI risk management is fundamentally dependent on access to sufficient information from third-parties to enable financial institutions to assess and manage the risks, and evidence compliance with their regulatory obligations. More explicit recognition of this challenge is particularly important given the report's acknowledgement that existing assurance mechanisms and compensatory measures have limitations and may not, in all circumstances, provide firms with the information needed to support effective oversight. Notably, some jurisdictions have domestic frameworks in place which may be utilized to help support this information sharing such as the UK's statutory framework for Critical Third Parties (CTPs)⁴; however, prudential regulatory frameworks often do not directly apply to third-party providers in other jurisdictions.

GFMA also encourages the FSB to more clearly reflect the reality that AI risk management operates across a value chain, in which responsibilities and control levers are not uniform across actors. Cloud providers, model developers, downstream integrators, and financial institution deployers each have differing levels of visibility and ability to mitigate risk. This has practical implications for what can be reasonably expected of deployers, particularly where third-party providers are not within the direct scope of prudential oversight in relevant jurisdictions.

In this context, the report could more explicitly recognize emerging "shared responsibility" approaches that seek to delineate responsibilities among AI vendors, their upstream suppliers, and deployers. Such approaches can support accountability without implying that financial institutions can unilaterally control risks embedded deep within the AI supply chain.

To help address the information asymmetry described above in a practical and scalable manner, GFMA suggests the report further acknowledge the role of standardized artifacts (where available and appropriate), such as model/system cards, audit artifacts, certifications, and standardized "data cards" or equivalent disclosures that can provide further information to better assess data provenance, quality, and key limitations. While the report highlights many of these as potentially useful tools, it should do so in the context of preserving flexible, outcomes-based approaches and avoiding prescriptive requirements that may not be feasible across all providers and deployment contexts. Additionally, the report should encourage third-party providers to utilize these artifacts and share relevant disclosures of material model changes in a timely manner.

³ See FSB's [Format for Incident Reporting \(FIRE\): final report](#) (April 2025)

⁴ See Bank of England, Financial Conduct Authority, and Prudential Regulation Authority joint "[Approach to the oversight of critical third parties](#)" (November 2024)

It is also important to note that financial institutions may face limitations in their ability to substitute AI systems and that any recommendation linked to exit strategies should be tempered realistically. Further, financial institutions may lack contractual bargaining power with AI providers and so using contractual compensating controls to mitigate risk may not be a realistic expectation. Some firms may also become increasingly reliant on a small number of AI providers (which further inhibits their contractual bargaining power) and these types of concentration risks should be recognized as systemic level risks rather than putting concentration risk management on firms at the individual level. GFMA recommends that the FSB encourage its members to further explore this issue.

Macroprudential Risk

The emerging, firm-specific practices described in this consultation will provide a solid foundation for evaluating whether there are additional macroprudential risks, such as: (i) the potential for correlated model failures during periods of market stress, (ii) herding behaviors arising from similar AI-driven investment or risk management decisions, and (iii) the amplification of procyclicality.

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

We consider the sound practices to be broadly comprehensive. However, some parts could be interpreted as more prescriptive than intended. We encourage the FSB to preserve the report's principles-based, proportionate, and technology-neutral approach, particularly in areas such as documentation, data governance, testing, and human oversight. For example, references to maintaining AI inventories through dedicated tools, detailed documentation expectations under Sound Practice 3, or specific cyber frameworks such as OWASP ML Security Top 10, MITRE ATLAS and the NIST adversarial machine learning framework under Sound Practice 11, should be clearly framed as examples rather than expected approaches. Similarly, case studies should be presented as illustrative examples, rather than preferred implementation models. This would be consistent with the FSB's stated intention to provide a non-binding menu of sound practices, rather than a prescriptive international standard.

GFMA also encourages the FSB to avoid language that could be interpreted as requiring "human oversight" in all circumstances or for all use cases. Oversight can take different forms and should be proportionate to the materiality and of the use case. For lower-risk and internal productivity use cases, technology-enabled oversight and monitoring mechanisms (e.g., automated controls, alerting, benchmarking against alternative models, and independent testing) may be effective and operationally more sustainable than requiring manual review steps that are not value-adding. The final report should preserve flexibility in how effective oversight is achieved, while maintaining clear accountability at the firm level.

The structure of the sound practices is useful, but the final report could improve clarity by more explicitly emphasizing that existing governance, model risk, data governance, technology risk, third-party risk, compliance, and operational resilience frameworks can be leveraged where they already achieve the intended outcomes. For example, at paragraph 3.9, the report notes that financial institutions should leverage existing regulatory frameworks as a foundation for managing third-party risks – however falls short of specifying that firms' existing third-party risk management frameworks remain the primary mechanism through which AI-related third-party risks should be managed, with targeted enhancements only where AI introduces genuinely distinct considerations. This approach would be consistent with established AI risk management frameworks that integrate AI oversight into enterprise risk management rather than creating a parallel regime.

The final report should also continue to emphasize proportionality, including that lower-risk and internal productivity use cases should not be subject to the same level of documentation, validation, and senior governance as high-materiality or



afme/

asifma

sifma

customer-impacting use cases. Furthermore, we agree that having lines of defense or other types of structures for defining clear roles and responsibilities for AI adoption is a sound practice for governing AI, it is worth clarifying that, for use cases presenting lower risk profiles, it may be appropriate for lines of business to conduct testing and monitoring without involvement of the second line.

GFMA supports the report's focus on AI-related cyber and ICT risks and welcomes the recognition that firms may need to incorporate AI cyber and ICT risk scenarios into resilience planning, testing, and exercises. In particular, firms and authorities will need to consider what constitutes a "severe but plausible" scenario for purposes of operational resilience planning. However, the final report should avoid encouraging an overemphasis on extreme tail-risk scenarios in a manner that could distort how firms prioritize and allocate cyber resources. Most industry guidance and regulatory statements on frontier AI cyber risk continue to recognize the foundational importance of cyber hygiene, including identity and access management, vulnerability management, secure configuration, monitoring, incident response, and other established preventative and detective controls. These fundamentals should remain the priority of the sector. While it is important to continue exploring strategies for recovering from an extreme AI-enabled cyber event, this remains a nascent area, and many potential solutions discussed by regulators are not yet market-ready or scalable for large and complex financial institutions. There is a real opportunity-cost risk if firms are encouraged to prioritize highly complex, untested solutions over proven preventative and detective controls where further improvement remains possible.

Similarly, while improving the speed and effectiveness of patching is a recognized priority given developments in frontier AI and compressed exploitation windows, firms must retain the ability to apply a risk-based approach to vulnerability classification and patch deployment. Patching alone will not address the full range of cyber risks facing the sector, and a "no-risk" approach to patching is not viable at scale. Such an approach could have negative implications if firms are diverted from prioritizing the vulnerabilities and patches that present the greatest risk to critical operations, sensitive data, customers, or market functioning. The final report should therefore emphasize risk-based vulnerability management, including compensating controls where immediate patching is not feasible or could itself create operational risk.

GFMA also welcomes the report's inclusion of a defense in-depth approach involving layered controls that strengthen a firm's security posture. However, the final report should frame controls such as segmentation and micro-segmentation by reference to the outcome they are intended to achieve—namely, limiting or preventing lateral movement by threat actors that have obtained an initial foothold in a network. While segmentation can be a valuable control, micro-segmentation in a large and complex environment can be highly complex to implement and may create risks of both control gaps, where legitimate traffic must continue to flow, and operational incidents that could affect service availability. Authorities should therefore allow firms to evidence how they manage lateral-movement risk within risk appetite through appropriate combinations of network architecture, access controls, monitoring, detection, response, and other compensating controls.

Finally, the report should clarify that references to using advanced AI models for cyber and ICT risk management—such as vulnerability identification and remediation—should not be read to imply that only the most advanced models are suitable for such work. While frontier models may offer advantages for certain tasks, firms should also be encouraged to use existing models and tools where they can deliver effective outcomes within appropriate security and governance controls. Many of the principles, architectures, and security controls needed for AI-enabled vulnerability identification, testing, and remediation apply across model types. Firms that delay developing expertise in this area while waiting for access to the most advanced models may be at a disadvantage. The final report could also give greater attention to the security controls needed before firms can safely use AI for vulnerability identification, testing, and remediation, including access controls, sandboxing, monitoring, human oversight, validation, logging, and controls to prevent misuse or unintended exposure of sensitive system information.

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?

GFMA considers that the sound practices generally strike an appropriate balance. They remain broadly applicable to all forms of AI while recognizing that certain risks and controls may be particularly relevant for GenAI and agentic AI where those technologies present distinct risks, such as hallucinations, limited explainability, autonomous action, prompt injection, memory poisoning, expanded system access, and challenges in real-time human oversight. This is the appropriate approach because many core risk and control frameworks—including governance, data quality, performance monitoring, human oversight, cyber controls, and third-party risk management—continue to provide the foundation for managing risks across AI technologies.

The report describes GenAI and agentic AI as presenting distinct performance, cyber, explainability, and oversight challenges, and includes additional human-oversight considerations specifically for GenAI and agentic AI. We believe it is important to recognize that different agentic systems have different capabilities and risk profiles, and that controls and oversight should be proportionate to those attributes. Further, many agent-governance practices remain nascent and are not yet uniformly implemented across industry deployments. Descriptions of sound practices for governing agentic AI should be sufficiently concrete to support responsible implementation, while not locking in specific technical approaches too prematurely. Human oversight should also be assessed based on whether it is effective in practice, rather than whether a human review step formally exists. Given the relative novelty of agentic AI, GFMA also believes that this area will require further monitoring and may evolve over time. The final report should therefore preserve a flexible and technology-neutral approach, recognizing that implementation approaches will continue to evolve as industry practices mature. Further, while some governance practices may resemble employee-related controls, firms should focus on the capabilities needed to support accountability, traceability, and risk-based oversight, rather than framing AI agents as “synthetic employees.”

GFMA also notes that some traditional Model Risk Management practices may not be directly transposable to GenAI and agentic AI systems. A significant example is data governance since the majority of large language models leveraged by financial institutions today are sourced from a small number of third-party providers that do not provide access to their training data, testing methodologies, or validation data sets. The final report should acknowledge this structural limitation and avoid framing data governance expectations in ways that implicitly assume firms will have direct access to such information. Where traditional data governance practices cannot be applied, the report should recognize the legitimacy of compensating controls and alternative validation approaches.

Furthermore, while much of the report appropriately emphasizes pre-deployment governance, GFMA believe the final report should also address post-deployment runtime governance, which is becoming increasingly critical as agentic AI adoption matures. This includes continuous runtime monitoring and observability of AI system behavior, and the auditability of multi-step agentic workflows. Pre-deployment approval processes alone are insufficient to govern systems whose effective capabilities may change dynamically after deployment.

A further point that could be more explicitly reflected is that AI is increasingly enabling business users to develop or configure solutions more autonomously, including through mature low-code and no-code tools. While the report rightly refers to governance gaps and “shadow AI,” this broader decentralization of AI development and deployment within firms deserves greater emphasis, as it may increase both innovation capacity and governance, oversight, and documentation challenges. However, this topic should receive the same flexible and technology-neutral approach as proposed for other emerging and evolving forms of AI.

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

We appreciate that the FSB avoided a “one-size-fits-all” prescriptive approach and believe that the sound practices are broadly flexible enough to accommodate newer types of AI and responsible adoption over time. This flexibility should be reinforced in the final report by ensuring that examples, case studies, and lists of controls are clearly framed as illustrative and optional. Any attempt to create rigid, technology-specific expectations would quickly become obsolete.

The practices are generally flexible enough to accommodate continued AI development because they are framed around governance outcomes, lifecycle controls, materiality, risk assessment, and adaptability rather than prescribing specific technologies or control mechanisms. The report also usefully recognizes that the practices are not exhaustive and may need to evolve as AI technologies, use cases, supervisory expectations, and industry practices develop.

To strengthen this point, the FSB should reinforce that firms need flexibility to update controls iteratively based on experience, testing, incidents, near-misses, vendor changes, and emerging industry practices. This is particularly important for AI agents, foundation models, and other rapidly evolving technologies where industry norms are still maturing. The final report could also encourage supervisors to support responsible experimentation through voluntary sandboxes, pilots, public-private fora, and other mechanisms that allow firms to test emerging use cases in controlled environments before broader deployment.

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.

Overall, the case studies are generally useful in illustrating how different financial institutions can benefit from responsible AI adoption. However, they should be clearly framed as illustrative examples rather than preferred or expected implementation approaches.

We believe expanding case studies with more detailed examples could also help further improve their relevance. We encourage the FSB to consider including additional examples reflecting practical approaches already being developed by financial institutions, including responsible use of AI for internal knowledge management, process automation, customer interfaces, fraud detection, and model monitoring and validation. Examples could also illustrate the application of AI in governance and risk management, such as risk-based tiering of AI use cases, ongoing performance monitoring, independent model validation, governance of third-party AI services, and data management/governance including ethics. These examples would further demonstrate how responsible AI adoption can be implemented proportionately across institutions with different business models, AI maturity, and risk profiles. Additionally, we believe the FSB could consider a case study on responsible wind-down or retirement of an AI system, where industry practices are still maturing.

The case studies are helpful, but they are clearly more bank-oriented, with limited coverage of insurers, asset managers, payment firms, financial market infrastructures, markets data/analytics providers, and other non-bank financial institutions. From a banking perspective, this makes the examples easier to relate to, but also makes it harder to judge whether other types of institutions may have different needs or implementation constraints that are not yet equally reflected in the report. Additional non-bank examples could include AI-enabled trade surveillance, liquidity risk monitoring, portfolio risk analytics, claims processing, fraud detection in payments, settlement exception management, and customer service or complaint triage. These would help demonstrate that responsible AI adoption is relevant across the financial ecosystem and not primarily a large-bank issue.

On page 47, the report includes a case study on the UK Cross-Market Operational Group (CMORG) which demonstrates a public-private collaboration. GFMA believes this type of collaboration will be important as AI continues to evolve and similar examples may be helpful in demonstrating benefits.

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

The case studies provide useful and actionable insights for financial institutions. The final report should clearly express that these are examples rather than preferred or expected approaches. This is important to preserve proportionality and avoid creating de facto supervisory expectations around specific tools, governance models, or control arrangements. For example, case studies describing dedicated AI inventories, specific documentation practices, or governance structures may be appropriate for some institutions but should not be interpreted as the FSB's preferred implementation model.

Many of the case studies provide actionable insights because they describe not only the use case but also the governance or control features that supported responsible deployment. Examples include risk-based documentation, human-in-the-loop review for fraud rules, staged deployment, developmental testing before production use, benchmarking against non-AI models, red-teaming, and enhanced vendor due diligence for third-party AI.

Notably, for the case study on documentation of AI use cases, we recommend clarifying that while AI inventory management may be a helpful practice, firms should retain flexibility to determine the appropriate scope, level of detail, and implementation approach for inventories or registries.

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

The glossary is generally clear and helpful, but the FSB should ensure that key terms are aligned with widely accessible international instruments to reduce definitional fragmentation and encourage its members to adopt the same key terms and definitions. For example, in the US, the Financial Services Sector Coordinating Council's (FSSC) Artificial Intelligence Executive Oversight Group (AIEOG) published an [AI Lexicon](#) in February 2026.

GFMA also recommends that the FSB consider adding definitions (or, at a minimum, illustrative criteria) for certain terms that are used throughout the report to calibrate proportionality and governance expectations. In particular, the report refers to "shadow AI" as a governance risk but does not define the term. A definition would be helpful to distinguish between (i) unauthorized or unmanaged AI use cases, (ii) decentralized development through sanctioned low-code/no-code tools, and (iii) the use of embedded AI functionality in third-party products without clear visibility. Additionally, the report discusses "concentration risk" in the context of AI and third-party providers, but does not provide a standalone definition. A formal definition may not be needed, but clarifying whether this concept refers to provider concentration, model concentration, infrastructure concentration (e.g., cloud dependency), or a combination thereof would support more consistent implementation and monitoring.

Finally, the report repeatedly relies on "materiality" to scale the intensity of controls across the sound practices, but it does not include a definition. GFMA agrees with this approach since prescribing specific materiality thresholds may contradict with the report's goal of providing high-level and flexible guidance. The final report should clearly recognize that materiality should be assessed in a risk-based manner, taking into account the nature, scale, purpose, degree of autonomy, and potential customer or market impact, within critical business processes. For this reason, we recommend the report clarify that firms

are best suited for developing their own approaches to assessing materiality which will in turn allow them to best apply the sound practices and ensure that governance expectations are calibrated to the actual risk and significance of the use case, rather than applied uniformly across all AI applications.

Below, please find further details and suggestions relating to key terms.

Agentic AI

GFMA suggests refining the definition of agentic AI to include key characteristics that drive distinct risk considerations, such as (i) goal-directed autonomy (including the ability to plan and reason iteratively), (ii) the ability to invoke tools or external systems to carry out tasks, and (iii) the ability to execute actions with real-world impact. Controls and oversight should remain proportionate to these capabilities and the materiality of the use case.

AI vs AI System

The FSB should consider aligning its definition of “AI” with the OECD’s distinction between “AI” as a field or discipline and an “AI system” as a deployed machine-based or engineered system. The FSB’s current definition of “Artificial Intelligence (AI)” closely tracks the OECD definition of an “AI system,” but labeling it as “AI” could create ambiguity.

This approach would also ensure alignment with the Council of Europe Framework Convention on Artificial Intelligence, which uses an OECD-aligned “AI system” definition and is the first international legally binding treaty focused on AI. This supports the concept of defining “AI system” separately, rather than using “AI” to mean both the broader field and a deployed machine-based system.

The FSB should also consider adding the OECD formulation that AI systems vary in their levels of autonomy and adaptiveness after deployment.

Additionally, we recommend that the report clarify the term “AI use case” as utilized in the report and how it is distinguishable from an “AI system.” Regarding this differentiation, GFMA recommends that the FSB consider utilizing the terms “AI use case,” “AI feature,” and “AI system” in the report. A common understanding of these terms would reduce fragmentation, support comparability, and facilitate more consistent supervisory engagement. For example, clarity on the distinction between an AI use case (a specific application of AI to achieve a defined business or operational objective), an AI feature (a component capability embedded within a broader product or service), and an AI system (as defined by the OECD) would help firms determine the appropriate level of governance and documentation for each.

Explainability

The definition of explainability may need to evolve to remain meaningful in the context of Generative AI and agentic AI systems. For traditional AI models, explainability has typically been understood to refer to the ability to articulate why a model produced a particular output, often through feature importance, rule extraction, or similar techniques. However, for GenAI and agentic AI systems, it is unclear whether capabilities such as chain-of-thought reasoning traces, action logs, decision audit trails, and step-by-step workflow records should be considered sufficient forms of explainability for regulatory and governance purposes. Regulators should continue to work with industry to develop a shared understanding on whether, and under what circumstances, these forms of process-level traceability satisfy explainability requirements, and should acknowledge that the applicable standard may need to be calibrated to the nature and risk profile of the system in question, rather than applied uniformly across all AI types.

Foundational Models and General-purpose AI Models

It may be helpful to ensure closer consistency between foundation models and the increasingly used regulatory terminology around general-purpose AI models.

Open Source Models

Current terminology increasingly distinguishes between “open-source” (typically referring to source code availability), “open-weight” (model weights available, without necessarily providing training code or trading data), and other forms of AI openness. GFMA suggests the report clarify these distinctions to improve precision and avoid conflating openness with particular risk outcomes. In particular, openness should not be treated as determinative of whether a model reduces vendor lock-in or increases risk, as outcomes depend on the specific deployment context, controls, and supply chain.

Traditional AI

For the definition of traditional AI, it is helpful to detail that traditional AI does not include GenAI and beyond, but further clarity on what it does include would be useful. Additionally, we recommend clarifying that many examples of what would typically be considered traditional AI are typically not considered to be high-risk. For example, when considering the EU AI Act, the ECB calls for “expressly excluding generalized linear models (e.g. linear or logistic regressions) from the scope of the AI Act’s definition of high-risk AI systems” on the basis these “are inherently interpretable and transparent.”⁵

⁵ See “[Opinion of the European Central Bank](#)” (March 2026)