

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

Response to Consultation

Binance

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?

Question 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 would encourage the final report to place greater emphasis on responsible AI adoption as part of the control environment, not only as a source of operational efficiency. This is increasingly relevant as financial institutions face threat actors using AI to scale fraud, impersonation, cyber exploitation and other forms of financial crime, while market and operational risks continue to evolve in real time. A balanced approach should therefore recognise that AI can be both a source of new risk and, when properly governed, an important tool for detecting emerging risks and responding more quickly and effectively.

This dual use framing is consistent with FATF's work on new technologies for AML/CFT and its recent analysis of AI and deepfakes, which together recognise that AI can strengthen detection and monitoring capabilities while also being misused to support more sophisticated fraud, impersonation and financial crime typologies. This distinction is important because a framework that focuses only on the risks of AI adoption may unintentionally discourage regulated institutions from deploying tools that are necessary to respond to AI-enabled threats. In some risk domains, failure to adopt appropriate AI capabilities may itself increase risk.

This is particularly relevant for digital asset markets, where activity is global, continuous and exposed to rapidly evolving financial crime and cyber typologies. In that environment, AI can help regulated firms combine blockchain analytics, behavioural indicators and cyber intelligence to identify suspicious patterns at a speed and scale that manual processes alone may not achieve. When implemented responsibly, these tools can improve customer protection, market integrity and the effectiveness of financial crime controls.

Governance expectations should remain proportionate to the materiality of the AI use case and the potential impact on customers, markets and financial stability. Lower-risk applications supporting internal productivity should not attract the same governance expectations as AI systems capable of making or executing consequential financial decisions.

The final report could more clearly distinguish between risks arising from an AI model and risks arising from the wider AI system into which that model is embedded. In practice, AI incidents may not result solely from inaccurate model outputs. They may also arise from:

poor data lineage;

excessive permissions;

weak access controls;

inadequate logging;

inappropriate workflow design;

insufficient human review;

unclear ownership;

vendor model changes; or

inadequate monitoring after deployment.

Sound practices should therefore focus not only on model validation, but also on system architecture, operational integration, accountability, monitoring, incident response and control design.

Third-party and concentration risks also deserve continued attention. Financial institutions increasingly depend on a limited number of cloud providers, model providers, data vendors, infrastructure partners, hardware supply chains and AI-enabled software providers. Individual firms can manage these risks through due diligence, contractual protections, monitoring, contingency planning and exit strategies. However, some concentration risks arise at the sector level and cannot be fully mitigated by individual firms acting alone. These risks may require supervisory monitoring, cross-sector exercises, information sharing and public-private coordination.

This is consistent with recent international examples observing that shared models, infrastructure and cross-border dependencies may create ecosystem level risks that are not fully visible, or fully manageable, through firm level controls alone. Recent policy work in the UK also illustrates how authorities are beginning to consider both systemic oversight of critical AI and cloud providers and more standardised assurance mechanisms for third-party AI systems, while preserving the responsibility of each firm to manage its own deployment and use of those systems.

Subject to these observations, Binance agrees with the overall balance of benefits and risks in the report. The final report would be strengthened by recognising that AI is increasingly becoming embedded in important financial control and operational functions, rather than simply operating as another technology tool used by financial institutions. As AI becomes more relevant to financial crime prevention, cyber defence, market surveillance and other consequential financial activities, governance frameworks should preserve the benefits of

responsible adoption while ensuring appropriate controls over the systems, decisions and processes that may affect customers, markets or the integrity of the control environment.

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

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

We consider the sound practices broadly comprehensive and appropriately principles-based. Their greatest strength is that they seek to complement existing governance and risk management frameworks rather than create a standalone AI compliance regime. The key implementation point is that governance expectations should remain proportionate to the materiality and risk of the AI use case. A low risk internal productivity tool supporting document drafting or information retrieval should not be subject to the same governance expectations as AI systems that may influence customer outcomes, execute transactions, restrict accounts, support financial crime decisions, conduct market surveillance or perform other consequential financial activities.

This risk-based approach is consistent with recognised AI governance frameworks such as the NIST AI Risk Management Framework and ISO/IEC 42001:2023, which support structured governance across the AI lifecycle. It would be helpful for the final report to include a more explicit materiality framework to support consistent implementation. Relevant considerations could include:

the extent of customer outcomes;

the importance of the function;

the sensitivity of the data;

the degree of autonomy;

the reversibility of outputs;

the scale at which the system operates; and

the firm's ability to monitor, override or change the system.

Governance expectations should also evolve throughout the AI lifecycle. Controls that are appropriate during experimentation, testing or limited deployment may need to be strengthened as an AI system becomes embedded in business-critical processes, affects a larger number of users or assumes a greater degree of autonomy. This would help firms preserve flexibility at the innovation stage while ensuring that more consequential uses are subject to appropriate governance, monitoring and escalation.

The final report should also apply expectations by reference to the institution's role and level of control. Different governance expectations should apply depending on whether a financial institution develops, fine-tunes, procures, configures, orchestrates or simply consumes AI functionality provided by third parties. Firms should remain responsible for their own use of

AI, but expectations should reflect their role in the AI value chain and what they can reasonably test, document, monitor, explain and change.

AI inventories can also support responsible adoption, provided they remain practical, risk-based and connected to the AI lifecycle. An inventory should help firms identify material AI use cases, ownership, purpose, risk rating, delegated authority, runtime controls, data sources, third-party dependencies, key controls and escalation points. It should also operate as a living governance tool that supports continuous monitoring, periodic reassessment and retirement of AI systems where appropriate, rather than as a static implementation record.

As AI systems become more autonomous or more closely embedded in important business processes, governance should extend beyond initial approval or validation. Firms should maintain appropriate runtime monitoring, access controls, change management and post-deployment oversight to ensure that AI systems continue to operate within approved parameters.

Governance expectations should support clear accountability without requiring boards or senior management to approve every AI use case. Board and senior management oversight should focus on strategy, risk appetite, the governance framework, material AI use cases, critical dependencies and escalation. Implementation and control responsibilities should then be allocated through the firm's governance model, including appropriate ownership by the relevant business or technology function, independent challenge by risk, compliance, legal and cyber functions where relevant, and assurance through internal audit or equivalent review mechanisms.

With these clarifications, the sound practices would provide a strong global reference point for responsible AI adoption while enabling responsible innovation and avoiding disproportionate governance burdens for lower-risk applications.

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?

Question 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?

We agree with the report's approach of applying sound practices across all forms of AI. However, we encourage the final report to make it even clearer that governance should be determined primarily by the materiality of the use case rather than by the underlying AI technology. Established machine learning models, statistical models or AI-supported decision tools can create material risk where they are used in customer outcomes or other consequential financial activities. Conversely, some GenAI tools may be relatively low risk where they are used for internal drafting, summarisation, research support or knowledge retrieval with appropriate controls and no direct customer or operational effect.

The final report should therefore continue to distinguish between the type of technology and the risk of the use case. The fact that a system uses GenAI should not automatically determine its risk rating. More relevant considerations include what the system is used for,

whether it affects customers or critical processes, whether it uses sensitive data, whether outputs are relied on without review, whether the system can interact with external tools or execute actions, and whether those outputs directly influence financial decisions or operational actions.

GenAI may require targeted controls for risks such as hallucination, prompt injection, data leakage, inappropriate disclosure of confidential information and overreliance on fluent outputs. Those controls should be strongest where GenAI is used in customer outcomes, regulated or operationally sensitive contexts. At the same time, GenAI-specific controls should not be applied mechanically to other AI systems where those risks are not present. The sound practices should therefore remain modular, technology-neutral and outcomes-focused, allowing institutions to apply controls that are proportionate to the risks presented by the particular use case.

Agentic AI warrants particular attention because the risk profile changes when an AI system moves from producing outputs to taking actions. For these use cases, the relevant issue is not only model accuracy, but also authority, permissions, autonomy, reversibility and accountability. As autonomy increases, governance increasingly shifts from model validation towards runtime governance, including continuous monitoring, bounded permissions, execution controls and real-time intervention capabilities.

For higher autonomy use cases, the FSB should encourage the concept of bounded autonomy. This would mean allowing AI systems to operate only within defined purposes, permissions and limits, with stronger safeguards for higher impact actions. Relevant controls may include access restrictions, transaction or action limits, approval gates, audit trails, controlled testing, runtime monitoring, continuous logging, rollback or revocation mechanisms, incident escalation and clear accountability for outcomes.

Emerging international work, including recent UK and Singapore publications, provides a useful way to frame this issue by identifying the foundations of agentic finance as including identity, authority, accountability, execution and liability. This framing is helpful because it recognises that safe delegation depends not only on the AI model itself, but also on the surrounding infrastructure that determines who or what is authorised to act, within what limits, with what records and with what responsibility if something goes wrong.

Recent public remarks by the Bank of England's Deputy Governor for Financial Stability, Sarah Breen, in *Agents of change*, similarly highlight that agentic AI may require governance approaches that go beyond traditional human review, including clear accountability and safeguards around autonomous activity in payments, markets and operational resilience.

This is also reflected in the UK Financial Services AI Adoption Plan, which identifies agentic payments as a practical near-term use case for autonomous financial systems and highlights the need for standards on legal accountability, agent identity, verification and machine-to-machine authentication. These issues reinforce the importance of treating agentic AI as a question of delegated authority, execution controls and accountability, rather than only as a model governance issue.

Digital asset markets provide a useful illustration of this transition, as agentic systems may increasingly interact with wallets, smart contracts, stablecoins, tokenised assets, payment flows or compliance workflows. The final report should therefore encourage firms and supervisors to focus on the actions an AI system can take, the systems it can affect and the safeguards around its authority.

The emergence of agentic AI should not be viewed as creating an entirely new regulatory paradigm. Rather, it reinforces the importance of applying established principles, including accountability, operational resilience, segregation of duties, least-privilege access and effective governance, to increasingly autonomous financial activities.

The appropriate intensity of governance should be determined by the materiality of the financial activity, the authority delegated to the AI system and the effectiveness of the surrounding control environment, not by whether the underlying technology is labelled as machine learning, generative AI or agentic AI.

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

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

Binance considers the proposed sound practices sufficiently flexible in principle. Their long-term effectiveness, however, will depend on remaining anchored in durable governance outcomes rather than current model terminology, deployment methods or specific technological architectures. A principles-based and technology-neutral approach is more likely to remain effective as AI capabilities continue to evolve than guidance tied to particular model types or implementation approaches.

Flexibility will be particularly important for financial institutions that operate across jurisdictions. Firms are already navigating AI principles, existing operational resilience, outsourcing, cyber resilience, model risk management, customer protection, financial crime and sector-specific financial regulation. The FSB's sound practices can add value by promoting common expectations and supervisory convergence, but they should remain interoperable with existing international frameworks and domestic regulatory regimes. This would also align with the OECD AI Principles, which emphasise interoperable, flexible and risk-based governance capable of adapting across jurisdictions and over time.

We also support a tiered approach to responsible AI governance. Firms should be encouraged to classify AI-enabled activities according to risk and materiality, with controls that increase in intensity as the potential impact of the system increases. This would allow the same governance framework to apply to internal productivity tools, traditional machine learning, GenAI, agentic AI and future model architectures, while still calibrating expectations to the risk presented by the use case. Responsible AI governance should itself be adaptive. As AI capabilities evolve, governance frameworks should support periodic reassessment of risks, controls and supervisory expectations rather than relying on static implementation assumptions established at the point of deployment.

The final report should also recognise that responsible AI adoption is a continuous process, not a one-time implementation exercise. AI systems may change after deployment through model updates, data drift, changes in user behaviour, new threat typologies, vendor updates or changes in delegated authority or integration with additional systems. Firms should therefore maintain processes for ongoing monitoring, periodic reassessment, incident review and control updates, including retirement or replacement where an AI system no longer performs as intended.

We would also encourage the FSB to promote supervisory cooperation and industry information sharing, including where appropriate through mechanisms for sharing AI incidents, near misses, emerging typologies and lessons learned. AI-related risks, including model vulnerabilities and shared dependencies, cyber, fraud, scam, third-party and concentration risks, may evolve quickly and affect multiple firms at the same time. This is especially relevant where risks arise from common vendors, shared infrastructure, or threat actors using similar AI-enabled techniques across the financial sector.

The final report could also encourage authorities to consider how AI-related vulnerabilities may be monitored at system level, including through targeted data collection, trend analysis, cross-border supervisory cooperation, supervisory intelligence, sector exercises and operational resilience testing. Recent policy work, including the UK Financial Services AI Adoption Plan, similarly points to the value of cross-sector dependency mapping, horizon scanning and scenario testing where AI-related risks may transmit across finance, cloud, digital infrastructure and other critical sectors.

Finally, the report should avoid discouraging controlled experimentation. Sandboxes, proofs of concept, pilots, limited deployments, red teaming and staged rollouts can help firms understand risks before wider implementation. A responsible framework should encourage this type of controlled learning, especially where the alternative is informal or ungoverned adoption by business users.

Accordingly, the final report should preserve flexibility not only in its drafting, but also in how the practices are implemented by supervisors and firms in a manner that remains proportionate, technology-neutral and responsive to future developments.

- 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.**

Question 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 provide useful practical illustrations of the proposed sound practices. However, the final report would be strengthened by including additional examples drawn from non-bank financial institutions. Many of the fastest-moving AI use cases are emerging in payment institutions, digital asset firms, financial market infrastructures and other technology-enabled financial service providers. Including these examples would improve

the report's relevance across an increasingly diverse financial ecosystem and demonstrate that the sound practices can be applied consistently across different business models.

We suggest that the FSB consider including the following additional case studies. These examples deliberately focus on AI applications that strengthen market integrity, customer protection, operational resilience, financial crime prevention and regulatory compliance, as these represent areas where responsible AI adoption can deliver clear public policy benefits while remaining consistent with existing regulatory objectives.

AI-supported financial crime monitoring in digital asset markets. A virtual asset service provider may use AI together with blockchain analytics, behavioural analytics and external threat intelligence to identify suspicious transaction patterns across wallets, chains and platforms, including structuring, mule activity, cross-chain layering, exposure to sanctioned or high risk addresses and rapid movement of funds following account takeover. Relevant controls would include clear ownership, data quality controls, alert validation, human escalation, quality assurance, false positive monitoring and audit trails.

AI-supported scam and fraud prevention. A digital asset platform or payment firm may use AI to detect unusual behaviour before funds are transferred, including indicators of AI-generated impersonation, deepfake-enabled scams, changes in device or login behaviour, unusual withdrawal timing, destination wallet risk or activity inconsistent with the customer's history. Controls should support proportionate intervention, escalation for high risk cases, review procedures, false positive and false negative monitoring, and urgent intervention where customer harm is likely.

AI for cyber defence and operational resilience. A nonbank financial institution may use AI to detect malicious traffic, abnormal API activity, phishing infrastructure, supply chain compromise, suspicious privileged access or abnormal system behaviour. This type of use case is particularly important for exchanges, custodians and other firms that operate critical digital infrastructure or safeguard customer assets. Responsible adoption would require AI-assisted vulnerability discovery, secure deployment, access controls, monitoring for adversarial manipulation, incident escalation and integration with existing cyber and ICT risk frameworks.

AI-assisted regulatory compliance. A financial institution may use AI to support regulatory horizon scanning, regulatory mapping, obligations management or regulatory reporting. These tools can help firms identify relevant regulatory changes, map obligations to internal policies and controls, and improve consistency in compliance workflows. Responsible adoption would require clear ownership, source validation, human review of regulatory interpretations, audit trails, version control and escalation where outputs may affect regulatory submissions or compliance decisions.

AI for market surveillance in digital asset and tokenised markets. A trading platform, market infrastructure provider or tokenisation platform may use AI to identify abnormal trading behaviour, coordinated activity, liquidity manipulation, wash-like activity, cross venue anomalies or other market abuse indicators. This would be consistent with the broader capital markets focus on using AI to support surveillance and compliance functions while preserving market integrity. This illustrates how AI may enhance, rather than replace,

existing market surveillance frameworks by improving the speed and effectiveness with which firms identify potential misconduct for human investigation.

Agentic AI in controlled payment or wallet environments. As AI agents become more capable, financial institutions may explore controlled use cases where an AI system can initiate or prepare actions within defined boundaries. In digital asset markets, emerging use cases may involve interaction with wallets, smart contracts, stablecoins or tokenised assets. Responsible adoption would require bounded autonomy, including permission limits, action limits, runtime monitoring, approval gates for high-impact actions, audit trails, testing in controlled environments, revocation mechanisms and clear accountability for outcomes.

These examples would help demonstrate how responsible AI adoption can support financial crime prevention, scam disruption, customer protection, market integrity and operational resilience outside the traditional banking sector. They would also show that the sound practices can be applied flexibly across different business models, provided the controls are calibrated to the risk profile, delegated authority and the degree of institutional control over the AI-enabled activity.

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

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

The case studies are helpful in illustrating how the sound practices may be applied in practice. Their implementation value, however, could be strengthened by making the underlying governance decisions, control choices and implementation considerations more explicit. Future case studies could be strengthened by using a common implementation template. This could identify the business objective, the AI capability, the materiality assessment, governance and accountability arrangements, key controls, oversight model, monitoring approach, residual risks and lessons learned. A consistent structure would help firms understand not only what AI can be used for, but how the relevant sound practices may be applied in practice without turning the case studies into prescriptive templates.

The case studies should also illustrate different models of human oversight. In some cases, human approval may be appropriate, particularly where the AI output directly affects customers or legal rights. In other cases, such as cyber defence, fraud detection or real time transaction monitoring, a human-in-the-loop model may be more effective, provided there are clear limits, escalation triggers, audit trails, monitoring and authority to override. Overall, the appropriate oversight model should reflect the materiality of the activity, the speed at which decisions must be taken and the effectiveness of alternative governance controls. The final report should avoid creating the impression that meaningful oversight always requires manual approval of every AI output.

The case studies could also address the trade-offs that firms often need to manage when adopting AI responsibly. For example, a more explainable model may not always be the most effective model for detecting complex fraud patterns, while fraud detection systems may deliberately favour higher recall at the expense of increased false positives. Similarly, a more automated cyber defence tool may improve response time but require stronger

permissioning and monitoring. Case studies should encourage firms to identify, document and manage these trade-offs, including the balance between customer experience, operational resilience and financial crime effectiveness.

We also recommend including examples of lower-risk AI adoption, such as internal productivity, software engineering assistance, translation, knowledge retrieval, regulatory research and analytics. This would help firms understand how to apply proportionate controls without creating unnecessary friction for lower-risk use cases.

Finally, the case studies should remain illustrative examples of good practice rather than operate as safe harbours. The same use case may present different risks depending on the firm's business model, data, deployment architecture, third-party dependencies, customer base and control maturity. Where appropriate, case studies could explain not only which controls were implemented, but why those controls were considered proportionate for the particular use case. Case studies could also demonstrate how governance evolves after deployment, including periodic reassessment, incident response, model updates and retirement where appropriate.

With these enhancements, the case studies would provide more actionable insights and help financial institutions translate high-level governance principles into practical implementation decisions across a wide range of business models and AI use cases.

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

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

The glossary provides a useful common vocabulary and is broadly aligned with current international practice. We suggest a small number of refinements that would improve implementation consistency and help the glossary remain relevant as AI technologies continue to evolve.

First, the definition of AI can remain broad and aligned with the OECD formulation of AI as a machine-based system that infers how to generate outputs. The practical issue is not to narrow the definition, but to avoid over-inclusion when firms identify AI use cases for governance purposes. A clearer distinction between AI systems and traditional deterministic software, rules-based automation or workflow orchestration would help firms maintain practical inventories and apply controls proportionately.

Second, the glossary should distinguish between an AI model and an AI system. Many operational, security and governance risks arise from the deployment architecture surrounding an AI model, including retrieval components, orchestration logic, permissions, APIs, monitoring and human workflows, rather than from the model itself. This distinction is important because many risks arise from how a model is deployed and integrated, not only from the model itself.

We support the proposed definition and suggest that it could be expanded to recognise that the principal governance consideration is not simply autonomous decision-making, but the

ability of a system to interact with tools, external systems and business processes in ways that can produce consequential financial or operational outcomes. The relevant risk driver is whether the system can use tools, interact with external systems, execute workflows, modify records, initiate transactions or affect customer or operational outcomes without step by step human instruction.

The glossary could also recognise that human oversight may take different forms depending on the use case, risk, autonomy, time sensitivity and reversibility of the AI output. Meaningful human oversight may therefore take the form of ex-ante approval, human-in-the-loop supervision, runtime monitoring or post-event review, depending on the materiality, reversibility and time sensitivity of the activity.

Fourth, the glossary should distinguish explainability from transparency. Transparency may relate to governance, documentation, disclosure, communication and accountability. Explainability relates more specifically to whether relevant stakeholders can understand the basis for an AI output or decision at an appropriate level of detail. The glossary could also clarify that explainability should also be proportionate to the audience and the purpose for which an explanation is required. For example, regulators, auditors, model validators and customers may require different forms of explanation. In certain financial crime, sanctions, cybersecurity and market surveillance contexts, excessive disclosure of model logic may increase the risk of evasion, gaming or adversarial behaviour.

Finally, the FSB could consider adding definitions for “AI system lifecycle”, “material AI use case”, “bounded autonomy”, “delegated authority”, “contestability” and “shadow AI”. As a suggestion, we propose the following:

“AI system lifecycle” should refer to the end-to-end lifecycle of an AI system, rather than only training or deployment.

“Bounded autonomy” could refer to an AI system operating within defined permissions, limits, approval gates and monitoring arrangements.

“Contestability” could refer to mechanisms that allow affected users, customers or internal stakeholders to review, appeal, correct or challenge AI-supported outcomes.

“Delegated authority” could refer to the scope of actions an AI system is authorised to perform within defined operational, legal or technical limits.

“Material AI use case” could refer to an AI system that may materially affect customers, markets, critical operations, financial crime controls, cyber resilience or financial stability.

“Shadow AI” could refer to the use of AI capabilities outside approved governance, procurement, security or risk management arrangements.

These refinements would make the glossary more useful for implementation and better able to accommodate future developments in AI, including GenAI, agentic systems and AI tools embedded into third-party financial infrastructure.

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

Re: Sound Practices for Responsible Adoption of Artificial Intelligence

Binance welcomes the opportunity to respond to the Financial Stability Board's consultation on Sound Practices for Responsible Adoption of Artificial Intelligence.

We support the FSB's work to develop a common set of sound practices for financial institutions adopting, using and innovating with artificial intelligence (AI). AI is already improving the quality, speed and scalability of financial services and, at the same time, it can create or amplify risks where systems are poorly governed, inadequately tested, relied on without appropriate challenge, deployed without clear accountability or connected to critical processes without appropriate controls.

This is consistent with the FSB's prior work on AI and financial stability. In our view, the proposed sound practices can usefully build on that work by helping firms manage AI-related vulnerabilities while preserving space for responsible innovation.

AI is also becoming increasingly embedded within core financial infrastructure rather than operating solely as a standalone productivity technology. Financial institutions are deploying AI across financial crime prevention, cyber defence, market surveillance, customer support, software engineering and operational resilience, while emerging use cases increasingly involve interaction with payment systems, digital wallets and other elements of financial market infrastructure. Sound governance should therefore focus not only on AI models, but on AI-enabled financial activities and the control environments within which they operate.

AI is not a single technology, deployment model or risk category. Governance expectations should therefore reflect the way an AI system is actually used, the impact it may have and the level of control the institution has over it. An internal productivity or support tool should not be treated in the same way as a system that may affect customers or the integrity of the control environment. Similarly, expectations should recognise the difference between firms that develop proprietary AI systems, firms that configure or integrate third-party tools and firms that use AI-enabled services provided by vendors or other regulated entities.

This distinction is particularly important for virtual asset service providers, which operate in global, continuous and data-rich markets exposed to fast-evolving financial crime, scam, cyber and market integrity risks. In that environment, responsible AI can help regulated firms detect patterns and respond at a speed and scale that manual controls alone may not achieve. Properly governed, AI should therefore be understood as part of the control environment itself, helping regulated firms respond to increasingly sophisticated fraud, cyber and financial crime threats while remaining subject to appropriate governance and oversight.

Against this background, our response focuses on how the final report can preserve proportionality, recognise different institutional roles and control levels, and reflect the role of responsible AI in strengthening financial crime, cyber and market integrity controls.

Our responses to the consultation questions are set out in Annex A below.

Yours faithfully,

The Binance Team

Conclusion

Binance supports the FSB's proposed sound practices and welcomes the balanced, proportionate and technology-neutral approach. The final report can provide a valuable global reference point for responsible AI adoption across the financial sector, provided it remains practical, risk-based and adaptable to different business models and future developments.

The most important refinement is to ensure that governance expectations remain anchored in the materiality of the financial activity, the authority delegated to the AI system and the effectiveness of the surrounding control environment, rather than the label attached to the underlying technology. This would help avoid both under-governance of high-impact AI use cases and unnecessary friction for lower-risk use cases, while preserving a technology-neutral framework that remains relevant as AI capabilities continue to evolve.

The final report should also recognise that responsible AI adoption can strengthen financial sector resilience. In areas such as fraud prevention, financial crime compliance, cybersecurity, market surveillance and scam disruption, regulated institutions increasingly require advanced AI capabilities to respond effectively to sophisticated and AI enabled threats. A sound framework should therefore support the responsible deployment of AI as part of the control environment, while ensuring that these capabilities remain subject to appropriate governance, testing, oversight, monitoring and accountability.

Looking ahead, AI is likely to become increasingly embedded within core financial infrastructure alongside tokenised assets, digital forms of money and other programmable financial services. Governance frameworks should therefore remain focused on enduring financial outcomes, including customer protection, market integrity, operational resilience and financial stability, rather than particular technologies or deployment models. By remaining principles-based, technology-neutral and proportionate, the FSB's sound practices can provide a durable foundation for responsible AI adoption across the global financial system as these technologies continue to evolve.