

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

Response to Consultation

B3 S.A. - Brasil, Bolsa, Balcão

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?

B3 S.A. — Brasil, Bolsa, Balcão (B3), a company incorporated under the laws of Brazil, is authorized as an operator of securities, commodities, and futures exchanges and of an OTC market, as well as clearinghouse, foreign exchange clearinghouse, central depository, and trade repository. B3 welcomes the opportunity to comment on this Consultation Report as a Financial Market Infrastructure (FMI). B3 has implemented a governed enterprise AI platform and maintains a corporate inventory of AI initiatives, supporting responsible adoption at scale under common governance, risk, technology, cyber, legal, and compliance standards.

As an FMI, B3 frames its responses from the perspective of financial market infrastructures. Because FMIs operate at the core of the financial system and interact with a broad range of market participants, the practices, benefits, and risks discussed below are directly relevant to responsible AI adoption by financial institutions more broadly, and B3's comments are intended to support that shared objective.

B3 broadly agrees with the report's characterization of the benefits and risks of AI adoption:

- Benefits identified in the report — operational efficiency, enhanced risk monitoring, fraud and anomaly detection, productivity gains from Generative AI, and improved cyber resilience.
- Risks identified in the report — model risk, data governance failures, explainability challenges, third-party concentration, cyber risks, and inadequate governance that may lead to the use of unapproved AI tools ("shadow AI"), among others.

Drawing on its experience as a regulated and systemically important FMI, B3 suggests the following additions.

Additional benefits B3 suggests reflecting in the report

- Platform-based responsible adoption at scale. A governed enterprise AI platform enables decentralized innovation while maintaining consistent controls over models, data, access,

lifecycle management, and operational resilience. For FMIs, this approach strengthens market integrity, financial stability, and confidence in critical market services.

- Agentic infrastructure as a strategic capability. For FMIs, agent-based automation can strengthen settlement, reconciliation, compliance monitoring, anomaly detection, and internal controls, while enabling future API-based services for market participants.

- Multidisciplinary governance as an enabler of speed. Tactical and strategic governance forums, combined with defined institutional roles and responsibilities (e.g., AI Owner, AI Builder, AI Ambassador), enable agile prioritization without sacrificing oversight.

- Workforce upskilling as a risk mitigant. Structured training programmes for executives, managers, and builders — including governance and digital-risk modules — materially improve responsible use and reduce ungoverned experimentation.

Additional risks B3 suggests reflecting in the report

- Cost and capacity governance. At high scale — where large volumes of tokens are consumed and numerous agents operate — inefficient model selection and uncontrolled agent proliferation can create operational, financial, and sustainability risks.

- Reputational and ethical decision-making risks. As a complement to the risks related to "Ineffective Data Governance and Poor Data Quality" and "Consumer Protection and Market Conduct Risks," B3 suggests considering reputational and ethical decision-making risks. Market participants, regulators, customers, and third parties may question the reliability of an institution if AI-related incidents reveal failures of systems and controls in governance, risk management, or internal controls. In addition, AI models and systems may conflict with organizational values and with principles of market integrity and fairness.

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

Overall, yes. The 12 sound practices provide a practical, well-structured menu, and B3 agrees with them. They enable financial institutions to structure responsible AI adoption in a clear and proportionate manner.

By way of illustration, B3's five-pillar governance approach is well aligned with the sound practices — Data & AI Governance; Risk, Ethics & Compliance; Architecture & Technology; People, Culture & Competencies; and Cyber & Resilience — as follows:

- SP 1–4 (strategic direction and oversight, accountability, risk framework, organizational adaptability) — can be operationalized through a centre-of-excellence structure, layered governance forums (tactical, operational, strategic), and executive training on AI governance.

- SP 5–10 (materiality, selection, data governance, explainability, performance, human oversight) — can be operationalized through a structured solution-intake process, risk-based classification, a homologated model stack, and an agent typology with differentiated guardrails.

- SP 11–12 (cyber/ICT and third-party risk) — this is reflected in controls such as restrictions on non-approved models, prompt and access controls, data security monitoring, and adversarial testing practices for selected AI systems.

To enhance clarity and facilitate implementation, B3 suggests the following improvements:

- More explicit guidance on proportionality thresholds — e.g., clearer criteria for determining material or critical AI use cases in FMIs (settlement, clearing, surveillance, market data, and critical operational services).

- Clearer mapping to existing frameworks — indicating how these practices complement model risk management, operational resilience, and third-party risk standards already adopted by FMIs.

- Practical templates — inventory fields, risk-assessment checkpoints, and documentation minimums by lifecycle stage would help institutions operationalize the practices faster.

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?

Yes, the balance is broadly appropriate. The report correctly treats the practices as technology-neutral while noting that certain practices — such as human oversight, explainability, performance testing, and cyber/ICT risk — are more relevant for GenAI and agentic AI.

B3's experience confirms that this balanced approach works in practice:

- For GenAI, B3 applies grounding techniques (e.g., retrieval-augmented generation), prompt controls, human review of material outputs, and model homologation.

- For agentic AI, B3 implements differentiated guardrails by agent type, role-based access control per agent, and human-in-the-loop for higher-risk workflows.

As agentic AI matures, B3 recommends that the report provide further guidance on agentic AI, including:

1. Levels of autonomy and permitted action boundaries;
2. Governance of tool execution and API integrations;
3. Identity, authentication, authorization, traceability, and accountability for agent actions;
4. Oversight of multi-agent orchestration; and
5. Pre-deployment testing and post-deployment monitoring.

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

Yes. The non-binding, menu-based design and the proportionality principle provide appropriate flexibility for evolving technologies.

B3's AI programme demonstrates this flexibility:

- Evolution from a centralized model (2023–2024) to area-owned AI products (2025–2026) under common institutional standards.
- Expansion from chat and LLM usage to agent creation, service APIs, integration infrastructure, and planned multi-agent orchestration.
- Risk-based guardrails that can be updated as agent typologies and regulatory expectations evolve.

To preserve flexibility and long-term relevance, B3 recommends that the FSB:

- Commit to periodic review and refresh of the practices and the glossary, ensuring that the framework remains aligned with the rapid evolution of AI technologies and related risks.
- Encourage institutions to maintain organizational adaptability (SP 4) through continuous learning, red-teaming exercises, and maturity assessments to strengthen AI governance and risk-management capabilities over time.
- Maintain technology-neutral principles while providing illustrative implementation reference patterns that can evolve alongside technological developments — for example, human-in-the-loop mechanisms, retrieval-augmented generation (RAG) grounding techniques, and compensating controls for AI systems with limited explainability.

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 are valuable and actionable, but predominantly focused on banking institutions. Coverage of FMIs, market infrastructures, and regulated non-bank entities is limited relative to their systemic importance.

B3 recommends adding FMI and non-bank case studies and offers the following illustrative examples from its programme.

Case Study A — Governed GenAI platform at FMI scale

B3 has implemented a governed internal AI platform that integrates approved models and common controls for access, lifecycle management, data protection, monitoring, and oversight. This approach allows decentralized innovation while preserving institutional governance and risk-management standards.

Insight: FMIs can scale GenAI responsibly through a central platform that enables decentralized innovation with unified controls.

Case Study B — Agentic AI for corporate risk assessment

B3's corporate risk function is developing agentic AI capabilities to identify, assess, and treat operational, technology, process, and new-initiative risks. The design includes human oversight, scalability beyond the risk function, and potential reuse by other governance forums when evaluating new AI proposals.

Insight: Agentic AI can strengthen governance itself when deployed with clear boundaries and expert review.

Case Study C — AI-native internal applications

B3 built an internal data-governance application with AI-assisted requirements generation, front-end development, and data support, reducing operational effort and enabling autonomous updates, with ownership and accountability maintained in a corporate AI inventory.

Insight: FMs can use GenAI to accelerate internal tooling while preserving centralized lifecycle governance.

B3 encourages the FSB to add similar examples for central securities depositories (CSD), central counterparties (CCP), securities settlement systems (SSS), trade repositories (TR), exchanges, and regtech and data-infrastructure providers.

Case Study D — AI and machine learning for market surveillance

BSM Supervisão de Mercados, self regulatory organization responsible for the supervision and inspection of the organized markets managed by B3, has progressively incorporated AI and machine learning across its supervisory processes, illustrating how surveillance activities at market infrastructures can benefit from responsible AI adoption:

- Analysis of historical patterns and behaviours — since 2022, supervised and unsupervised machine-learning models have been used in the construction of surveillance alerts.
- Transcription and analysis of recorded conversations — in 2023, a pilot was launched using audio-transcription tools to analyse telephone recordings obtained from market participants, supporting surveillance for potential misconduct; the generation of written inquiries based on alert activity and surveillance data is a potential future use.
- Reduction of false positives — in 2024, supervised machine-learning models were adopted in the continuous improvement of existing surveillance alerts.
- Meeting summaries, report drafting, and compliance documentation — in 2024, meeting-summary practices began through corporate productivity tools, while report drafting and compliance documentation are in development and testing together with business areas.
- Extraction of data from complex, unstructured documents — in 2024, a first pilot transformed unstructured documents obtained from market participants (for example, operating-parameter rule documents) into structured formats for easier analysis and database population.

- Software-development assistance — since 2024, AI coding-assistant tools have supported software development.

In addition, governed conversational-AI “rooms” support business analysts in the market-surveillance area with insights for investigating potential irregularities — price manipulation, artificial market conditions, coordinated trading, and anti-money-laundering. Examples include:

- Daytrade Ranking – triage and analysis — automates the triage of the “Daytrade Ranking” monitoring alert, selecting the investors that warrant closer analysis based on accuracy, frequency, and account-type criteria, thereby focusing analyst effort on the cases of highest potential risk and making monitoring more efficient and targeted.

- Registration data — queries participant and investor registration data, prioritizing the most recent records, and acts as a curated data room of active, valid accounts registered at B3, enabling fast and consistent queries for the area.

- Relationships between investors (listed markets) — identifies non-obvious relationships between investors based on registration links — family, economic group, same manager, same administrator group, same postal code, same e-mail, and links with the issuer — highlighting relationships in common among investors involved in cases under analysis by the market-surveillance area.

Insight: Market infrastructures and their surveillance functions can materially strengthen market integrity and anti-money-laundering defences by combining supervised and unsupervised machine learning with governed conversational AI, while keeping human analysts in control of investigative decisions.

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

Yes. The case studies effectively illustrate proportionality, human oversight, compensating controls for explainability gaps, and phased deployment — all of which are directly relevant to and consistent with B3's own practice.

The most actionable elements identified by B3 were:

- Human oversight before operational impact — approval of fraud rules, review of material outputs, and validation of agent-executed tasks.

- Compensating controls where explainability is limited — RAG grounding, benchmarking against non-AI models, and enhanced monitoring.

- Phased rollout and sampling-based monitoring for agents and LLMs.

- Cross-functional involvement — engaging compliance, risk, and subject-matter experts in design and validation.

To further enhance actionability, B3 recommends that each case study include a concise implementation checklist covering key elements such as materiality classification, minimum

data-governance requirements, oversight arrangements, testing and validation plans, third-party dependencies, and contingency measures.

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 consistent with industry usage. Definitions of GenAI, agentic AI, foundation models, LLM hallucinations, RAG, data/model drift, and explainability are particularly useful.

B3 suggests adding or clarifying the following terms:

- AI agent / agentic workflow — distinguish prompt-based assistants from autonomous agents with tool execution.
- Tool calling / function calling — relevant for operational-automation risks.
- Model Context Protocol (MCP) / AI gateway — increasingly used for governed integration of AI with enterprise systems.
- Shadow AI — use of unapproved AI tools outside institutional controls.
- AI red teaming — structured adversarial testing of AI systems.
- LLMOps / agent lifecycle management — complement to traditional MLOps.

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

No additional comments.