



FINANCIAL
STABILITY
BOARD

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

Response to Consultation

UBS

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?

UBS broadly agrees with the benefits and risks of AI adoption as described in the report. UBS considers the report to have given a fair and balanced treatment of the benefits and risks and also welcomes recognition that not adopting AI itself may create risks.

UBS 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.

UBS supports additional risk considerations as outlined in the trade association responses, in particular GFMA, IIF and BPI, including on fragmentation risks, AI-specific considerations, cyber and ICT risk including incident reporting, third-party risk and macroprudential risk.

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

UBS considers the sound practices to be broadly comprehensive and sufficiently clear. FIs are already responsibly adopting AI, and AI is already strictly regulated by regimes including privacy, data governance, and operational risk regimes. As such, it is not a question of “enabling” responsible AI adoption as much as encouraging alignment between FI supervisors, and of fostering those sound practices over time.

UBS agrees that responsible adoption of AI can be achieved most effectively by integrating AI governance into existing firm-wide governance and risk management frameworks, rather than by creating parallel AI-specific regimes or rules.

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. 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.

UBS support GFMA/IIF calls encouraging 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. The final report should preserve flexibility in how effective oversight is achieved, while maintaining clear accountability at the firm level.

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 customer-impacting use cases.

We also think it is important to distinguish between requirements that apply to an AI system and those that genuinely need to be repeated for each use case (deployment, use and application of an AI System to achieve specific business objectives), particularly where requirements would be duplicative or disproportionate for lower risk use cases. Some of the practices (5, 6, 12) could be re-phrased to account for this important difference as not all recommendations may be relevant at AI use case level (e.g. risk tiering).

With regard to Practice 5, an approach to assessing the materiality and risk of AI can also be achieved by effective and systematic risk assessments in line with the governance and controls of the underlying risk taxonomies. A single overarching Inherent Risk Rating (IRR) may support governance, triage, and portfolio oversight, but it cannot replace taxonomy-specific risk assessments nor meaningfully determine the depth of 2nd Line of Defense reviews.

At UBS AI risks span over multiple risk taxonomies which ensure that risks are identified, measured and controlled in line with the underlying control frameworks, such as model risks, data quality issues, cyber risks, or third-party dependency. UBS would like to highlight that i) it is operating upon globally consistent frameworks and controls with clear roles and responsibilities across all three Lines of Defense ("LoD"); and ii) it is important that the framework and practices remain based on the practice of proportionality (e.g. for low-risk cases) and practicality (e.g. foundational model provider). In the spirit of this, UBS would appreciate regulatory efforts to promote global consistency and convergence in evolving regulation.

For addition considerations, including related to AI-related cyber and ICT risks, third-party risk management, governance and literacy, we refer to the aforementioned trade association responses to which we contributed and which we support.

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?

UBS 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.

UBS is broadly supportive of the AI Shared Responsibility model. This type of model would be useful in delineating accountability among AI vendors, their upstream suppliers, and the FIs that deploy the technology. This model also formally acknowledges that FIs cannot and should not be solely responsible for risks embedded deep within the AI supply chain.

For further details, including on GenAI and agentic systems, please refer to aforementioned trade association responses to which we contributed and which we support.

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.

For further details please refer to aforementioned trade association responses to which we contributed and which we support.

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.

For further details, including on case studies, please refer to aforementioned trade association responses to which we contributed and which we support.

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. 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.

For further details please refer to aforementioned trade association responses to which we contributed and which we support.

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

As a global standard-setting body, the FSB is positioned to address regulatory fragmentation on key AI concepts. 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.

As mentioned previously, we think it is important to distinguish between requirements that apply to an AI system and those that genuinely need to be repeated for each use case (deployment, use and application of an AI System to achieve specific business objectives), particularly where requirements would be duplicative or disproportionate for lower risk use cases. Some of the practices (5, 6, 12) could be re-phrased to account for this important difference as not all recommendations may be relevant at AI use case level (e.g. risk tiering).

For further details, including on definitions of AI systems and AI use cases, concentration risks, shadow AI, materiality, etc., please refer to aforementioned trade association responses to which we contributed and which we support.

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

Secretariat to the Financial Stability Board
Bank for International Settlements
Dr. John Schindler
Centralbahnplatz 2
4002 Basel

22 July 2026

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

Dear Dr. Schindler

UBS welcomes the opportunity to provide feedback to the Financial Stability Board (FSB) on its consultation report on sound practices for responsible adoption of artificial intelligence (AI). Following the publication of your report UBS has performed a comprehensive review to assess both our internal readiness and public policy aspects. As a global systemically important bank headquartered in Switzerland, UBS has a great interest in the clarity, design, and operability of AI frameworks globally. From our June bilateral exchange, we understand the FSB is not only interested in our views on the principles, but even more so in how UBS is adopting and deploying AI. We will cover both aspects in this letter.

Response to the consultation

Financial institutions have significant experience employing artificial intelligence and machine learning (AI/ML) tools in their work over an extended period of time. AI is subject to existing regulation, including data protection, model risk management, consumer protection, operational resilience, third party risk, etc. The global landscape for AI regulation is complex and rapidly evolving as various regulators are introducing new or enhanced AI laws and regulations. Due to its experience in handling sensitive data, cloud technology, and the employment of traditional uses of AI, the financial services industry is well-positioned to contribute to broader policy development for AI transparency and trust-building. We think regulators and the industry share common goals, i.e. to safeguard consumers (data) and provide fair, safe, legal, ethical and high-quality services. Although novel forms of AI, incl. GenAI, don't create new risks, they potentially exacerbate existing risks (e.g. across model risk, cyber and information security, third party risk, data privacy and ethics, data management, etc.). We also see an increase in the use of AI by regulators themselves, in some cases even combined with requests for access to supervised entities data and systems. While this may increase supervisory efficiency and effectiveness, it also creates risk and additional vulnerabilities. As result, regulators themselves should follow a consistent set of principles, which are ideally aligned with those applied to financial institutions, to avoid loopholes and reduce risks across institutional boundaries.

In this context, UBS appreciates the structured and practice-based direction provided by the FSB. UBS is either already aligned to or is in the process of aligning with FSB's direction of travel. We support the high-level practices, which are generally well founded and appropriate, as well as the FSB's balanced assessment of AI's benefits and risks. The benefits and risks described in the report are comprehensive and consistent with observed practices across financial institutions.

In addition to this letter, we have focused our response effort on contributing to the responses by the International Institute of Finance (IIF), the Global Financial Markets Association (GFMA), the Institute of International Bankers (IIB), the Bank Policy Institute (BPI), as well as other trade associations of which UBS is a member. We support their responses, which reflect our views as well as those of the broader industry. Three elements also included in those responses are of particular importance to UBS and are outlined in the following paragraphs.

First, UBS agrees with all sound practices, with some suggested improvements to Practice 5 ('Materiality and risk assessment'). Generally, we think it is important to distinguish between requirements that apply to an AI system and those that genuinely need to be repeated for each use case (deployment, use and application of an AI System to achieve specific business objectives), particularly where requirements would be duplicative or disproportionate for lower risk use cases. Some of the practices (5, 6, 12) could be re-phrased to account for this important difference as not all recommendations may be relevant at AI use case level (e.g. risk tiering). With regard to Practice 5, an approach to assessing the materiality and risk of AI can also be achieved by effective and systematic risk assessments in line with the governance and controls of the underlying risk taxonomies. A single overarching Inherent Risk Rating (IRR) may support governance, triage, and portfolio oversight, but it cannot replace taxonomy-specific risk assessments nor meaningfully determine the depth of 2nd Line of Defense reviews.

Second, at UBS AI risks span over multiple risk taxonomies which ensure that risks are identified, measured and controlled in line with the underlying control frameworks, such as model risks, data quality issues, cyber risks, or third-party dependency. UBS would like to highlight that i) it is operating upon globally consistent frameworks and controls with clear roles and responsibilities across all three Lines of Defense ("LoD"); and ii) it is important that the framework and practices remain based on the practice of proportionality (e.g. for low-risk cases) and practicality (e.g. foundational model provider). In the spirit of this, UBS would appreciate regulatory efforts to promote global consistency and convergence in evolving regulation.

Third, UBS agrees that responsible adoption of AI can be achieved most effectively by integrating AI governance into existing firm-wide governance and risk management frameworks, rather than by creating parallel AI-specific regimes or rules.

UBS approach to AI

UBS continues to focus on responsible adoption of AI and accelerating our Group artificial intelligence (AI) strategy to achieve impactful outcomes more quickly and incrementally. Based on the Strategy we aim to become an AI-enabled institution where all benefit from AI tooling, and where we are building AI-powered capabilities to maximize impact for our clients, employees and shareholders in a responsible and sustainable way. Building on more than a decade of experience in this field, our AI strategy is focused on adoption, agent-readiness, data strategy, partnerships and value delivery, all underpinned by safe, responsible, governed and trusted AI.

In 2025, we appointed a Chief Artificial Intelligence Officer to support UBS's transformation into a fully AI-enabled institution. We are moving quickly to adopt AI throughout the Group to reshape business capabilities and enhance employee productivity, and we are constantly developing and exploring new opportunities, from more traditional automation to agentic AI solutions. As part of our approach to deliver the AI strategy, there are nine transformational AI initiatives (so-called "Big Rocks"). These initiatives are designed to have a broad impact across the Group, applying AI to transform key business and control processes. In practice, they focus on enhancing client and employee outcomes, modernizing operational processes, strengthening risk management and control capabilities, improving software engineering and data foundations, and supporting the responsible adoption of AI at scale. For example, we are implementing the next generation of software development, which includes the widespread

deployment of AI-augmented development tools, enabling engineers to deliver faster, more innovative and scalable solutions and streamlining UBS's software development cycle.

Currently, we already have over 560 live use cases of AI across the bank with around 900 use cases in development and are progressing on our 9 large-scale, end-to-end transformational initiatives. Our strategic approach to applying AI at scale to support Financial Advisors through our Smart Technologies and Advanced Analytics Team (STAAT) was recognized by the Financial Times at the Professional Wealth Management Wealth Tech Awards, which named UBS as the Best Wealth Management Firm for Use of AI in the US and by the Celent Model Wealth Manager Award for Data, Analytics and AI. Nearly 90% of FA teams use the platform, powering millions of AI-driven client interactions.

As with any large-scale transformation, this requires securely integrating AI capabilities into existing systems and control environments, developing specialist skills, managing third-party dependencies, addressing evolving cybersecurity and information security risks and ensuring that outcomes are transversal and reusable across the firm.

In 2026, we also established a hub-and-spoke model, where the hub is the central AI office—staffed by technology, operations and risk teams—that sets standards for responsible AI use and brokers external partnerships. The spokes are within the bank's separate business divisions and group functions, each with a dedicated AI head who oversees delivery and drives experimentation. Binding it all together are more than 1,000 AI champions and a network of regional "AI factories" that pressure-test ideas and keep ambition tethered to the bank's risk appetite. That architecture helps managers decide what to scale—and spot when a solution built for one purpose can also serve another. For more details, our Group Chief Operating Officer Bea Martin recently outlined our approach in the Economist in an article titled "How UBS is rewiring for the age of AI"¹.

To support our employees, we are upskilling and reskilling employees in AI and adjacent technologies through specific AI training to ensure innovation is embedded responsibly and effectively. For example, over 1,000 senior Leaders have completed the AI Senior Leadership Journeys with Oxford.

We also continue to invest in partnerships with leading academic institutions worldwide and other key players, to develop ideas, drive positive outcomes across the Group and foster pioneering AI research. For example, we have announced the launch of the Oxford-UBS Centre for Applied Artificial Intelligence, focusing on three key research areas (society, economy and the future of AI). In addition, we also continue to expand our partnership ecosystem through collaborations with hyperscalers, frontier AI and FinTechs. This enables us to explore pioneering applications and practical solutions that can be implemented at scale across the Group.

AI Framework and Policy

These efforts are underpinned by an AI framework and policy that have been designed to safeguard the interests of clients, employees and stakeholders. Our Group AI policy sets out our governance standards and requirements for the development, adoption and operation of AI based on principles of respecting human autonomy, preventing harm, and ensuring fairness and transparency.

As of 1 April, we've updated our Group AI Policy to reflect the rapid evolution of artificial intelligence. With the update we are moving from risk managing AI at the individual AI Use Case level to a System-based approach. This means we'll oversee entire AI Systems, not just single projects. It's a more effective way to manage risk, reduce duplication and help teams bring AI into their work more quickly and consistently. All AI must comply with UBS policies on data protection, data ethics, cybersecurity, third-party risk, data management and model governance. And the responsible use of AI remains everyone's responsibility.

We hope this additional input is helpful as you develop the final report ahead of twenty-first meeting of the Group of Twenty (G20) in the United States of America this fall.

¹ Economist - [How UBS is rewiring for the age of AI | Enterprise AI in action](#)

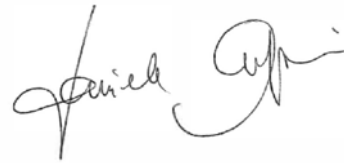
Should you have any questions or like to discuss further our approach to implementing AI at UBS, please do not hesitate to contact Oliver Wunsch or Daniele Magazzeni. We would also be available to participate in outreach activities.

Yours sincerely

UBS AG



Oliver Wunsch
Head Governmental & Regulatory Affairs



Daniele Magazzeni
Chief Artificial Intelligence Officer