

28 July 2026

BSA RECOMMENDATIONS FOR THE PUBLIC CONSULTATION PAPER ON THE AI GOVERNANCE BILL

On behalf of the Business Software Alliance (**BSA**),¹ we appreciate the opportunity to provide our recommendations on Malaysia's Public Consultation Paper on the AI Governance Bill (**Consultation Paper**) to the National AI Office (**NAIO**). BSA is the global trade association of the enterprise software industry. Our member companies are leaders in artificial intelligence, cybersecurity, cloud computing, and other cutting-edge technologies.

Close consultation between policy makers and the affected industry and other stakeholders is critical for effective policy making. BSA offers our extensive global experience in technology policy to serve as a resource, and we hope that our recommendations below will be helpful to the Government of Malaysia as you develop an AI Governance Bill that promotes the responsible and trustworthy development, deployment, and use of AI in Malaysia.

We commend NAIO for conducting this public consultation at the pre-drafting stage. Our recommendations below are directed at refining the principles- and risk-based approach in the Consultation Paper for carrying through to the legislative text and subsidiary instruments that will follow.

BSA Resources

We provide BSA resources relating to AI policy and governance that you and your team may find helpful as you consider the principles underpinning the AI Governance Bill:

- [**BSA Policy Solutions for Building Responsible AI:**](#)² A comprehensive set of recommendations for policymakers worldwide to address the most important AI policy issues and advance the adoption of responsible AI across the economy
- [**BSA's Framework to Build Trust in AI:**](#)³ BSA's bias risk management framework for AI
- [**Crosswalk Between BSA Framework to Build Trust in AI and NIST AI Risk Management Framework:**](#)⁴ An analysis to illustrate the alignment between the two frameworks

¹ BSA's members include: Adobe, Alteryx, Amadeus, Amazon Web Services, Asana, Atlassian, Autodesk, Avalara, Bentley Systems, Box, Cisco, Cloudflare, Cohere, Cohesity, Dassault Systemes, Databricks, Datadog, Docusign, Dropbox, Elastic, EY, Graphisoft, HubSpot, IBM, Kyndryl, MathWorks, Microsoft, Notion, Okta, OpenAI, Oracle, PagerDuty, Palo Alto Networks, PTC, Rubrik, Salesforce, SAP, ServiceNow, Shopify Inc., Siemens Industry Software Inc., TrendAI, TriNet, Veeam, Workday, Zendesk, and Zoom Communications Inc.

² BSA Policy Solutions for Building Responsible AI, at <https://ai.bsa.org/bsa-policy-solutions-for-building-responsible-ai/>

³ BSA's Framework to Build Trust in AI, at <https://www.bsa.org/reports/confronting-bias-bsas-framework-to-build-trust-in-ai>

⁴ Crosswalk Between BSA Framework to Build Trust in AI and NIST AI Risk Management Framework, at <https://www.bsa.org/policy-filings/us-crosswalk-between-bsa-framework-to-build-trust-in-ai-and-nist-ai-risk-management-framework>

Risk-Based Approach

We support a risk-based approach to assigning responsibilities to different entities in the AI development and deployment value Chain. The Consultation Paper recognises that regulatory obligations should be proportionate to the nature, context, and level of risk posed by an AI system, and that lower-risk AI applications should not be unnecessarily burdened. The risks of AI are inherently use-case specific. Any regulatory obligations should focus on specific applications of the technology that pose higher risks to the public and must be flexible enough to account for the unique considerations implicated by specific use cases.

The scope of any regulatory obligations should be a function of the degree of risk and the potential scope and severity of harm. Many AI systems pose extremely low, or even no, risk to individuals or society, while creating significant benefits like helping to organise digital files, auto-populate common forms for later human review, or improve a company's ability to forecast supply chain issues.

Focus on high-risk use cases rather than classes of AI systems

The Consultation Paper proposes that the Central AI Authority may issue requirements “By Class of AI Systems,” classifying AI by functional capability and technical characteristics, with the Developer as the primary duty holder. We are concerned that this approach, classifying AI systems by capability and technical characteristics, rather than focusing on the actual uses of AI systems, will impose unnecessary obligations on the developers and deployers of AI systems used for low-risk purposes.

Risk should be assessed by whether a particular AI system is used in a way that presents a material risk of harm to users, not based on the type of AI system or the sector in which it is deployed. The context of use and the potential for harm are generally known only at the deployment stage: a general-purpose capability may serve use cases ranging from no risk to high risk, and class-based technical controls risk capturing large numbers of low-risk applications. For the same reason, blanket obligations should not be imposed on general-purpose AI models without reference to their use cases.

Recommendation: Focus governance requirements on specified high-risk use cases, defined by consequential outcomes. For example, AI systems may be high-risk if they are used to make consequential decisions that determine an individual's eligibility for, and result in the provision or denial of, employment, credit, housing, education, healthcare, or insurance. Requirements issued “By Class of AI Systems” should be reserved for narrow circumstances and should not become a vehicle for regulating general-purpose technologies irrespective of the risk of the use case.

Narrow the harm categories within the risk framework

The Consultation Paper bases the risk framework on four categories of harm: death, bodily injury, unlawful deprivation of fundamental liberty, and contravention of any written law. We are concerned this approach does not ensure that obligations focus on truly high-risk uses of AI. For example, if contravention of “any” written law is treated as high risk, it may treat violations of laws that prevent bodily harm the same as violations of laws that regulate how government bodies purchase office supplies. We strongly recommend removing this category, to ensure the consultation takes a truly risk-based approach to AI.

Recommendation: Remove or substantially narrow the category of “contravention of any written law,” and define harm by reference to consequential effects that are legally or similarly significant for individuals. Define the high-risk category with an exhaustive, closed list of clearly articulated use cases, limited in scope, and subject to periodic review through a transparent and consultative process.

Define the prohibited tier by reference to specific practices, not intent

The Consultation Paper proposes that Tier 1 (unacceptable risk) would capture any AI system developed or deployed “with an intent to cause harm.” The “intent to cause harm” standard is challenging for a variety of reasons and using this standard to potentially prohibit entire AI systems may diverge significantly from the risk-based approach described above. We further posit that an action with the “intent to cause harm” is already within the scope of existing law (e.g., under the Penal Code) without being constrained by the tool used to carry out the harm. Rather than creating a prohibited category, we encourage the adoption of a robust, use-case oriented risk-based approach to regulation that would mitigate harms to consumers without necessarily prohibiting entire classes of AI systems.

Recommendation: Delete the unacceptable risk category and adopt a use-case oriented risk-based approach to regulating AI systems. If the unacceptable risk category is retained, this risk tier should be defined by an exhaustive, closed list of specific prohibited practices, objectively defined by the nature of the use and its effect on individuals. The Bill should also clarify that the prohibition attaches to the party engaging in the prohibited practice, and that a Developer or provider of a general-purpose AI system is not subject to sanctions due to a third party’s misuse of that system. Similar to the list of high-risk practices, the list of prohibited practices should be limited in scope and subject to periodic review through a transparent and consultative process.

Set Definitions in Line with International Understanding

Given that AI systems are developed and deployed in an international context, regulations and standards that apply to AI should operate across different jurisdictions to facilitate and promote further adoption and use of AI technologies. Definitions pertaining to AI should be aligned across jurisdictions to ensure that all stakeholders have a common understanding of AI.

Definition of AI

The Consultation Paper describes AI as “a functional capability or a set of methods that enables a system to perform functions that resemble human cognitive functions.” To support global adoption of AI, regulations and standards should be internationally aligned. BSA recommends that the AI Governance Bill adopts the OECD’s updated definition of AI which defines AI as “a machine-based system that for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments. Different AI systems vary in their levels of autonomy and adaptiveness after deployment.” This definition has been referenced by regulators worldwide, including the European Union (EU).

Recommendation: Adopt the OECD definition of AI in the AI Governance Bill. Because AI regulations should focus on high-risk uses of AI, rather than AI broadly, we also strongly recommend identifying a specific set of high-risk uses of AI that are the focus of any regulation.

Roles and Responsibilities of AI Actors

We welcome the Consultation Paper's recognition that parties should have varying degrees of accountability based on their degree of control over the AI system, and the Consultation Paper's recognition of the distinction between Developers and Deployers. AI legislation should reflect the different roles and responsibilities of different organisations in the AI supply chain. The OECD's Recommendation states that effective AI policies must necessarily account for "stakeholders according to their role and the context" in which AI is being deployed. The AI supply chain often includes several different actors: developers who build original models, integrators who incorporate general purpose AI into software applications, and deployers who put systems into use in the real world.

As currently proposed, however, the definition of "Developer" is very broad as it includes the party that trains the model, the party that adapts it for a specific use case, the party that integrates it into a wider system, and the party that later modifies it after deployment. Grouping these actors into a single category does not reflect their materially different levels of control over, and information about, an AI system. The Consultation Paper also does not address the allocation of responsibility where a Deployer takes an AI system that was not designed for high-risk use and implements it in a way that creates a high-risk use case.

Recommendation: Consistent with a focus on high-risk uses of AI, define developers of high-risk AI and deployers of high-risk AI specifically:

- **Developers of high-risk AI systems** are entities that: (1) design an AI system specifically intended to be used as a high-risk system; (2) substantially modify high-risk AI systems, or (3) substantially modify non-high risk AI systems so that they become high risk AI systems.
- **Deployers of high-risk AI systems** are entities that use a high-risk AI system.

Recommendation: Clarify that where a Deployer implements an AI system that was not designed for high-risk use in a way that creates a high-risk use case, the Deployer — rather than the Developer — is responsible for the risk management obligations associated with that high-risk use.

AI Incident Reporting

BSA cautions against introducing broad, standalone AI incident reporting requirements. Incident reporting can be an important way in which users and regulators are made aware of circumstances that may put consumers at risk and warrant mitigating steps to be taken. However, AI-specific incident reporting requirements risk duplicating existing reporting mechanisms, such as personal data breaches, cybersecurity incidents, system outage reports, and other such reporting obligations already in place. Rather than introducing another layer of reporting, BSA encourages ensuring current and existing incident reporting mechanisms are updated, if needed, to reflect the increasing use of AI. As such, we welcome the Consultation Paper's stated intention to leverage existing sectoral reporting mechanisms.

The reportable incidents as currently described — failures, weaknesses, misuse, unexpected effects, and near misses — are very broad and appear to apply across all AI systems regardless of risk tier. Further, the Consultation Paper does not define thresholds for reporting or notification timelines. Applying mandatory reporting obligations to low-risk use cases would be inconsistent with a risk-based approach and would be highly resource-intensive for both reporting organisations and the Central AI Authority without commensurate benefit to the safe and responsible use of AI.

BSA members already operate under incident and breach reporting obligations in Malaysia, including under the Personal Data Protection Act, the Cyber Security Act, and sector-specific regimes such as those administered by Bank Negara Malaysia. An AI reporting channel that is not coordinated with these existing systems risks requiring multiple reports on the same underlying event to different authorities under different criteria and timelines.

Recommendation: Limit mandatory incident reporting obligations to high-risk use cases and place the reporting obligation on the appropriate party. Typically, the deployer is best placed to identify whether harm has occurred. The developer should support the deployer with relevant information rather than be subject to reporting requirements for the same incident. This ensures that reporting obligations are proportionate to the likelihood and severity of potential harm, consistent with the risk-based approach of the Bill and allows both organisations and the Central AI Authority to focus resources on the incidents most relevant to regulatory oversight and policy development.

Recommendation: Incident reporting requirements should be based on clearly defined and objectively measurable criteria, including specific classifications of reportable incidents tied to the defined categories of harm, materiality thresholds, and notification timelines that reflect internationally adopted best practices. These principles should be established in the Bill or its subsidiary regulations to avoid open-ended compliance expectations and excessive reporting of low-value events. Notification timelines should start when the organisation confirms that a reportable incident has occurred, rather than from initial detection or suspicion, and should allow a reasonable period for reporting that does not divert resources from, or otherwise jeopardize, containment and remediation efforts.

Recommendation: Include an explicit non-duplication provision so that where an incident is already reportable under another law — such as the Personal Data Protection Act, the Cyber Security Act, or a sectoral regime — a single report satisfies the overlapping obligations, with information shared between the relevant authorities as needed.

Transparency, Explainability, and Protection of Confidential Information

We support Principle 2 of the proposed AI Governance Principles, which calls for transparency and explainability proportionate to risk and impact, and Principle 3's emphasis on clear accountability and effective redress. The implementation of these principles must be workable in practice. Explainability obligations should account for the capabilities of the technology, the need to protect businesses' sensitive proprietary information, and relevance to the audience. Fulfilling such obligations should be accomplished by explaining how the AI system operates to produce outcomes, rather than granular disclosures of training data or explanations of algorithms. In addition, redress mechanisms should be operationally feasible and narrowly tailored to circumstances that warrant individual review.

Similar considerations apply to the Central AI Authority's proposed investigation and enforcement powers, which include the power to compel the production of records. Compelled disclosure of source code, model weights, or other confidential technical information would create significant risks for the protection of trade secrets, intellectual property, and cybersecurity.

Recommendation: Ensure that instruments implementing Principles 2 and 3, and the exercise of investigation and enforcement powers, do not require the routine disclosure of source code, proprietary algorithms, model weights, or other confidential technical information. Any transparency or documentation requirements should be proportionate to risk, and information obtained through investigative powers should be subject to robust statutory confidentiality safeguards.

Consultation on Legislative Text, Subsidiary Instruments, and Transition

We commend the NAIO for holding a public consultation at the pre-drafting stage. As the Consultation Paper reserves the operational requirements of the principles for subsequent regulation, codes, standards, and guidelines, much of the framework's practical impact on organisations will be determined by instruments that have yet to be published. Meaningful consultation on the actual legislative text, and on the subsidiary instruments that will follow, is essential to regulatory certainty and ensures that requirements are practical to implement.

Recommendation: Conduct further public consultation on the draft legislative text of the AI Governance Bill before it is tabled, and enshrine stakeholder consultation and proportionality as statutory principles, including a requirement to consult affected stakeholders before issuing or revising binding subsidiary instruments.

Recommendation: Phase in compliance obligations with a transition period of at least 12 to 18 months following the publication of the relevant implementing instruments and provide grandfathering arrangements for AI systems and contractual arrangements structured in compliance with laws in force at the time of the Bill's enactment.

Conclusion

BSA thanks the NAIO for considering our recommendations on the Public Consultation Paper. We stand ready to support the Malaysian government and wish to act as a resource on international developments and best practices for AI governance policy. **Additionally, we would like to continue the discussion with the NAIO.** Please do not hesitate to contact me at waisanw@bsa.org or +65 9729 1253 to make arrangements. Thank you once more for your time and consideration.

Yours sincerely,

Wong Wai San

Director, Policy – APAC