

ANNEXURE A

Annexure A: Response to the Consultation Paper's Guiding Questions

Submission of Edwin Lee & Partners, with a contribution from Global Law Office, Beijing
Public Consultation on the Proposed AI Governance Bill

Edwin Lee & Partners, Kuala Lumpur

with a contribution from Global Law Office, Beijing, People's Republic of China (Annexure C)

Dated 31 July 2026 | Filed at policy@ai.gov.my per the Consultation Paper

[FEEDBACK] AI Governance Bill - Public Consultation - EDWIN LEE & PARTNERS - 31 JULY 2026

1. Scope of the Bill

Question	Our position	Developed at
1. Alignment of the five definitions with our understanding	Yes, in substance. Two adjustments: the closing limb of the AI System definition reads more widely than the Paper's own AI/ordinary-software distinction intends; and the AI Lifecycle would be far more useful if mapped against duties Malaysian organisations already owe at each stage. Developer and Deployer track controller and processor closely but not perfectly, and the mapping should be express.	¶23 · R3, R6
2. A common definition across the public and private sectors	Yes, unequivocally. The question is more consequential than it appears: section 3(1) of Act 709 excludes the Federal and State Governments from the PDPA, so for public-sector AI there is no data-protection floor beneath the Bill and the Bill must supply the substantive standard itself. A citizen refused a benefit and a customer refused a loan face the same category of harm and should be described by the same vocabulary.	¶6–7, 20–21 · R5
3. Exemptions by sector, domain or use-case	We do not recommend further categorical exemptions; risk tiering is the better instrument, and exemptions granted by sector will not track where harm actually occurs. Two clarifications to the exemptions already proposed: personal use should attach to the individual's use only: never the provider, never a public authority delivering services; and national security should attach to the <i>use</i> rather than the system, so that dual-use systems are regulated in respect of their non-security use.	¶21–22 · R4, R5
4. Additional considerations or exemptions	Three, all structural. The extraterritorial scope needs an enforcement anchor, and its information demands need a conflicts mechanism: a foreign Developer may be restrained by its home law from producing the documentation the Authority demands, and the framework should route such conflicts through regulatory cooperation. The Bill should state whether it binds the Federal and State Governments; Act 854 s. 2 shows how, and the Paper's Enablement function already assumes public-sector application. And the Bill should not be drafted on the assumption that a lawful basis exists for training AI on personal data, because under Act 709 that question is unresolved on three independent grounds.	¶8–10, 20–21, 24 · Annexure C · R5, R7, R10
5. Additional definitions	Eight: <i>AI harm</i> ; <i>AI incident</i> ; the high-risk threshold; <i>automated decision-making and profiling</i> ; <i>personal data</i> (by reference to Act 709 s. 4); <i>synthetic content</i> ; <i>substantial modification</i> ; and <i>general-purpose AI</i> . Each is a term on which an operative duty already silently depends, and for most a Malaysian instrument supplies the content. Annexure C addresses the general-purpose AI point with worked configurations.	¶23 · R8
6. Standards to guide the definitions	Internationally, the OECD definition or ISO/IEC 22989 for AI, and the Paper's own ISO/IEC 5338:2023 reference for the lifecycle. Domestically, and we would press this equally, Act 709 s. 4 for personal data and the ADMP Guideline for automated decision-making and profiling. Interoperability with Malaysian law matters as much as interoperability with foreign law. At the international level the Council of Europe Framework Convention on AI (CETS No. 225) supplies a human-rights baseline against which Principle 1 can be tested, and ISO/IEC 42001 is the certifiable management-system standard most likely to serve as the practical multi-jurisdiction vehicle. Malaysia's founding membership of the World AI Cooperation Organization (16 July 2026), alongside its ASEAN commitments and its businesses' exposure to the EU regime, means the definitions must now work across frameworks anchored in different regulatory traditions, a further reason to anchor to neutral international formulations rather than any single jurisdiction's.	¶19A, 23, 31 · R6, R8, R12

2. The AI Governance Architecture

Question	Our position	Developed at
1. Our current practices in engaging regulators and managing multi-regulator compliance	As advisers, we act for organisations that engage the Commissioner, BNM, the Securities Commission, MCMC and NACSA in respect of the same technology deployments. The recurring difficulty we observe is not the number of regulators but that each applies its own assessment and reporting expectations to one system, which is the practical foundation of R1 and R16.	¶18 · R1, R16
2. Additional functions, powers or safeguards	Three. A defined internal review and an independent appeal route, absent from the Paper and unusual by Malaysian regulatory standards given the combination of rule-making, investigation and penalty in one body. A statutory coordination protocol with the Commissioner rather than administrative practice; the National AI Action Plan 2026-2030 (29 July 2026) already commits to formalised coordination mechanisms and protocols covering incident reporting and sandboxes (initiatives E11, E12), and the Bill is their natural home. An express detection-capability limb within the Enablement function, paired with a knowledge-based reporting trigger. And clarification of the institutional map following the 28 July 2026 launch of AI Malaysia Berhad as NAIIO's successor: which body is the Central AI Authority; vesting of coercive powers in a statutory authority rather than a company; and defined status under the Bill for assessments by the Malaysian AI Safety Institute.	¶18–19, 18A, 41 · R1, R2, R16
3. Anticipated implementation challenges	Duplicate assessment of one system; multiple incident clocks on one event; two legal characterisations of one organisation; and classification that is indeterminate where the Paper's three axes disagree. Each is avoidable at drafting stage and expensive to retrofit.	¶18, 36, 37–38, 40–41 · R1, R15, R16
4. Additional guidance, tools or implementation support	Four deliverables that would convert our recommendations into usable instruments: a joint NAIIO–Commissioner guidance note mapping Developer/Deployer to controller/processor; a DPIA-to-AI-assessment crosswalk so one document demonstrably serves both regimes; a single combined incident-notification form accepted by the Authority, the Commissioner and NACSA; and a simplified assessment template for SMEs, which the Paper's own proportionality commitment (p. 3) implies. Two further asks: recognition of certification against an established AI management standard as evidence of due regard, so multinational organisations are not certified repeatedly for the same thing; and guidance on the agentic deployment pattern, for which regional material already exists.	¶31, 38, 41–42 · R12, R16
5. Powers and functions of Sectoral Leads	These should be specified in the Bill rather than left wholly to the instrument of delegation: issuing sectoral instruments; receiving and triaging incident reports for the sector; conducting or accepting risk assessments; co-supervising sandbox cohorts; and a stated primacy rule where the sectoral and central positions differ. The two clearest candidates are the Commissioner, for AI processing personal data, and MCMC, for AI-generated content harms.	¶16–18 · R1, R19
6. Local and international benchmarks	Korea's AI Basic Act (Act No. 20676, in force 22 January 2026) is the closest architectural comparator: a framework statute delegating detail to subordinate instruments, classifying its elevated category by sector and use rather than intent, applying across public and private sectors, and requiring foreign operators to designate a domestic representative; it independently reached four of our recommended positions. Vietnam's Law on AI No. 134/2025 is ASEAN's first binding standalone AI statute. Singapore is the closest comparator to Malaysia's <i>present</i> position: AI governed through data protection law plus sectoral regulation and central assurance infrastructure. The Council of Europe Framework Convention (CETS No. 225) is the human-rights baseline. Malaysia's own founding membership of the World AI Cooperation Organization (16 July 2026) gives the	¶18A, 19A, 40

Question	Our position	Developed at
	Authority's international cooperation function concrete treaty-level content; see ¶19A. Domestically, Act 854 already runs a central-authority-plus-sector-lead architecture in a cognate field and shares part of the Bill's vocabulary; and the National AI Action Plan 2026-2030 supplies the domestic policy anchor, committing in terms to central oversight with sector empowerment, a principles-based framework aligned with international standards, and formal cross-regulator coordination (E11, E12).	

3. AI Governance Principles

Question	Our position	Developed at
1. How the five principles are reflected in our practices	We answer as an organisation that itself deploys AI. <i>Human agency (P1)</i> : no AI-assisted output reaches a client without review by a qualified lawyer, and our client-facing compliance tools are deliberately deterministic so that their entire logic is human-written and reviewable. <i>Transparency (P2)</i> : proportionate to a professional-services context; our deliverables are issued under an identified lawyer's responsibility rather than presented as machine output. <i>Accountability (P3)</i> : every deliverable is attributable to the reviewing lawyer, never to the tool; our professional obligations do not permit responsibility to be displaced onto a system, which is Principle 3's own rule operating through professional duty. <i>Safety and security (P4) and data governance (P5)</i> : client information entering any AI system is governed by our confidentiality and PDPA obligations, which operate as the binding framework the Bill's principles describe. Our practice illustrates our broader submission: for organisations already under professional and data-protection duties, the principles largely formalise existing obligations; the value the Bill adds is making them uniform.	¶2, 23
2. The framework currently guiding us	For our own use and for the organisations we advise, the operative framework is the PDPA and the Commissioner's 2024–2025 instruments, supplemented for AI specifically by the National Guidelines on AI Governance and Ethics (2024), which are voluntary, and by professional duties of confidentiality and competence. The absence of any binding AI-specific standard is precisely why this Bill matters.	¶3
3. Challenges in operationalising the principles	Principally that "due regard" is unfalsifiable without a floor: an organisation can honestly believe it had due regard and have nothing to show for it. Second, that Principles 1–3 already have operative Malaysian content in the ADMP Guideline which the Bill has not yet claimed. Third, that the phased voluntary-then-codify approach creates a transition problem: organisations investing early against voluntary instruments need assurance the investment will count when those instruments harden. Fourth, and most substantially: Principle 5's "properly sourced" and "provenance" language presupposes a lawful basis for training AI on personal data which is unresolved under Act 709 at three separate levels , whether the Act applies to open-web collection at all, what counts as effective anonymisation, and, if it applies to identifiable data, which of the narrow section 6 bases is available.	¶8–10, 26–31 · R9, R10, R12
4. Other principles to consider	One: fairness and non-discrimination. The gap matters more in Malaysia than elsewhere, because discriminatory private-sector algorithmic outcomes will often contravene no written law and so escape the risk framework as well; and the National Guidelines on AI Governance and Ethics (2024) already name fairness, so adopting it restores continuity rather than importing anything. Two further candidates were considered and are not recommended: a standalone environmental principle (a matter for existing energy regulation) and a standalone innovation principle (already carried by the Bill's proportionality structure and the sandbox).	¶29 · R11

Question	Our position	Developed at
5. Guidance or governance tools to assist implementation	A documentation template for demonstrable due regard at Tier 2; statutory adoption of the ADMP Guideline's no-sole-factor rule and AI-use notice rather than fresh formulations of the same ideas; a capability standard for human oversight (understanding of what the system optimises for, access to evaluative information, authority to depart from the output) expressed as a gradation of involvement rather than as presence or absence; recognition of ISO/IEC 42001 certification as evidence of due regard; and sectoral codes issued through Sectoral Leads.	¶27, 30–31 · R9, R12

4. AI Risk Framework

Question	Our position	Developed at
1. How we identify, assess, classify and manage AI risk	Our most consequential AI risk decision was architectural. When we built compliance tools for clients' direct use, we deliberately made them deterministic (fixed logic and scoring rules written by our lawyers, calling no model at runtime) because for an assessment delivered under a law firm's name, output unpredictability is itself the principal risk, and we chose to design it out rather than manage it. For the generative systems we do deploy, the control is a human checkpoint: AI-assisted output is reviewed by a qualified lawyer before it reaches a client. We offer this as a small illustration of the Bill's own logic working in practice: risk classification driving proportionate controls, with the strictest control (architectural elimination) applied where the output most affects third parties.	¶2, 23
2. Risk management frameworks currently used	Across the organisations we advise, the DPIA under the Commissioner's guideline is the assessment instrument actually in operation today, which is the practical basis of the mutual-recognition recommendation at R16. AI-specific frameworks, where present, are voluntary adoptions of the National Guidelines on AI Governance and Ethics (2024) or internal AI-use policies.	¶37 · R16
3. Risks and unintended consequences of the proposed framework	Five, none requiring anyone to behave badly. An unqualified "contravention of any written law" category that no Developer or Deployer can triage against. An intent gate at Tier 1 that excludes catastrophic indifference while duplicating criminal law for what it does catch. Duplicate assessment cost, falling hardest on the SMEs the Paper undertakes to protect. Indeterminate classification where the three implementation axes disagree about the same system. And, less obviously: because the fourth harm category is contravention of any written law, the unresolved lawful-basis question for AI training means a large class of <i>development</i> activity could register as harm-relevant from day one, not because a system is dangerous but because an input question is unsettled. That makes the gravity threshold at R13 more important, not less.	¶9, 33, 35–38 · R10, R13, R14, R15, R16
4. Additional risk categories, safeguards or governance measures	Three harm categories, each already recognised in the Commissioner's own materials rather than imported: significant financial harm; unjustified discriminatory treatment; and significant reputational harm or denial of access to essential services, employment or education. Plus a gravity threshold on the fourth existing category, so it keeps its valuable integrating function without collapsing the tiering. On safeguards: the role mapping should accommodate parties added to the processing chain at run time, which is how agentic deployments assemble, and oversight should be graduated rather than binary.	¶31–34 · R3, R12, R13
5. Implementation challenges in	Duplicate assessment; the absence of a named internal owner, where the mandatorily appointed Data Protection Officer is the natural candidate; the absence of a substantial-modification threshold, so that a classification made once is treated as permanent and the	¶23, 36, 37–38 · R8, R15, R16

Question	Our position	Developed at
applying the framework	obligation quietly expires on the first model update; and tiering that is functionally binary, which argues for graduated obligations within Tier 2.	

5. AI Incident Reporting

Question	Our position	Developed at
1. How AI incidents are currently identified, classified and managed in our organisation	In our own operations and across the organisations we advise, AI-related incidents surface today as personal data breaches and are managed under the Circular No. 1/2025 process; no AI-specific incident taxonomy exists, and events with no personal-data dimension (a model behaving harmfully without a breach) currently have no reporting home at all. That gap is precisely what the Bill's incident framework fills, and why the single-window design at R16 matters: for most organisations, the AI incident and the data breach will be the same event.	¶39–42 · R16
2. Coordination between the Central AI Authority and Sectoral Leads	Designate the Commissioner's breach-notification regime and the Act 854 notification regime as recognised existing mechanisms; provide for single-window filing shared between the Authority, the Commissioner and NACSA; state a primacy rule; designate the Deployer as primary reporting party, with responsibility allocation between Developer and Deployer resolved after filing; and provide for the case where a criminal investigation under the Cybercrimes Bill 2026 runs alongside the Authority's technical fact-finding, including evidence preservation.	¶40–41 · Annexure C · R16
3. Challenges in establishing cross-sector reporting	Three clocks on one event, which will degrade the incident response itself in the hours that matter. An occurrence-based trigger that would impose strict liability for incidents an organisation never saw, punishing incapability rather than concealment. And differing definitions of the reportable event across the engaged regimes.	¶40–41 · R16
4. Information to be collected	The particulars already prescribed by Circular No. 1/2025 should form the base (time of knowledge, nature and type, detection method and believed cause, numbers affected, systems implicated, likely consequences, chronology, remedial and mitigating steps, steps for affected persons, contact point) with four AI-specific additions: the system and version that acted; the risk classification and whether an assessment existed; whether human oversight was designed in and exercised; and whether an automated decision about a person, or synthetic content, was involved.	¶42 · R16
5. Additional reporting mechanisms or safeguards	A three-layer structure for the duty itself: incidents reportable mandatorily on a knowledge trigger; near misses reportable voluntarily with safe-harbour protection, without which that limb will not be used; and mere weaknesses, no harm occurred or reasonably foreseeable, attracting a mitigation duty rather than a filing. And routing of the public-complaint channel so one complaint by an affected data subject can engage both the AI incident process and the individual's PDPA rights, noting a third complaint channel already exists under the Online Safety Act 2025.	¶41–42 · Annexure B and C · R16, R19

6. AI Sandbox

Question	Our position	Developed at
1. Whether we participate in a	No: neither we nor most organisations we advise currently participate in a sandbox; BNM's Financial Technology Regulatory Sandbox is the familiar domestic reference	¶43–45 · R17

Question	Our position	Developed at
sandbox, and what environments would benefit us	point. The environments most useful to the organisations we advise would be cohorts permitting supervised testing on anonymised or synthetic personal data with the Commissioner's involvement, and exit outcomes carrying defined regulatory effect; see Q2 to Q4 below.	
2. Operationalising the sandbox: conditions, support, incentives	The decisive incentive is legal effect on exit: a provisional risk classification or time-limited safe harbour for the system as tested. Conditions for personal-data cohorts should include a DPIA at entry, standing data-minimisation and purpose-limitation terms, expressly unaffected breach-notification duties, and anonymisation and synthetic-data standards set with the Commissioner's input. Framed this way the requirements are a benefit: the participant exits with a validated data-protection posture alongside its AI assessment. Beyond conditions, the framing itself is an incentive: designed as a regional entry pathway with ASEAN-recognised exit outcomes, the sandbox becomes an inbound-investment instrument, not only a domestic testbed.	¶43–45, 44A · Annexure C · R17, R24
3. Barriers to effective participation	Post-sandbox regulatory uncertainty, which the Questionnaire itself identifies and which is the barrier that matters most. Uncertainty about whether PDPA obligations are relaxed inside the sandbox: they are not, and the drafting should say so. Cost and duration for SMEs, which free or subsidised SME access addresses. And the risk that participation is itself read by regulators as an admission of risk, which only a defined safe harbour answers. Annexure C addresses barriers specific to foreign participants.	¶44–45 · Annexure C · R17
4. Additional safeguards, functions or operational considerations	Separate the technology-sandbox functions from the regulatory-sandbox function in the drafting, because only the latter needs statutory power to relax requirements, and no sandbox can relax obligations under other statutes unless the Bill so provides. Add Commissioner co-supervision for personal-data cohorts; deletion of sandbox personal data at exit; exit reports feeding the central risk and incident repository; and express preservation of third-party civil liability, so relief from administrative penalties is not mistaken for relief from claims by those harmed.	¶43–45 · R17
5. Sandbox models Malaysia should consider	Domestically, BNM's Financial Technology Regulatory Sandbox as the established template, and the Securities Commission's comparable arrangements. For the data-protection dimension, the UK Information Commissioner's regulatory sandbox, the leading example operated by a data-protection regulator. For AI-specific testing, Singapore's IMDA assurance programmes, instructive because they are technology rather than regulatory sandboxes. And for statutory design, EU AI Act Articles 57(10), 57(12) and 59, which resolve in terms the three questions the Paper leaves open: the data protection authority must be associated with the sandbox's operation (Art 57(10)); the conditions for personal data in the sandbox are exhaustively specified, including deletion at exit (Art 59); and good-faith participants are protected from administrative fines, extending to other law where that law's regulator supervised (Art 57(12)). Two further EU provisions complete the design: exit reports must be taken positively into account in subsequent conformity assessment (Art 57(7)), and sandbox participation in one Member State is mutually recognised with the same legal effects across the Union (Art 58(2)(g)), the internal precedent for the ASEAN mutual-recognition initiative we propose at R24.	¶43–45, 44A · Annexure C · R17, R24

This document is prepared by Edwin Lee & Partners. It is submitted in response to a public consultation and does not constitute legal advice to any person. Specific advice should be sought for particular circumstances.