



CXO INTELLIGENCE SERIES ENTERPRISE AIGOVERNANCE

DOI: 10.6084/m9.figshare.33137627

Strengthening the World's First Agentic AI Framework

Nine Recommendations for Singapore

*A constructive submission to IMDA on the Model AI
Governance Framework for Agentic AI (v1.5)*

SUBMISSION ON	Model AI Governance Framework for Agentic AI (v1.5)
PREPARED FOR	Infocomm Media Development Authority (IMDA), Singapore
SCOPE	9 recommendations · gaps, impact and proposed changes
AUTHOR	Prashant Akhawat akhawat.com akhawat@isecol.com
DATE	19 th July 2026
DOI	10.6084/m9.figshare.33137627

Contents

Executive Summary 3

Why This Submission, and Why Now 4

What Version 1.5 Already Gets Right..... 5

Nine Recommendations to Strengthen the Framework 6

 1. Resolve liability apportionment, do not just assign accountability 6

 2. Make Agent Identity Cards runtime-verifiable, not just a disclosure format..... 6

 3. Add an authority-creep trigger to the graduated autonomy taxonomy 7

 4. Publish quantitative acceptance thresholds and finalize agentic testing standards 7

 5. Introduce risk tiering with proportionate expectations for high-stakes agents 8

 6. Build an assurance and attestation pathway for agents 9

 7. Lead on cross-border interoperability, not just national guidance 9

 8. Create a structured channel for agentic incident reporting and shared learning.....10

 9. Give affected end users a defined path to contest an agent's action11

The Principle That Ties Them Together 12

Summary of Recommendations 13

Frequently Asked Questions..... 14

References and Further Reading..... 15

Executive Summary

Singapore has built the most advanced agentic-AI governance framework in the world. This is not a critique of it; it is a contribution to it: nine recommendations, offered in the same spirit of pragmatic leadership that produced the framework, for where the next version can go.

In January 2026, Singapore's Infocomm Media Development Authority (IMDA) published a governance framework built specifically for AI agents, systems that plan, reason, and act on their own. In May 2026 it released Version 1.5 with real-world case studies and feedback from more than 60 organizations, and it has explicitly invited feedback to refine the framework. This submission takes that invitation seriously.

The recommendations that follow are deliberately consistent with Singapore's philosophy: pragmatic, pro-innovation, and light-touch. None calls for heavy regulation. Each is designed to make an already-excellent framework more enforceable, more measurable, and more interoperable, while preserving the space for innovation that has made it a model for the world.

At a glance

- **The framework is the strongest of its kind.** This submission builds on its strengths rather than re-litigating them.
- **Nine issues remain genuinely open after v1.5.** From legal-liability apportionment, runtime-verifiable identity, and an authority-creep trigger to incident learning, end-user recourse, and cross-border interoperability.
- **Every recommendation is proportionate.** Each preserves the framework's voluntary, pro-innovation character.
- **One organizing principle unites them.** Accountability must scale with autonomy, an idea the framework already embraces in spirit.

Guiding principle: *accountability must scale with autonomy. The framework can now make that real.*

Why This Submission, and Why Now

The best time to strengthen a framework is while it is still being written, and IMDA has made clear that the MGF is a living document. Version 1.5 already absorbed feedback from dozens of organizations and closed several gaps that existed at launch, adding guidance on multi-agent systems, third-party agents, and automation bias. That responsiveness is exactly why targeted, specific input is worth offering now, before practice hardens, and while the testing guidelines and legal-responsibility work are still in train.

The wider ecosystem this submission builds on

This submission is written with Singapore's broader instrument stack in view, and is careful to connect to it rather than duplicate it: CSA's Securing Agentic AI addendum (announced October 2025, finalized 2026), which establishes security controls including immutable logging of prompts and tool calls, kill switches and fallbacks, and testing at autonomy levels; IMDA's Starter Kit for Testing LLM-Based Applications and the agentic testing guidelines being developed on top of it; the Global AI Assurance Pilot, which paired specialist testing firms with real deployments; the Model AI Governance Framework for Generative AI (2024), whose nine dimensions include incident reporting; the PDPC's proposed guidelines on personal data in generative AI (June 2026); and Singapore's AI Safety Institute together with its leadership of the ASEAN Working Group on AI Governance. Several of the recommendations below are, deliberately, proposals to extend and connect these existing instruments for agents.

What Version 1.5 Already Gets Right

Credible recommendations begin with an honest account of the framework's strengths, and they are considerable.

- **First and most complete.** The world's first governance framework built specifically for agentic AI, structured around four dimensions: assess and bound risks upfront, make humans meaningfully accountable, implement technical controls, and enable end-user responsibility.
- **A real risk model, not slogans.** Eight agent core components, the action-space versus autonomy axes, a four-level graduated autonomy taxonomy, and a named risk taxonomy give practitioners something concrete to work with.
- **Multi-agent risk taken seriously.** Version 1.5 explicitly addresses sequential, supervisor, and swarm patterns, and names agent sprawl, collaborative failure, conflicting objectives, collusion, and emergent behavior.
- **Agent Identity Cards.** A standardized disclosure format specifying an agent's capabilities, limitations, authorized action domains, and escalation protocols, a genuine contribution to the field.
- **Grounded in practice.** Updated with feedback from more than 60 organizations and over 10 real-world deployment case studies, plus a companion discussion paper on legal responsibility and testing guidelines in development.

The recommendations that follow treat these as the foundation to build on, not ground to revisit.

Nine Recommendations to Strengthen the Framework

Each recommendation names a specific open gap, ties it to the framework's own four dimensions, explains why the gap falls short in practice, and proposes a concrete, proportionate change.

01 Resolve liability apportionment, do not just assign accountability

FRAMEWORK DIMENSION

Dimension 2: Make humans meaningfully accountable

GAP THEME

Legal liability apportionment

PRIORITY

CRITICAL

THE GAP

The framework rightly insists humans remain accountable, but legal commentators note it leaves unresolved how liability is apportioned among the model provider, the agent developer, the deployer, and the end user, and the open question of whether an agent's autonomous action can legally bind its principal.

WHY IT FALLS SHORT

Without a default apportionment model, a single autonomous failure leaves provider, developer, deployer, and user pointing at one another, and courts, insurers, and boards with no shared starting point for who pays.

RECOMMENDATION

Extend the Legal Responsibility for AI Agents work into a default apportionment model: a rebuttable presumption that the deploying entity is liable, with defined conditions under which responsibility shifts to a provider or integrator, and clear treatment of when an agent's action binds its principal. Even as guidance, a reference allocation would give courts, insurers, and boards a shared starting point.

02 Make Agent Identity Cards runtime-verifiable, not just a disclosure format

FRAMEWORK DIMENSION

Dimension 1 & 3: Bound risks / Technical controls

GAP THEME

Runtime-verifiable identity

PRIORITY

HIGH

THE GAP

Agent Identity Cards are an excellent idea, but as a disclosure format they describe an agent rather than constrain it. Identity researchers note that current standards do not yet support the delegated authorization, transitive trust, and cross-domain federation that real agentic deployments require.

WHY IT FALLS SHORT

A card that only describes an agent cannot stop it: at the moment an agent acts across an organizational boundary, nothing checks that the action is within its authorized domain, so the disclosure and the behavior can silently diverge.

RECOMMENDATION

Evolve the Identity Card toward a cryptographically verifiable credential that a system can check at the moment of action, binding an agent's permitted action domain to enforcement, not just to documentation. Align it with emerging agent-identity and delegated-authorization standards so an agent's authority is machine-verifiable across organizational boundaries.

03

Add an authority-creep trigger to the graduated autonomy taxonomy

FRAMEWORK DIMENSION	GAP THEME	PRIORITY
Dimension 1: Assess & bound risks upfront	Authority-creep over time	CRITICAL

THE GAP

The four-level autonomy taxonomy sets governance requirements at deployment, but nothing in a voluntary framework arrests the slow, reasonable escalation of autonomy over time, the pattern in which approval steps are relaxed until a high-risk system is still run like the pilot it once was, and accountability quietly dissolves.

WHY IT FALLS SHORT

Because the taxonomy only sets requirements at deployment, autonomy can be relaxed step by step until a high-risk system is still governed like the low-risk pilot it began as, and accountability erodes without anyone deciding that it should.

RECOMMENDATION

Require, as a named practice, a re-assessment trigger tied to autonomy change: any increase in an agent's autonomy level, action space, or data access re-opens the upfront risk assessment and the accountability assignment. This is the governance complement to what CSA's Securing Agentic AI addendum already establishes at the security layer, where safety constraints must be verified before higher autonomy is enabled: the security tests gate the capability, while this trigger re-opens the risk assessment and, crucially, the question of who is answerable. Make raising autonomy a governed event, not a silent configuration change.

04

Publish quantitative acceptance thresholds and finalize agentic testing standards

FRAMEWORK DIMENSION

Dimension 3: Implement technical controls

GAP THEME

Measurable acceptance thresholds

PRIORITY

HIGH

THE GAP

The framework is strong on qualitative practice but offers few measurable acceptance criteria, and the agentic-AI testing guidelines are still in development. Without thresholds, meaningful oversight and end-user trust rest on judgment that varies by organization.

WHY IT FALLS SHORT

With no measurable acceptance criteria, 'safe enough' means something different at every organization, oversight rests on unaudited judgment, and results cannot be compared across deployers or trusted by end users.

RECOMMENDATION

Finalize the agentic testing guidelines being developed on top of the Starter Kit for Testing LLM-Based Applications, with concrete, tiered acceptance thresholds, for example maximum acceptable rates of erroneous or unauthorized actions by autonomy level, minimum logging-completeness and reconstructability standards, and monitoring-coverage expectations for agent-to-agent interactions. Pair the guidance with a reference test battery so results are comparable across deployers.

05

Introduce risk tiering with proportionate expectations for high-stakes agents

FRAMEWORK DIMENSION

Cross-cutting: Proportionality

GAP THEME

Risk tiering by stakes

PRIORITY

HIGH

THE GAP

The framework applies broadly regardless of stakes, so an agent handling routine queries and an agent moving money or affecting health receive the same framing. Peer regimes such as the EU AI Act tier obligations by risk; the MGF does not.

WHY IT FALLS SHORT

Applying the same expectations to a trivial agent and one that moves money or affects health either over-burdens low-stakes use or under-protects high-stakes use, and gives sector regulators no clear hook to act where consequences are greatest.

RECOMMENDATION

Introduce a light risk tier that scales expectations with consequence, sitting naturally on top of the existing autonomy and action-space axes. For a defined set of high-stakes domains, elevate a subset of recommendations from suggested to expected, giving sector regulators such as the Monetary Authority of Singapore and the Ministry of Health a clear hook without abandoning the framework's pro-innovation, voluntary character. Reinforce the tier through procurement: require MGF alignment for agentic systems supplied to the Singapore government, a proven soft-enforcement lever that creates real market incentive while keeping the framework itself voluntary.

06 Build an assurance and attestation pathway for agents

FRAMEWORK DIMENSION	GAP THEME	PRIORITY
Dimension 3: Implement technical controls	Assurance & attestation	MEDIUM

THE GAP

The framework tells organizations what good looks like but provides no mechanism to demonstrate, to a counterparty or a regulator, that an agent actually meets it. Singapore already has assurance tooling for traditional AI through AI Verify; agents have no equivalent.

WHY IT FALLS SHORT

Guidance on what ‘good’ looks like cannot be demonstrated to a counterparty or regulator without an attestation mechanism, so a well-governed agent and a poorly-governed one look identical from the outside and the market cannot price the difference.

RECOMMENDATION

Extend AI Verify and the Global AI Assurance Pilot, which has already paired specialist testing firms with real deployments, to test and attest agentic systems against the MGF, producing portable, evidence-backed assurance that a deployer can share with partners and supervisors. Independent attestation is what turns best-practice guidance into something the market can price and trust.

07 Lead on cross-border interoperability, not just national guidance

FRAMEWORK DIMENSION	GAP THEME	PRIORITY
Cross-cutting: Interoperability	Cross-border fragmentation	HIGH

THE GAP

An agent that satisfies Singapore's framework may still fail to meet the EU AI Act's obligations or China's labeling rules, and the jurisdictions differ at a structural level about what agent governance is even for. Fragmentation, not absence, is now the dominant obstacle for organizations operating across borders.

WHY IT FALLS SHORT

An agent that satisfies Singapore's framework can still fail the EU AI Act or China's labeling rules; fragmentation, not absence, is now the dominant obstacle, and it falls hardest on the cross-border deployers Singapore most wants to attract.

RECOMMENDATION

Use the vehicles Singapore already operates, its AI Safety Institute and its leadership of the ASEAN Working Group on AI Governance, alongside the convening credibility demonstrated at the World Economic Forum, to drive a minimum interoperable core: shared definitions, a common agent-identity schema, and mutual-recognition pathways so that meeting the MGF satisfies a defined baseline elsewhere. First-mover leadership on interoperability would be as significant as the framework itself.

08 Create a structured channel for agentic incident reporting and shared learning

FRAMEWORK DIMENSION	GAP THEME	PRIORITY
Cross-cutting: Ecosystem learning	Incident reporting & shared learning	MEDIUM

THE GAP

The framework tells each organization how to govern its own agents, but creates no mechanism for the ecosystem to learn from failures. When an agent takes an erroneous or unauthorized action, or a multi-agent chain fails in a novel way, that knowledge stays inside one deployer. Notably, Singapore's own Model AI Governance Framework for Generative AI (2024) names incident reporting as one of its nine dimensions, and the EU AI Act mandates serious-incident reporting for high-risk systems, yet the agentic framework does not carry an ecosystem-level incident-learning mechanism forward.

WHY IT FALLS SHORT

When each deployer learns only from its own failures, the same cascading multi-agent error is discovered independently many times over, and IMDA loses the empirical evidence base it needs to sharpen future versions of the framework.

RECOMMENDATION

Carry the generative-AI framework's incident-reporting dimension into the agentic framework, operationalized as a voluntary, safe-harbor channel for reporting significant agentic incidents to IMDA, with anonymized patterns published periodically, in the spirit of aviation-style incident learning rather than enforcement. Deployers gain early sight of failure modes such as cascading multi-agent errors before they experience them, and IMDA gains an empirical evidence base for future versions of the framework. Confidentiality and a clearly non-punitive posture are what make such systems work.

09

Give affected end users a defined path to contest an agent's action

FRAMEWORK DIMENSION	GAP THEME	PRIORITY
Dimension 4: Enable end-user responsibility	End-user recourse	HIGH

THE GAP

Dimension 4 rightly equips end users to use agents responsibly, but the reverse direction is thin: a person affected by an agent's autonomous action, a denied application, a mis-executed transaction, a wrongly closed account, has no defined route to contest it, appeal it, or have it corrected. Accountability that the affected person cannot invoke is accountability in name only.

WHY IT FALLS SHORT

Accountability that the affected person cannot invoke is accountability in name only: a wrongly denied application or mis-executed transaction has no defined route to be contested, appealed, or corrected, leaving the framework's promise incomplete.

RECOMMENDATION

Add an expectation that deployers of consumer-facing agents provide an accessible recourse path: a way to reach a human, a defined timeframe for review, and a documented correction-and-remedy process when an agent's action is found to be wrong. Pair it with the Agent Identity Card so the disclosure a user sees includes how to challenge the agent's decision. This completes the framework's accountability loop from the perspective of the people agents act upon.

The Principle That Ties Them Together

These nine recommendations are not a grab bag. They share a single organizing idea, one the framework already embraces in spirit and could now express more fully: accountability must scale with autonomy.

Making Agent Identity Cards verifiable rather than merely descriptive, adding a trigger that re-opens governance when autonomy rises, publishing measurable thresholds, tiering expectations by stakes, building an attestation pathway, creating a shared incident-learning channel, giving affected users a route to contest an agent's action, and leading on cross-border interoperability are all expressions of that one principle. Singapore has already given the world the best map of what responsible agentic AI looks like. The opportunity now is to make that map enforceable, and to lead again.

Accountability must scale with autonomy.

The framework can now make that real.

Summary of Recommendations

A one-line index of the nine recommendations, the framework dimension each addresses, and its indicative priority.

#	Recommendation	Framework dimension	Priority
1	Resolve liability apportionment, do not just assign accountability	Dimension 2: Make humans meaningfully accountable	CRITICAL
2	Make Agent Identity Cards runtime-verifiable, not just a disclosure format	Dimension 1 & 3: Bound risks / Technical controls	HIGH
3	Add an authority-creep trigger to the graduated autonomy taxonomy	Dimension 1: Assess & bound risks upfront	CRITICAL
4	Publish quantitative acceptance thresholds and finalize agentic testing standards	Dimension 3: Implement technical controls	HIGH
5	Introduce risk tiering with proportionate expectations for high-stakes agents	Cross-cutting: Proportionality	HIGH
6	Build an assurance and attestation pathway for agents	Dimension 3: Implement technical controls	MEDIUM
7	Lead on cross-border interoperability, not just national guidance	Cross-cutting: Interoperability	HIGH
8	Create a structured channel for agentic incident reporting and shared learning	Cross-cutting: Ecosystem learning	MEDIUM
9	Give affected end users a defined path to contest an agent's action	Dimension 4: Enable end-user responsibility	HIGH

Frequently Asked Questions

What is Singapore's MGF for Agentic AI?

It is the Model AI Governance Framework for Agentic AI, best-practice guidance from Singapore's IMDA, first launched in January 2026 and updated to Version 1.5 in May 2026. It is the world's first governance framework written specifically for AI agents, organized around four dimensions: assess and bound risks upfront, make humans meaningfully accountable, implement technical controls, and enable end-user responsibility.

Is the Singapore MGF mandatory?

No. It is voluntary best-practice guidance with no standalone penalty. Organizations remain legally accountable for their agents under existing law, and sector regulators may reference it when assessing incidents, but adopting the framework itself is not legally required.

What does Version 1.5 add over the January 2026 launch?

Version 1.5, published on 20 May 2026 and updated on 5 June 2026, incorporated feedback from more than 60 organizations and added over 10 real-world case studies, along with expanded guidance on multi-agent systems, third-party agents, and automation bias. It also introduced Agent Identity Cards and a four-level graduated autonomy taxonomy.

What are the main gaps still open in the MGF after v1.5?

By the assessment of practitioners and researchers, the genuinely open issues are: how legal liability is apportioned among providers, developers, deployers, and users; whether Agent Identity Cards can be made runtime-verifiable rather than merely descriptive; the absence of a trigger that re-opens governance when autonomy is raised over time; the lack of quantitative acceptance thresholds and finalized testing standards; the absence of risk tiering; the lack of an assurance and attestation pathway for agents; the absence of an ecosystem-level incident-reporting and learning channel; the lack of a defined recourse path for people affected by an agent's action; and cross-border fragmentation.

What is an Agent Identity Card?

It is a standardized disclosure format introduced by Singapore's MGF that specifies an AI agent's capabilities, limitations, authorized action domains, and escalation protocols. It is a genuine contribution to agentic AI governance; the recommendation in this submission is to evolve it from a disclosure format into a cryptographically verifiable credential that systems can check and enforce at the moment an agent acts.

What is the single principle behind these recommendations?

Accountability must scale with autonomy. Each recommendation, from verifiable identity to an authority-creep trigger to risk tiering, is a way of ensuring that as an agent is given more freedom to act, the governance, verifiability, and answerability around it grow in step rather than lagging behind.

References and Further Reading

1. IMDA (Singapore), Model AI Governance Framework for Agentic AI, launched 22 January 2026 at the World Economic Forum; Version 1.5 published 20 May 2026 and updated 5 June 2026; IMDA discussion paper on Legal Responsibility for AI Agents (May 2026).
2. Global Policy Watch and Baker McKenzie, on the Version 1.5 update incorporating feedback from more than 60 organizations and adding multi-agent, third-party-agent, and automation-bias guidance (2026).
3. Lexology and Stephenson Harwood, on open legal-liability questions, including apportionment among stakeholders and whether an agent's autonomous action can bind its principal (2026).
4. OpenID Foundation and AI-identity research, on delegated authorization, transitive trust, and cross-domain federation, and the insufficiency of current standards for agentic deployments (2026).
5. Research on machine identity and agent governance, describing Agent Identity Cards and the four-level graduated autonomy taxonomy (2026).
6. Analyses of cross-jurisdiction fragmentation, including the EU AI Act Article 50 and differing national labeling regimes (2026).
7. IMDA AI Verify, Singapore's AI governance testing and assurance framework, and the Global AI Assurance Pilot pairing specialist testing firms with real deployments.
8. Cyber Security Agency of Singapore, Securing Agentic AI: An Addendum to the Guidelines and Companion Guide on Securing AI Systems (announced at SICW, 22 October 2025; consultation to 31 December 2025; finalized 2026).
9. IMDA, Starter Kit for Testing LLM-Based Applications for Safety and Reliability (2025), and the agentic-AI testing guidelines in development on top of it.
10. IMDA and AI Verify Foundation, Model AI Governance Framework for Generative AI (May 2024), including its incident-reporting dimension; PDPC, Proposed Advisory Guidelines on Use of Personal Data in Generative AI (2 June 2026).

ABOUT THIS SUBMISSION

Prepared by Prashant Akhawat as part of the CXO Intelligence Series on Governance. Offered as a constructive, voluntary submission to IMDA. References to the ISO/IEC 42001 standard and to third-party regimes are for alignment and context and are paraphrased, not quoted. This is an educational and policy contribution, not legal advice.

Prashant Akhawat · akhawat.com · linkedin.com/in/prashantakhawat

Cite As

Akhawat, P. (2026). *Strengthening the World's First Agentic AI Governance Framework: Nine Recommendations for Singapore's Model AI Governance Framework for Agentic AI*. figshare.
<https://doi.org/10.6084/m9.figshare.33137627>