

JULY 2026 | ISSUE 2

COUNTERPOINT INDO-PACIFIC

A Publication of the Centre on Asia and Globalisation



Image Credit: ChatGPT

AI in the Indo-Pacific: A Path to Gains or Instability?

By Miguel Gomez

In his speech during the 2023 Singapore Conference on AI for the Global Good, then Singapore Deputy Prime Minister and Minister for Finance Lawrence Wong **noted** that “breakthroughs in AI have sparked renewed interest – about the potential of AI, its risks, and its implications for humanity. Regardless of our individual perspectives on these issues, one thing is clear – AI now stands as a pivotal game-changing technology that impacts us all.” These sentiments reflect the dual perspectives that color the beliefs surrounding artificial intelligence – such as the **utility and threat posed by frontier models** in the **domain of cybersecurity**. On the one hand, we stand at a technological



Counterpoint Indo-Pacific is published regularly by the Centre on Asia and Globalisation at the National University of Singapore's Lee Kuan Yew School of Public Policy. It seeks to answer major questions of strategic significance for Indo-Pacific by bringing in diverse voices from around the region. Each issue will tackle one question from different perspectives.

Centre on Asia and Globalisation

+65 6516 7113

cag@nus.edu.sg

469A Bukit Timah Road,

Tower Block Level 10,

Singapore 259770

<https://lkyspp.nus.edu.sg/cag>

cont'd p2

threshold wherein AI is increasingly embedded into the socio-technical fabric of modern society – a critical component of modern life. On the other, the specter of uncertainty concerning its social, political, economic, and technological risks that follows from our growing dependence on the technology is a reality faced by states across the globe as policymakers continue to grapple with the benefits and consequences that this emergent, and at times disruptive, technology brings.

Within the Indo-Pacific, the urgency to understand the implications of artificial intelligence is especially pronounced. Surveys and subsequent reports from the [University of Melbourne](#) and the [Pew Research Center](#) suggest a distinct divide between several countries in the region and their counterparts in Europe and North America. Whereas the latter appear less trusting of the technology, the former seem more optimistic about the benefits that it brings. Nevertheless, this divergence serves as one of a handful of reasons that motivate this issue of Counterpoint Indo-Pacific in which we ask **whether AI in the Indo-Pacific offers a path towards gains or instability**.

To answer this question, we frame our approach across four interrelated themes that tackle AI's impact on **(1) competition and conflict, (2) socio-economic realities, (3) technological governance, and (4) ethics and norms**. Keen-eyed observers will note

that these reflect the prevalent debates that transcend regional divides and shape policy preferences amongst those tasked with managing the use and effects of this technology. Furthermore, these choices echo our conscious attempt to treat artificial intelligence not as a material artefact that independently exerts influence at an individual, domestic, and international level. Instead, the essays contained in this issue approach AI as a socio-technical ensemble that emerges from the advancing technological know-how that is employed in and through dynamic, and often uncertain, social and political systems. To aid us in this task, this issue of Counterpoint Indo-Pacific engages with thought leaders who bring these themes into the context of the region's ongoing relationship with this technology. They presented their views in the [2nd Counterpoint Indo-Pacific \(CIP\) webinar](#).

This issue starts with **Michael Raska** discussing if and how the adoption of artificial intelligence across militaries in the region will bring increased (in)stability. The growing relevance of Military AI is often seen as a key determinant for geopolitical stability across the globe. In it he asserts that these outcomes are contingent on the interplay between regional countries' technical/economic capacities and the flexibility of defense institutions to absorb and integrate these technologies. As such, instability is not a foregone conclusion but depends greatly on if and how countries apply AI in pursuit of their strategic objectives.

Complementing this argument, **Mark Mananatan** takes a broader view to establish how countries, as a whole, are likely to pursue AI capabilities considering persistent technological limitations and interdependencies. In so doing, he proposes that regional countries are likely to follow three possible scenarios: (1) AI Fragmentation, (2) Sovereignty, or (3) Resilience. While none of these is a foregone conclusion, his essay surfaces the necessary compromises that each scenario brings and that, ultimately, it is unlikely that any one country can pursue these capabilities independently.

Moving away from the state-centric view of these debates, **Shaleen Khanal** foregrounds the prominent role of technology companies (i.e. Big Tech) in governing this technology. In so doing, his essay surfaces the unique character of AI in which governments do not hold a monopoly on governance but, instead, rely on the cooperation of the private sector for both sustained development and appropriate risk mitigation strategies.

Finally, we close this issue with an essay from **Massimiliano (Max) Cappuccio** on the ethical challenges that artificial intelligence brings. Specifically, his essay moves beyond simply calling for dialogue on ethics, demanding real and sustained collaboration across stakeholders instead. Noting national-level variance on how the technology is perceived, success hinges on going beyond mere discourse; instead, seeking means to overcome ideological

differences vis-à-vis AI is necessary if we are to overcome the ethical pitfalls that follow the use of AI.

While our contributions to this debate remain necessarily modest by virtue of both its format and the scale of the issue itself, it is nonetheless necessary. As artificial intelligence continues to develop and embed itself within modern society, academics and practitioners alike can no longer afford to adopt a “wait-and-see” attitude. Despite the obstacles that it faced over its history dating back to the pivotal workshop at Dartmouth College in the 1950s, artificial intelligence will continue to play a crucial role across the region and the world. Seeking greater understanding of this technology and how it interacts in the social, economic, and political environment is fundamental to managing its risks.

Miguel Alberto Gomez is a Senior Research Fellow at the Centre on Asia and Globalisation (CAG), Lee Kuan Yew School of Public Policy (LKYSPP). His research explores how emerging technologies such as cyberspace and artificial intelligence reshape security, strategy, and political behaviour. His work appears in leading journals such as the International Studies Quarterly, Journal of Peace Research, and International Interactions. He earned a PhD in Political Science from the University of Hildesheim, an MA in International Security from the Barcelona Institute of International Studies, and a BSc in Computer Science from De La Salle University.

Military AI Trajectories and the Future of Strategic Stability in East Asia

By Michael Raska



Image Credit: Wikimedia Commons/Tech. Sgt. Dave Nolan

In the twenty-first century, strategic stability in East Asia can be characterised as a condition of **contested peace**. In a contested peace, strategic stability has been undermined by recurring “hybrid” actions below the threshold of armed conflict. These actions are coercive in both military and non-military forms and aim to alter the political and strategic status quo. In particular, defence policymakers and militaries in East Asia must contend with expanding and persistent forms of **“grey-zone” operations**, including the deployments of maritime militia and coast guard operations, cyber and information campaigns, and increasingly intensive air and naval operations around disputed territories and **maritime domains**. Such actions have essentially blurred the boundaries between peace and conflict in the region, heightening the risks of potential military miscalculation, mistrust, and **inadvertent escalation**.

The key question now, however, is how the

diffusion and military integration of artificial intelligence (AI) will affect the region’s already fragile condition of contested peace. Will AI strengthen deterrence and regional stability by improving situational awareness, operational decision-making, and overall resilience? Or will it intensify existing military competition and rivalries, compress decision-making timelines, and enable the use of autonomous weapons systems, creating new pathways towards miscalculation and inadvertent escalation through confrontations in the grey zones?

Diffusion of Military AI in East Asia

Military AI in East Asia cannot be understood apart from the broader defence transformation already underway across the region. As **Richard Bitzinger noted**, this transformation has extended well beyond conventional force modernisation or the

acquisition of new military platforms. Instead, regional militaries have pursued the convergence of advanced capabilities, including long-range networked precision strike, autonomous systems across the air, land, maritime, and undersea domains, advanced C4ISR networks, cyber and space capabilities, and increasingly AI-enabled systems and applications for command-and-control, intelligence analysis, targeting, logistics, cyber operations, and autonomous systems.

The diffusion of AI-enabled military capabilities across East Asia, however, has been **neither linear nor uniform**. While the underlying technologies expand through global commercial innovation ecosystems, their military adoption in the region has been conditioned by differing technological capacities, defence-industrial bases and resources, organisational structures, strategic cultures, alliance relationships, and operational requirements. In other words, **the military “AI waves” in the region** have been progressing along divergent adoption and adaptation trajectories, distinct military innovation paths and patterns.

In this context, the varying national AI innovation ecosystems and military integration capacity have **strategically significant implications** for the region’s security and stability. On one hand, varying military AI capabilities could alter threat perceptions and influence how states make decisions during crises. At the same time,

divergent AI trajectories may reshape regional alliance dynamics by **strengthening cooperation** in data, computing infrastructure, intelligence and capability development, while creating novel technological dependencies and interoperability gaps between allies with differing capacities to integrate and operationalise AI.

Three Tiers of Military AI Trajectories

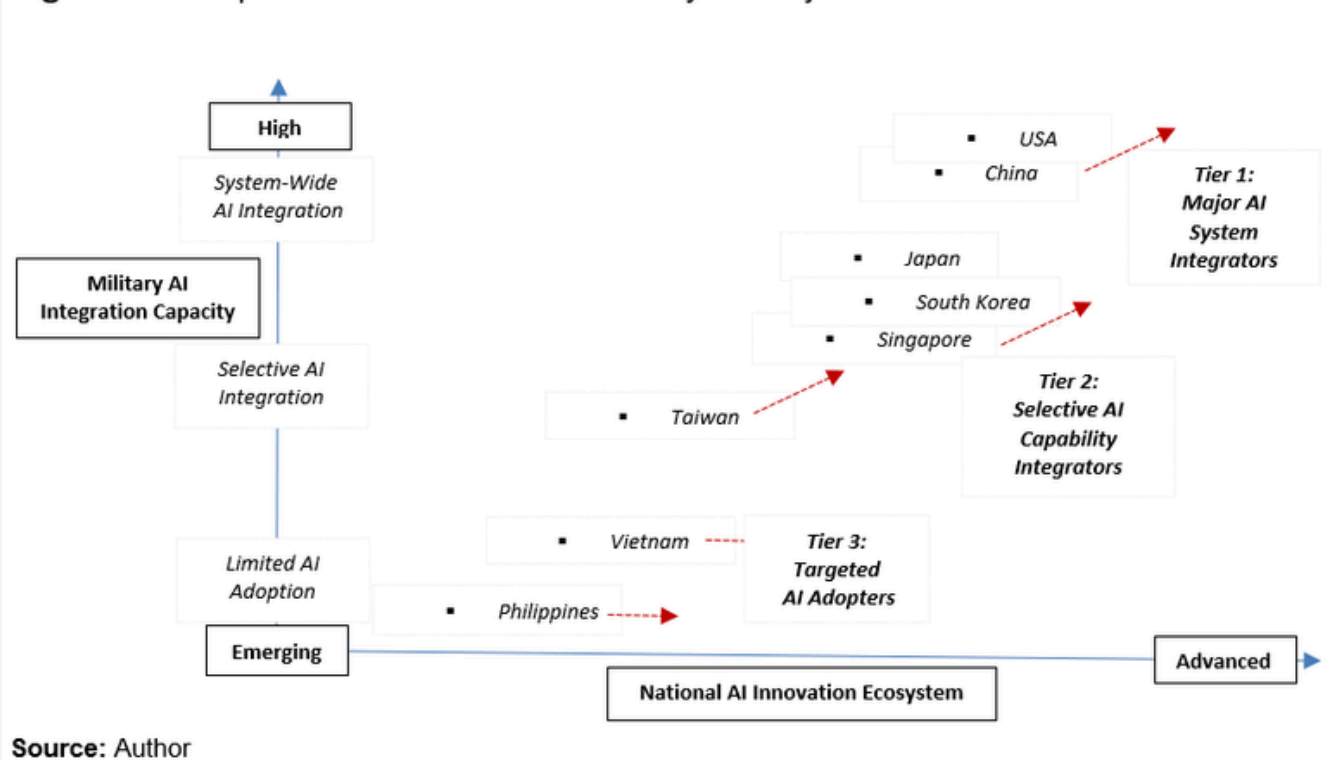
Conceptually, the varying military AI trajectories can be projected in a comparative framework through the interaction of two major systemic variables and their indicators: (1) the maturity of a country's **civilian/commercial AI innovation ecosystem**, and (2) its **military capacity to absorb, integrate, and operationalise AI-enabled capabilities**. Other **conditioning variables** also include civil-military and public-private relations, AI governance, legal and regulatory frameworks, data governance, and overall national approaches to human control over AI-enabled systems.

Using this framework, one can delineate three broad tiers of military AI trajectories, shaping the strategic environment in East Asia. The first tier comprises major AI system integrators (the United States and China), which combine advanced civilian AI ecosystems with the institutional and technological capacity (albeit varying) and strategic requirements to potentially

integrate AI across the “full spectrum” of military capabilities. While China has made remarkable progress over the past decade in developing AI systems at scale through its Military-Civil Fusion strategy and the People’s Liberation Army’s (PLA) pursuit of “intelligentisation” strategy, its military absorption capacity for AI may face constraints relative to the US. These constraints may include uneven organisational integration across the PLA services, and more importantly, lack of operational combat experience in employing AI-enabled systems and capabilities in crises. That said, both China and the US militaries have technological, organisational, and financial resources to pursue military AI transformation at scale. The US has demonstrated the operational use of its AI systems at scale during its 2026 military campaign against Iran.

The second tier can be characterised as selective AI capability integrators, i.e. Japan, South Korea and Singapore, in which advanced national AI ecosystems support increasingly sophisticated, but still selective, patterns of military AI integration. For example, Japan’s Ministry of Defense AI Policy, adopted in 2024, shows seven priority AI application areas, including target detection, intelligence analysis, command decision support, logistics, unmanned systems and cybersecurity. Similarly, South Korea’s Defence Innovation 4.0 strategy emphasises AI-based command systems, manned-unmanned teaming and autonomous platforms. South Korea’s advanced semiconductor and electronics industries provide a strong foundation for these efforts. However, the translation of industrial capacity into operational military capabilities has been constrained by inter-service coordination

Figure 1. Comparative Framework of Military AI Trajectories



problems, established organisational interests, and a strategic culture that has traditionally favoured platform-centred modernisation.

Singapore represents a more concentrated approach to selective AI integration. The establishment of the Digital and Intelligence Service in 2022 consolidated the Singapore Armed Forces (SAF)'s intelligence, C4I, cyber and digital capabilities within a dedicated service. Under the broader SAF 2040 transformation agenda, Singapore has increasingly applied AI to command and control, intelligence fusion, military training and autonomous systems. For example, the SAF's AI-enabled command-and-control systems analyse sensor data and recommend weapon-to-target pairings while retaining human command authority. Meanwhile, Taiwan occupies a transitional position between selective integration and targeted adoption. It dominates the leading-edge chip manufacturing but has translated little of that advantage into broad military AI integration.

The third tier consists of emerging or targeted AI adopters, including Vietnam, Malaysia, Indonesia and the Philippines. In these countries, military AI development remains focused on specific operational requirements, particularly maritime domain awareness, border surveillance, cyber defence, uncrewed systems, and command and control. Their AI trajectories, however, are not uniform. Vietnam has placed greater emphasis on nationally directed technological development, whereas

the Philippines remains more dependent on alliance-enabled access to platforms, intelligence and digital infrastructure. Across this tier, the pace and direction of military AI adoption will depend heavily on foreign technology partnerships, commercial off-the-shelf solutions, defence cooperation with external powers, and the institutional capacity to adapt civilian digital technologies to specific military requirements.

Implications for Strategic Stability

The divergent paths and patterns of military AI adoption and integration in East Asia raise three interrelated issues for the future of strategic stability. The first is the distribution of military balance. In particular, will the diffusion of AI-enabled capabilities reinforce existing military asymmetries, or will they enable smaller states to offset conventional disadvantages? For example, autonomous systems, distributed sensing and networked precision strike could strengthen Taiwan's capacity for asymmetric deterrence by making its forces more dispersed, resilient and difficult for the PLA to detect and suppress. Conversely, China's advantages in scale, data, computing infrastructure and defence-industrial capacity may allow it to integrate AI across a broader range of military functions, further widening existing capability gaps.

The second issue relates to crisis stability. The integration of AI into command and control, intelligence, targeting and autonomous systems may accelerate military decision cycles during crises in the conflict faultlines

across the region, i.e. the Taiwan Strait, the Korean Peninsula and the South China Sea. AI-enabled sensor-to-shooter networks could improve situational awareness and reduce uncertainty, but they could also create pressure to act before information is verified or a perceived operational advantage disappears. The central risk lies in a widening gap between the speed of AI-enabled military operations and the time required for political leaders to evaluate information, interpret adversary intentions and control escalation. In contested “grey zone” environments, tactical actions recommended by AI-enabled systems could therefore be misinterpreted as deliberate strategic escalation.

The third issue concerns strategic opacity and interoperability. States may be unable to determine how extensively an adversary relies on AI, how much authority has been delegated to autonomous systems, or whether a particular action reflects political intent, machine-optimised recommendations or system failure.

Divergent approaches to AI governance and human control could complicate security and defence cooperation among allies. While AI could strengthen intelligence sharing, integrated sensing and joint command networks, differences in data standards, system reliability, legal constraints and decision-making procedures could create new dependencies and interoperability gaps.

Taken together, the diffusion of military AI will not have an inherently stabilising or destabilising effect on East Asia’s security. Its evolving dynamic and consequences will depend on how particular capabilities interact with existing military asymmetries, deterrence and defence strategies, alliance structures and the use of force. The central challenge for all three tiers will be to embed AI-enabled systems within resilient command-and-control arrangements, rigorous testing and assurance processes, effective political control and meaningful human oversight. Without these safeguards, divergent military AI trajectories may reinforce existing tensions, increase uncertainty in strategic decision-making during future crises and conflicts.

Dr. Michael Raska is Senior Researcher and Section Head for Security and Defense in Europe at the Center for Security Studies at ETH Zurich. He is also a Non-Resident Senior Research Fellow at The Hague Centre for Strategic Studies and a member of the Expert Advisory Group of the Global Commission on Responsible Artificial Intelligence in the Military Domain. Previously, he was an Assistant Professor in the Military Transformations Programme at the S. Rajaratnam School of International Studies at Nanyang Technological University. His research focuses on the drivers and trajectories of defense and military innovation, examining how emerging technologies - particularly

artificial intelligence - shape defense planning, force development, intelligence assessment, and the future character of warfare, including hybrid conflict and digital warfare. He is the author of *Military Innovation and Small States: Creating Reverse Asymmetry* (Routledge, 2016) and co-editor of several volumes, including *The AI Wave in Defence Innovation* (Routledge, 2023) and *Defence Innovation and the 4th Industrial Revolution* (Routledge, 2022). His scholarly work has appeared in journals such as the *Journal of Strategic Studies*, *Strategic Studies Quarterly*, *PRISM: Journal of Complex Operations*, and *SIRIUS – Zeitschrift für Strategische Analysen*, among others. His contributions also include chapters in edited volumes, policy reports, and commentaries, including collaborative projects with leading academic, military, and policy research institutions focusing on defense innovation, emerging technologies, and future of warfare.

Who Controls the Stack: Three Strategic Scenarios for AI Governance in the Indo-Pacific

By Mark Bryan Manantan

Over the past few months, the US and China have both demonstrated their capability to wield export control measures over their respective frontier AI developers. On June 12, 2026, the US Department of Commerce (DOC) ordered Anthropic to suspend access to its two most advanced models—Claude Fable 5 and Mythos 5—for any foreign national, including the company’s own employees. With no reliable way to verify nationality in real time, Anthropic disabled both models for its entire global user base within hours. DOC lifted the controls on June 30, in exchange for company commitments on risk detection, standards development, and reporting of malicious activity. Meanwhile, on April 27, 2026, China’s National Development and Reform Commission (NDRC) blocked the USD 2 billion Manus-Meta deal under its Foreign Investment Security Review. The decision extended Beijing’s regulatory toolkit beyond tech transfer, data, and talent to completed overseas transactions.

The reversal does not restore the status quo ante; it confirms that the status quo has changed. For eighteen days, a model widely regarded as the world’s most capable was disabled for every non-US user by a single



Image Credit: Flickr/Ethen Rera

administrative instrument and restored only through a negotiation between one firm and one government in which no other jurisdiction had standing. Access to frontier AI is now a policy variable rather than a commercial one, and the precedent is portable.

Neither move is unprecedented. Over the past decade, both superpowers have refined their export control playbooks. Washington has moved through successive postures: from decoupling to derisking to friendshoring. Beijing has pursued a dual track. It has prioritised self-reliance while expanding research collaboration and diffusing its open-weight models, especially in Southeast Asia. Both converge on the same economic security agenda: who controls the stack, controls the future. Drawing on empirical analysis, this article conducts strategic foresight to outline three scenarios for the Indo-Pacific’s AI governance, to help

scholars, policymakers, and practitioners navigate a more contested and unpredictable region.

Scenario 1: AI Fragmentation

The first and most feasible scenario is AI fragmentation: two competing spheres of influence, one anchored by the US, the other by China. This bifurcation is already visible.

In December 2025, Washington announced Pax Silica, a coalition to streamline the AI supply chain from energy and critical minerals to semiconductors and data centres. Japan, South Korea, and Australia have embraced it, deepening their integration with the US ecosystem.

The window of tech neutrality for Southeast Asia—including close US partners like Singapore and the Philippines—is narrowing. US models remain widely used, but cheaper and strongly multilingual open-weight alternatives from DeepSeek, Alibaba, and Baidu are closing the performance and adoption gaps. The June suspension sharpened this dynamic: Zhipu AI launched its GLM-5.2 flagship the following day. Across the region, enterprises publicly reassessed their dependence on a single American provider. Every disruption in the US stack is an adoption event for the Chinese one. Southeast Asia's deepening reliance on Huawei and ZTE in 5G and data centre's compounds this dynamic.

Both capitals are consolidating jurisdiction. Beijing is expanding its capacity to restrict Chinese firms' overseas investments, transactions, and research partnerships. Washington's AI Diffusion Rule places Singapore and the Philippines under a validated end-user scheme for advanced chips, while the rest of Southeast Asia faces heightened scrutiny and ambiguity.

The Fable–Mythos episode exposed a quieter tiering mechanism operating above the chip layer. Project Glasswing, Anthropic's restricted cybersecurity initiative, expanded in June to around 200 vetted organisations across more than fifteen countries, granting access to vulnerability discovery capabilities unavailable on the open market. Reported participants include NATO, the European Union Agency for Cybersecurity (ENISA), and South Korea's Samsung. Northeast Asia's national champions sit inside the perimeter; Southeast Asia is almost entirely outside it. Frontier cyber defence is now allocated by security vetting and alliance proximity rather than market demand: a validated end-user logic for models rather than chips, with far less regulatory visibility.

In short, both powers can now flex jurisdictional capacity to constrain Southeast Asia's hedging and, at worst, force a choice between Pax Silica and Pax Sinica. Japan and South Korea face the same terrain, exposed to retaliation should they ban Chinese models outright. Partial decoupling results—but interoperability erodes, raising governance

and compliance costs and leaving less secure applications with embedded bias.

Fragmentation and sovereignty assertion are not sequential responses. They are simultaneous. The region's most consequential actors are pursuing both at once.

Scenario 2: AI Sovereignty

The second scenario points to accelerating sovereign AI capabilities. Recognising their exposure to US and Chinese export policy, many countries are investing heavily in building their own.

The Fable–Mythos episode redefined AI sovereignty. For a decade, debates in the Indo-Pacific centred on infrastructure and data residency. The directive to Anthropic restricted access by nationality, irrespective of user location or the data processing jurisdiction. A model can run in an approved region, on compliant infrastructure, over lawfully localised data, and still be unavailable to a government's own officials because of who they are. Data residency is no hedge against eligibility risk—and most regional strategies are built almost entirely around residency.

Northeast Asia is moving fastest, fast-tracking industrial policy to build national champions. Japan's AI Promotion Act, enacted in May 2025, commits nearly JPY 10 trillion to AI and semiconductor infrastructure. South Korea backs its sovereign push with a USD 75 billion

national fund that designates AI data centres as social infrastructure.

Taiwan occupies the most anomalous position. As home to TSMC, it manufactures the chips every sovereign AI programme depends on yet cannot formally join Pax Silica due to PRC sensitivities: indispensable but unrecognised. Taipei has responded with the AI Basic Act, the Tainan Cloud Center, and TAIDE, a sovereign model trained on Taiwanese data to counter the assumptions embedded in Chinese systems. It is building the full stack from below.

Southeast Asia is not idle. Sovereign ambitions are materialising across three layers—legal, infrastructural, and model-level. Vietnam enacted the region's most comprehensive AI law in March 2026, with a risk-based classification system and a National AI Development Fund. Malaysia imposed an AI-only data centre mandate while building ILMU, the region's first sovereign LLM, which outperforms GPT-4o on Malay benchmarks. Indonesia partnered with India on Sahabat-AI, a 70-billion-parameter model supporting Javanese and Sundanese. Southeast Asia is not choosing between Washington and Beijing; it is building a third option.

Scenario 3: AI Resilience

The final scenario is the middle ground. Japan, South Korea, Australia, Taiwan, and ASEAN form ASEAN+Tech Plus—a coalition of small and medium powers pursuing a middle pathway between fragmentation's zero-sum logic and sovereignty's isolation.

Australia, the Philippines, and Indonesia anchor the sourcing and processing layer for critical minerals, while Taiwan, Japan, South Korea, Australia, Taiwan, and ASEAN form ASEAN+Tech Plus—a coalition of small and medium powers pursuing a middle pathway between fragmentation’s zero-sum logic and sovereignty’s isolation. Australia, the Philippines, and Indonesia anchor the sourcing and processing layer for critical minerals, while Taiwan, Japan, South Korea, Malaysia, and Singapore collaborate on advanced chips and data centres. The caveat is structural dependency: US and Chinese frontier models still underpin sovereign LLM development, NVIDIA still holds the most advanced compute, and AWS, Microsoft, Google, Huawei, and Alibaba dominate cloud.

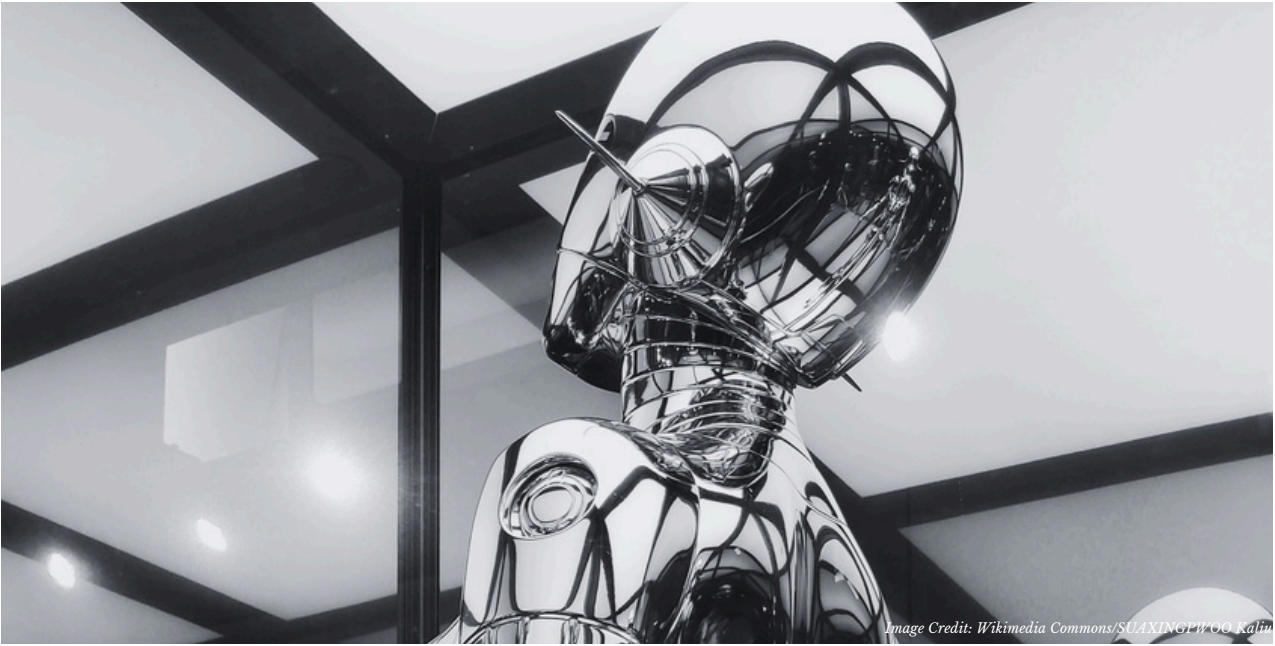
Realising this requires a coordination instrument, and the anticipated ASEAN Digital Economy Framework Agreement (DEFA) provides one. DEFA’s digital trade and data governance architecture can anchor an ASEAN-led AI standards body—with Australia, Japan, and South Korea as full partners and Taiwan as an observer—harmonising sovereign cloud certification and aligning safety benchmarks. The Fable-Mythos resolution makes this urgent. The terms on which the models returned included commitments on pre-release government access, safeguard information sharing, and an industry jailbreak-severity framework drafted with Amazon, Microsoft, and Google. These will function as de facto global standards, and no Indo-Pacific middle power was party to the negotiation that produced them.

Conclusion

The evolving US-China export control rivalry is weaponised interdependence made visible: countries without frontier AI capabilities are structurally exposed to external shocks that threaten both national interest and digital economic development. The three scenarios outlined here represent three different answers to the same underlying question—who controls the stack. Fragmentation accepts the answer as Washington or Beijing. Sovereignty tries to rewrite it nationally. Resilience argues that small and middle powers can collectively redefine it on their own terms. For Northeast Asia’s advanced economies, doubling down on sovereign AI ecosystems is both viable and underway. For Southeast Asia, the capacity gap is real—but not permanent.

The most actionable path forward is deliberate multilateral coordination: diversifying technology partnerships, pooling compute investments, and building the institutional architecture that converts national initiatives into a genuinely interoperable regional AI order. The Fable-Mythos episode offered a preview of how that question gets answered by default: in a private negotiation between one company and one regulator, with every other jurisdiction learning the outcome after the fact. Who controls the stack is still an open question. The window to shape the answer is now.

Dr. Mark Bryan Manantan is research fellow at La Trobe University's Centre for Global Security. He is also the Co-Founder of Balai Labs and a member of UNESCO's AI Experts Without Borders.



AI Governance in the Age of Big Tech

By Shaleen Khanal

In 2023, the Singapore Government launched the AI Verify Foundation to develop verification tools for testing responsible use of AI. The highly acclaimed initiative has led to the development and launch of AI Verify – an open-source toolkit designed for organizations to test and assess their AI stack. What stood out in the [announcement](#) was that five of the seven pioneering premier members of the Foundation were large American tech companies, namely IBM, Google, Microsoft, Red Hat and Salesforce (as of today, Amazon is also part of premier members).

The Foundation's case illustrates the growing role of private actors in AI governance. Governance of emerging technologies, particularly those that are considered highly transformative, has traditionally been regarded as a state responsibility. Private

entities are expected to lead research, development, and commercialisation, while governments are typically expected to set the rules, standards and institutions that govern them.

AI systems challenge this conventional governance model. Similar to other emerging technologies, AI systems' research and deployment are led by private actors. However, these entities, widely known as Big Tech, are no ordinary companies. They are "Big" because they rank among the world's largest companies, with their market capitalisation exceeding the GDP of some advanced global economies. They also serve billions of users worldwide. This scale gives them unparalleled access to talent, high-quality data, computational power, and perhaps most importantly, capital.

Another crucial distinction to note here is that Big Tech are not merely AI companies but digital platform companies. Network effects in the platform economy drive these firms toward monopoly and sectoral expansion. As such, we see Google's presence in health services or Microsoft providing smart farming services.

Owing to their control over the aforementioned resources that are central to developing AI systems, the same companies dominate the AI ecosystem as well. A few moonshot developers such as OpenAI or Anthropic have emerged, but they operate within the same digital ecosystem and depend heavily on Big Tech.

Therefore, Big Tech are highly influential in shaping AI governance. Understanding the influence of any actors in the governance of AI systems requires examining three questions: who defines the problems that demand attention, which policy tools are chosen to address them, and who has the political influence to shape those choices. Given the abundance of technical, financial and policy tools available at their disposal, Big Tech are shaping all three answers. Big Tech, over time, have played a crucial role in shaping the public imagination and scientific knowledge of AI systems. Public surveys across the Indo-Pacific suggest that trust in AI is closely tied to trust in Big Tech, underscoring these firms' influence over the public's perceptions of the technology. They have also positioned themselves as key suppliers of AI solutions and the digital infrastructure that supports them, leaving

governments with few viable alternatives. Their political influence was equally evident in the prominent presence of tech leaders at Donald Trump's inaugural ceremony, underscoring Big Tech's proximity to centres of political power.

The most tangible effect of such overarching influence can be seen in terms of AI ethics washing. Big Tech have been largely successful in persuading the public and governments that self-regulation and technical improvements are the best solution to the majority of social and political challenges that AI systems pose.

It is important not to treat Big Tech as a technological or political monolith with homogenous interests or strategies. Over the past few years, at least two competing yet complementary AI ecosystems have emerged: the American and the Chinese. The American ecosystem's strength lies in frontier models, while the Chinese ecosystem has established a lead in embodied and autonomous AI systems. These technological differences also shape the incentives and material interests of companies within each ecosystem. Frontier models, data, and compute do not need to be localised. They can be built and served from almost anywhere. Most American firms do not face entry barriers globally and therefore do not really "need" Asia. This is not the case for Chinese firms. Technologically, autonomous systems require physical deployment. More importantly, Chinese firms are increasingly being excluded from much of the West. This makes Asia, and Southeast Asia in particular, an important gateway to global

markets. As a result, Chinese firms are likely to pursue deeper engagement with local and national governments in Southeast Asia than their American counterparts that can adopt a more arm's length approach.

Regardless or perhaps because of the differences in technologies, interests and approaches, both ecosystems and subsequently both sets of companies are crucial for the region, economically, and politically. Big Tech companies from both ecosystems are now leveraging the ongoing AI rivalry to position themselves as standard bearers of their respective governments' national AI ambitions. Governing AI in the Indo-Pacific is therefore not simply a matter of understanding technological risks. It also means managing influential, resourceful, and politically astute firms that draw, however informally, on the backing of a great power with its own stake in the outcome.

Dr Shaleen Khanal is a Research Fellow at the Centre for Trusted Internet and Community at the National University of Singapore. His research lies at the intersection of organisation theory and policy design, with a particular focus on how governments and public organisations build the capacity to govern emerging technologies. His current work examines the governance of innovations such as robotics, autonomous vehicles, and artificial intelligence.



Who Controls AI? Nationalising Trust in the Indo-Pacific and Beyond

By Max Cappuccio

Intelligent technologies cannot be effectively governed in the absence of public trust. In Western societies, public anxiety about AI is typically construed as the fear that humans may become unable to control intelligent machines and their societal effects. But people do not fear AI simply because it is powerful; they distrust the organisations that develop, own, and decide how to use that power. Trust reflects less technological literacy than confidence in government, private enterprise, national development narratives, and the authorities expected to manage technological risks.

Cultural specificities influence how nationally coordinated AI programs are perceived, soliciting greater or smaller appetites for state-directed interventions for technology development and adoption,

industry centralisation, and infrastructure control. Western publics prevalently tend to foreground privacy, dignity, employment, and democratic accountability; East Asian debates often give greater weight to technical reliability, collective security, social coordination, and national capability.

Confidence in AI is significantly higher in China, India, Singapore, and parts of Southeast Asia than in the United States, Australia, Japan, and most European countries. These trust differences are especially relevant in the Indo-Pacific, where national AI programmes reflect different political economies, innovation ecosystems, cultural narratives, and geostrategic ambitions.

Trust tends to decrease wherever AI ambition is perceived as an expression of corporations

inherently geared toward profit maximisation, market dominance, data extraction, labour reduction, and political influence. Because private agendas need not reflect sovereign goals or answer public concerns, fear of AI is engendered by distrust of the socio-technical and politico-economic systems influenced by corporate power. Before asking whether AI is safe or reliable, the public asks who governs it, who benefits, and who bears the risks.

Conflicting and opaque narratives increase their distrust. Workers are told that automation will raise productivity but are rarely assured that it will improve job security or working conditions.

Professionals are sold tools trained on their knowledge and designed to monitor, deskill, or replace them. Citizens encounter facial recognition, behavioural prediction, and mass-surveillance from the same systems entrusted with their security. The prevalent fear is not that machines will escape control, but that private actors will control them to pursue goals that do not represent public values, tailoring AI to their own interests.

This is why ethical discourse alone often fails to generate trust. Corporate commitments to fairness, transparency, privacy, responsibility, and human-centred design **remain abstract** when disconnected from ownership, labour protection, democratic oversight, and social justice. Ethics washing occurs when institutions debate how AI should behave while excluding who should control AI and what purposes it should serve. A company may

adopt an ethics charter while automating employment, appropriating data, expanding surveillance, lobbying against regulation, or controlling strategic infrastructure.

These concerns intensify when technology firms exert influence over the defence institutions they work for/with. Autonomous weapons, decision-support systems, cyber capabilities, and surveillance technologies are prevalently developed within dual-use ecosystems integrating civilian industry, universities, intelligence agencies, and armed forces. In this context, companies involved in cloud computing, data analytics, autonomous systems, and generative AI are no longer merely suppliers. Some present themselves as strategic actors responsible for national security, capacity resilience, and military superiority. Recent controversies surrounding **Palantir** and **Anthropic** illustrate this shift. Palantir's leadership portrays technological development as politically non-neutral and links innovation to hard power, civilisational competition, and military recruitment. **Critics interpret this as an attempt to elevate corporate-state power over democratic deliberation**, observing that major firms are moving from supplying technical means to defining political ends.

Proposals to nationalise AI, or to place critical parts of its infrastructure under public ownership and strategic direction, respond directly to this anxiety. Nationalisation reduces neither to expropriation nor to centralisation of productive capabilities, but it may implicate public equity, sovereign wealth funds, state compute infrastructure,

strategic procurement, binding access requirements, or government-led coordination of research and development. These arrangements differ in purpose, but all recognise that AI has become too consequential to be treated as an ordinary private commodity.

Three approaches illustrate the range.

Bernie Sanders's social-democratic model would acquire a substantial public stake in leading AI companies through a sovereign wealth fund, giving citizens voting power and a share of the wealth generated by automation. **Donald Trump's nationalist approach** is less concerned with redistribution than with strategic leverage: public stakes, procurement power, and negotiated partnerships would align frontier firms with national industrial, security, and geopolitical objectives. **Xi Jinping's five-year plan** is more comprehensive and integrates **political governance**, public investment, universities, private enterprise, and military-civil fusion within a long-term programme of national modernisation and **capability expansion**.

East Asian countries have spearheaded state programmes that embed AI development within coordinated industrial, research, and security strategies. Their approaches, however, have diverged.

In China, AI is presented as part of societal modernisation, strategic sovereignty, administrative coordination, and military transformation, with ample space for private entrepreneurship and decentralised

capabilities. Its development is embedded in a close relationship among central and local governments, universities, private enterprises, and the armed forces. Public optimism may therefore be linked to the expectation that corporate goals must align with a broader political vision and national values. This model can accelerate capability development and reduce foreign dependence, but some Western commentators have raised concerns about surveillance, authoritarian control, aggressive militarisation, and the subordination of individual rights to state-defined goals.

South Korea offers a **developmental-security model**. AI is connected to economic growth, technological leadership, demographic pressures, and the continuing threat from North Korea. State-supported firms, research institutions, and defence organisations treat autonomous systems as force multipliers capable of compensating for personnel shortages, strengthening deterrence, and supporting a globally **competitive defence industry**. Legitimacy and trust are primarily associated with competence, preparedness, strategic necessity, and technological reliability but also with cultural cohesiveness.

In Japan, demographic anxieties, cultural familiarity with robots and narratives presenting machines as companions and co-workers coexist with constitutional constraints that limit dual-use solutions. Narratives stressing a presumed affinity of Japanese people for robotics **don't necessarily reflect unrestricted acceptance** and have often

been mobilised to support divisive policies, including the [adoption of robots in eldercare](#) and as a response to labour scarcity. Japan therefore shows that cultural familiarity and national need can't assure public legitimacy if citizens distrust their institutions.

These cases suggest that the regional AI race has encouraged strategic synergies among governments, corporations, universities, and defence institutions. Globally, nationalisation may contribute to establish similar synergies using different solutions, ranging from public equity and sovereign funds to state-led research programs and military-civil coordination. Each redistributes authority differently.

Integrated ecosystems can win popular trust insofar as they make technological power publicly accountable, politically contestable, and directed toward collective goals such as social stability, human flourishing, strategic autonomy, and protection of identities.

Public involvement can increase trust when it makes the benefits of innovation tangible, protects employment and public services, and gives citizens or national institutions meaningful influence over technological priorities.

Military leaderships appreciate national programmes' strategic advantages: sovereign access to compute, data, models, cloud infrastructure, testing facilities, and skilled personnel; continuity of service during

crises; stronger supply-chain security; and reduced dependence on foreign vendors whose commercial or political decisions may constrain national action. Owning a weapon platform is insufficient if the decisive model, data pipeline, cloud layer, or update process remains controlled by a foreign or unaccountable contractor. This contributes to explain why military AI has yet to fully gain the trust of citizens, [service members](#), and [military leaderships](#). National programs may help mitigate their concerns, preserving sovereign access, auditability, operational continuity, and freedom from vendor lock-in.

Yet nationalisation is not automatically ethical, let alone democratic. Without transparency, contestability, and institutional checks, nationalisation may [“accelerate the corporate-government fusion we’re already sliding toward,”](#) promoting state monopolies that pursue surveillance and militarisation without factually reducing tech oligarchies' influence. State ownership is not rewarding in itself but only insofar as it promotes publicly accountable technological sovereignty.

Nationalising trust, not merely assets, will prove key to secure sovereign capability while ensuring that technological power remains politically legitimate, socially beneficial, and open to public scrutiny.

Dr Massimiliano (Max) Cappuccio currently is a senior researcher affiliated with the Creative Robotics Lab of UNSW Sydney.

THE CENTRE ON ASIA AND GLOBALISATION

The Centre on Asia and Globalisation is a research centre at the Lee Kuan Yew School of Public Policy, National University of Singapore. It conducts in-depth research on developments in the Asia-Pacific and beyond, and aims to provide academics, decision-makers, and the general public with objective analysis on issues of regional and global significance. The Centre's motto "Objective Research with Impact" reflects its commitment towards ensuring that its analysis informs policy and decision makers in and about Asia.

OTHER CAG PUBLICATIONS



**Counterpoint
Southeast Asia**



China-India Brief



ASEAN Bulletin



Compiled and sent to you by Centre on Asia and Globalisation
and the Lee Kuan Yew School of Public Policy, National University of Singapore

Counterpoint Indo-Pacific is supported by Wilmar International Limited

Feedback or comment?

Contact our team: scarlet@nus.edu.sg

Subscribe: cag-lkyspp.com/CIP-Subscribe