

Capacity Lacking Consequence: How Twenty AI Middle-Powers Govern General-Purpose AI

Four governance functions, twenty jurisdictions, and a chain whose links are rarely joined.

Josephine Schwab [Research Team Lead], Nathan Naidoo, Ferruccio Barazzutti, and Sheryn Lee
with Caio V. Machado (The Future Society) • August 2026

Access the paper: <https://www.arcadiaimpact.org/aigt/research/s26-gpai-mapping-report>

Summary

- **Convergence on subject matter, divergence in structure.** Twenty AI middle-powers¹ address the same governance functions which we take as together forming an accountability chain – identifying a model's systemic risks, testing claims about it, prohibiting certain uses and outcomes, and reporting on serious incidents – and six legal architectures recur. However, where one jurisdiction may cover a particular function, another leaves it empty. (Table 1)
- **Divergence in force.** Half of the jurisdictions hold no binding GPAI provision at all, and of those that exist across the sample, the European Union and Brazil hold more than half between them. GPAI definitions are also carried more in soft law and strategy documents.
- **One structure repeats.** The majority of institutions carrying GPAI governance were built for something else, so obligations land on the actors they already reached, and the bodies with the capacity to evaluate are not the bodies with the authority to act on what they find.
- **These are drafting gaps in domestic law.** These gaps do not need to wait on international agreement about which AI capabilities should be off-limits.

Governing a technology whose limits are set elsewhere

The most capable general-purpose AI (GPAI) models are trained in two countries but the risks they carry – while they are bought, deployed, and depended on – land globally. That leaves a large group of states, including advanced economies, regional rule-makers, substantial markets, which we call "AI middle-powers", trying to govern a technology whose limits are set elsewhere by companies they can neither license nor tax.

¹ The jurisdictions: Australia, Brazil, Canada, Chile, the European Union, France, Germany, India, Israel, Japan, Kenya, Nigeria, Peru, Singapore, South Africa, the Republic of Korea, Switzerland, Taiwan, the United Arab Emirates, and the United Kingdom.

However, these AI middle-powers have not been idle. Over the past three years they have produced statutes, codes of practice, national strategies and regulatory guidance addressing GPAI. While existing comparative inventories mainly outline which governance instruments a jurisdiction holds and whether they are currently in force, few extend to the provision level, which our research is based on.

Across these twenty AI middle-powers, we read the primary text of 101 governance instruments, provision by provision, producing 587 records: 382 of these impose or recommend something, and 205 record a governance function as absent.

We asked: who has to identify the systemic risks a model presents, and does that assessment reach anyone who can act on it? Who has to test whether claims about the model hold, and to what standard? What is prohibited outright, and who is charged with noticing a breach? Should a serious AI incident unfold, who must tell whom, and what follows if they do not? Read in order, these governance areas form a chain, which works only if the links are joined.

"Jurisdictions are most precise about what GPAI is exactly where precision carries the least consequence."

The governance exists, but the links between them do not

Systemic risk assessment

Fewer than a quarter of risk assessment records require the result to reach a named authority, and 65 percent expressly do not. The duty falls almost equally on public authorities and on providers – where a company assesses its own model, the analysis stays with the company. And what gets assessed is mostly not what a model can do: cross-cutting risks, fundamental rights and cyber lead the assessed-for risk domains, while loss of control and CBRN uplift appear much less frequently.

Evaluation is nearly universal, and almost never binding

Sixteen of the twenty jurisdictions require some form of GPAI model evaluation, but only six of recorded provisions sit in law that binds and is in force. Only two provisions in the entire sample actually reach the model developer: Article 55(1)(a) of the EU AI Act, and a Korean certification scheme which developers may use (but are not required to). Where evaluation is required, the party doing the testing is usually the party being tested, such as a developer that checks what it built, or a government agency checks what it has bought. Where a duty to test exists, the instrument names the method (red-teaming, controllability testing) but says nothing about the standard: no benchmark, no threshold, no failure condition. Outside the EU, nothing says what a bad result obliges anyone to do.

Prohibitions target conduct, not capability

Over 70 percent of recorded prohibitions forbid a use of AI, fewer than 19 percent a capability. That decides when a rule can be enforced: a capability can be tested before release, a use only after deployment. The international red-lines debate is built on the smaller category, and domestic law is not positioned for it. Prohibitions cluster where existing law already forbade the conduct and had someone to enforce it – election integrity, children's safety, human rights – while in the domains that worry AI safety researchers most the cupboard is close to bare. One prohibition in the sample concerns CBRN uplift, and it sits in the voluntary EU GPAI Code of Practice. None addresses offensive cyber capability, and none prevents an AI system taking autonomous control of critical infrastructure.

Serious incident reporting

Half the jurisdictions impose no reporting duty of their own, most of the provisions fail to state when a report is due, or what follows if none is made, and every duty that exists asks an actor to report on its own systems. Nobody is required to look across them, so a failure appearing in several models at once has no authority designated to receive the report.

The institutions were built for something else

Of the 144 public bodies named across these instruments, 113 held their mandate before the AI governance question arrived, and nine are dedicated AI regulators. Most are ministries and sectoral supervisors: financial regulators, data protection authorities, online safety bodies.

Obligations land where those institutions already have reach, such as in procurement, public-sector use, what online platforms carry, and what supervised firms rely on. Whilst eleven jurisdictions have set up an AI safety institute or equivalent, all but one lack the mandates to do something about an evaluation result, and were constituted to advise.

"These institutes were built to measure, but constituted without the power to act on what they find."

Recommendations

Safety institutes: amend the mandate first

Verification capacity is being built. What is missing is the wire between a finding and a consequence – a duty to act, a report to a named authority, a bar on market access. On our reading, closing that is a question of what existing bodies are permitted to do rather than of resourcing, though the mapping shows the pattern and cannot establish the cause.

Sectoral regulators: adapt the trigger, not only the channel

Where GPAI duties travel down channels built for financial crime or data breaches, a model failure producing neither may activate nothing. Israel's anti-money-laundering authority shows inheritance is not inevitable – it instructs reporting bodies to file where suspicion arises through generative AI or deepfake technology, so the channel is borrowed and the trigger adapted.

Reporting is where coordination starts closest to scratch

Half the jurisdictions impose no duty of their own, and no governance body anywhere in the sample is charged with holding the composite picture. If a coordinated post-incident response is ever needed, that gap will be discovered at the worst possible moment.

	Upstream risk assessment	Evaluation & verification	Red lines & monitoring	Incident reporting
European Union	H	H	H	H
France (domestic)	H	S	S	—
Germany (domestic)	—	H	H	—
Brazil	S	S	H	S
Republic of Korea	H	H	H	S
Australia	S	H	S	H
New South Wales (subnational)	—	H	H	H
United Kingdom	S	H	S	S
India	S	H	H	H
Singapore	H	H	H	H
Japan	H	H	H	H
Kenya	H	H	H	H
Taiwan	H	—	S	—
Canada	H	H	H	—
Israel	—	H	H	H
United Arab Emirates	H	H	H	—
Dubai (emirate) (subnational)	H	—	H	—
Dubai Intl. Financial Centre (subnational)	—	H	H	—
Peru	—	—	H	—
Switzerland	—	—	H	—
Chile	H >	H >	H >	—
Nigeria	H >	H >	—	—
South Africa	—	—	—	—

■ Hard law
■ Soft law
■ Non-binding
■ Strategy scaffold
■ Confirmed absence

Table 1: GPAI governance coverage by jurisdiction and governance area

Work with what has been built

Among AI middle-powers, the GPAI governance mechanisms that exist diverge in legal form, in trigger, and in who bears the duty. Harmonisation strategies that work with what each jurisdiction has already built, may have more to work with than waiting for governance templates to align.

About the authors



Josephine Schwab

Research Team Lead, Arcadia Impact

<https://jreschwab.carrd.co>



Nathan Naidoo

Research Associate

<https://www.linkedin.com/in/nathanrنايدو>



Ferruccio Barazzutti

Research Associate

<https://www.linkedin.com/in/ferrucciobarazzutti>



Sheryn Lee

Research Associate



Caio V. Machado

Expert Partner, The Future Society

<https://www.linkedin.com/in/caiocvm>

With additional thanks to the experts across all twenty jurisdictions who checked our findings.