

The UN Scientific Panel on AI's Preliminary Report Does Not Establish Its Independence

JASON TUCKER, VIRGINIA DIGNUM, RACHELE CARLI, PETTER ERICSON, TATJANA TITAREVA / JUL 8, 2026

The authors are researchers at the [AI Policy Lab](#) at Umeå University.



The United Nations headquarters in New York. [Shutterstock](#)

Share



Over the coming weeks, there will be no shortage of analysis of the contents and priorities put forward by the [UN's Independent International Scientific Panel on Artificial Intelligence \(IISPAI\) Preliminary Report](#). Yet, the report leaves unresolved a prior and more fundamental question: the nature of its independence, or more accurately, how that independence is understood, secured and demonstrated.

Independence is what gives the Panel its value. Without it, it is hard to distinguish the Panel from the many interest groups already vying for space in AI governance.

This independence here has two faces, and neither the report nor the documents related to the organization of the Panel address either one. The first is independence from industry. The second, and the more classic for a body of this kind, is independence from states. The IISPAI was established by a General Assembly resolution, its members nominated through a Secretary-General shortlist and appointed by Member States, and its reports are timed to feed the Global Dialogue on AI Governance on a political calendar the Panel does not control. A comparison with the Intergovernmental Panel on Climate Change (IPCC) is instructive: its sharpest independence problem is not funding but its Summary for Policymakers, approved by governments line by line. Whether IISPAI outputs face comparable state filtering, or whether the Panel even sets its own agenda, is unclear. For a body whose value rests on standing apart from the interests it assesses, these are first-order questions, not procedural detail.

As a first report, one would expect the mechanisms that secure independence to be set out clearly in the document: how content is determined when no consensus exists, and how donor funding and individual members' interests are disclosed and managed. Members serve in a personal capacity, and the funders are named: the governments of Germany, Japan and Spain, and the Omidyar Network Fund. But naming donors is not the same as managing their influence. There are no amounts, no split between financial and in-kind support, no conditions, and no declarations of individual members' interests. To be clear, this is not evidence of capture, but it is precisely why individual, itemized transparency is needed, rather than acknowledgement alone. The implicit request is that the public should trust the IISPAI as scientists. The experience of climate science, where trust and credibility had to be earned through transparency rather than assumed, shows why that is not enough.

The panel reproduces a weakness it identifies elsewhere: a lack of consensus and transparency. Its mandate is explicitly to document scientific consensus and disagreements, yet the report presents a smoothed broad consensus noting only that "no member endorses every point". Disagreement is to be expected in a body of forty experts across disciplines and regions, not a defect to be managed away. The value of an independent scientific panel lies precisely in making contested understandings visible,

engaging a broad range of stakeholders on how to navigate them. These tensions are what the panel should foreground.

The same tension runs through the evidence itself. The report rightly warns that many safety assessments are conducted by the developers building the systems, yet its own capability claims rest heavily on developer-produced and developer-adjacent evidence, from benchmark scores to forecasts such as the estimated doubling time of the tasks agentic systems can complete. A panel cannot flag developer self-assessment as a structural weakness and then rely on that same evidence base without stating how it corrects for it. Independence is not only about who sits on the Panel and who funds it, but of whose evidence is admitted and how its provenance is weighed.

This also matters for how priorities are set. It is unclear why the report foregrounds frontier AI and catastrophic risk, while the distributed harms and structural inequalities of existing AI systems receive much thinner treatment. These emphases mirror the research agendas of some of the Panel's most prominent members. The concern is not the integrity of these members, which we do not doubt, but method: without a transparent account of how topics were selected and ranked, readers can not tell whether prominence reflects a considered judgement of salience or the composition of the Panel. Where a report shapes a global agenda, the ranking of priorities should itself be explained and open to scrutiny.

It is also unclear why the report focuses on the inadequacy of largely untested governance frameworks, rather than building on instruments that are already in operation, such as the [EU AI Act](#), the [UNESCO Recommendation on the Ethics of Artificial Intelligence](#), the [Council of Europe Framework Convention on Artificial Intelligence](#), and [China's binding rules on algorithms and generative AI](#). By casting existing governance mechanisms as unfit for frontier models, the report obscures the value of existing regulation and risks pushing states towards reckless, discretionary interventions, as in June 2026 when a US Department of Commerce [export-control directive](#) compelled Anthropic to disable its Fable 5 and Mythos 5 models worldwide, an 18-day suspension imposed and then [reversed](#) through executive discretion rather than any statutory process.

Some of these tensions surface when reading the report. AI is presented as a solution to the most pressing societal challenges, yet the report also argues that significant societal transformation is required in order to enable these solutions. Essentially, we are expected to work for AI rather than AI for

us. The point is sharpened by the acknowledgement that current AI models are trained and optimized for only a very small number of the world's approximately 7,000 languages, leaving much of the world's population expected to adapt to systems that do not serve their own languages or cultural contexts.

The report has real strengths. It recognizes the structural concentration of AI power and insists that AI is neither inherently good nor bad, its outcomes depending on human and institutional choices. This recognition is what makes the independence issue awkward: a body that warns about the concentration of AI power in a handful of advanced economies and firms, and about the exclusion of the Global South, is underwritten by three high-income governments and a private fund rooted in technology-sector wealth.

Declarations of interest are standard in other international scientific committees. Measures adopted range from full transparency with searchable databases on individual members' disclosure forms (European Medicine Agency), to disclosure of members' interests ahead of meetings and reports (World Health Organization) to confidential individual disclosure reviewed by a standing committee (Intergovernmental Panel on Climate Change). Given the report's own acknowledgement of the risks of concentrated power, the individually itemized end of this spectrum is the most prudent and the most in line with the IISPAI's mandate.

In sum, we propose that the Panel adopt at least the following procedures and disclosures. These recommendations are addressed to different actors. The first three lie within the Panel's own remit; the last two rest with the Member States and the UN Secretariat that established and support it, and are no less essential.

1. Publish a clear methodology for reaching conclusions. Set out how evidence is assessed, how priorities are selected, and how decisions are made when members disagree. This should explain how minority positions are handled and what threshold constitutes consensus.

2. Increase transparency of the deliberative process. Document and publish the process by which the report is developed, including meeting summaries, records of deliberation, evidence submissions, and responses to expert feedback.

3. Make disagreement visible rather than smoothing it over. Where consensus does not exist, present competing scientific perspectives and the basis for each, rather than presenting contested issues as settled. Use calibrated confidence and uncertainty language so that the strength of evidence behind each position is legible.

4. Strengthen the management of interests and funding, not just their disclosure. Publish comprehensive declarations of interests for all panel members, disclose all funding sources and donor relationships, including amounts and conditions. Disclosure alone is insufficient: establish the architecture that makes it meaningful, a standing conflict-of-interest committee with recusal rules and attention to balance in author teams. Individual-level transparency is necessary to avoid accusations of undue industry influence.

5. Insulate the Panel structurally, not only through disclosure. Fund the Panel through ring-fenced or pooled contributions rather than earmarked national donations, so that no single government's support maps onto specific outputs, and establish a clear firewall between the Panel's scientific work and its secretariat, which currently sits within the UN Office for Digital and Emerging Technologies, a body with its own policy mandate. Transparency reveals influence; structural separation is what limits it.

Disclosure: One of the authors (Virginia Dignum) was a member of the United Nations High-Level Advisory Body on Artificial Intelligence (2023–2024), whose report “[Governing AI for Humanity](#)” recommended the creation of the body assessed in this piece. The views expressed are the authors' own. Work at the [AI Policy Lab](#) is supported by multiple grants from the European Union and the Wallenberg Foundation across diverse projects.

Editor’s note: Tech Policy Press has received grant funding from the Omidyar Network.

Authors



JASON TUCKER

Dr. Jason Tucker is a Researcher at the Institute for Futures Studies, Stockholm, and an Adjunct Associate Professor at the AI Policy Lab, Department of Computing Science, Umeå University. He is a researcher and policy expert working at the intersection of AI, healthcare, and geopolitics, with a bro...



VIRGINIA DIGNUM

Dr. Virginia Dignum is a Full Professor of Responsible Artificial Intelligence at Umeå University, where she leads the AI Policy Lab. A leading voice in global AI policy, she advises governments and organizations worldwide, including UNESCO, OECD, the EU, and the UN, and co-chairs major initiatives ...



RACHELE CARLI

Dr. Rachele Carli is a Postdoctoral Researcher in the AI Policy Lab at Umeå University, who studies the social, legal, and ethical impact of highly interactive AI technologies on individuals and the broader social fabric, with a particular focus on re-conceptualizing human vulnerability and analyzin...



PETTER ERICSON

Dr. Petter Ericson is a Staff Scientist at the AI Policy Lab, Department of Computing Science, Umeå University. A computer scientist by training, he is working in the intersection of AI, Policy, and AI Policy, studying the social effects of different AI technologies and implementations, and how thes...



TATJANA TITAREVA

Dr. Tatjana Titareva is a Staff Scientist at the AI Policy Lab, Umeå University. She researches AI in education, policy, leadership, and governance, with a broader focus on bridging research and AI policy to understand how AI technologies and their implementation reshape people, institutions...