

—Addressed to: President of the European Commission, Ursula von der Leyen and Executive Vice-President of the European Commission, Henna Virkkunen—

EU enforcement powers to ensure frontier AI models are safe and secure

As August 2nd 2026 approaches, we—AI researchers and representatives from civil society, academia and independent experts—express our support to the Commission as it prepares to **enforce the AI Act’s rules for general-purpose AI models with systemic risk** (GPAISR). The Commission oversees a world-leading framework that ensures frontier AI models deployed in Europe are safe, secure and fit for a trustworthy AI market. With systemic risks beginning to concretely materialise¹, and jurisdictions around the world (including the US) following the EU’s approach², we call on EU leaders to use its state-of-the-art powers.

One year on from the finalisation of the Code of Practice – signed by all GPAI providers advancing the frontier³, and therefore acting as a *de facto* industry standard – rapid developments in the systemic risk landscape⁴ have made the urgency of enforcement stronger than ever. Widely reported **cyber-offence capabilities from Anthropic’s Mythos model**⁵, as well as critical thresholds fast approaching for **biology**⁶ and **loss of control**⁷, reflect the specified systemic risks set out by the Code of Practice⁸. Given that providers have already had a year to directly engage with the AI Office in preparation for enforcement, we call on the AI Office’s Safety Unit A3 to **concentrate its resources on these specified systemic risks**, avoiding dual-accountabilities with existing Union law, including the DSA.

The signatories note that important sovereignty-related parallel steps are being taken under the European Technological Sovereignty Package⁹, especially in the wake of Europe’s initial exclusion from Project Glasswing¹⁰ and the US government’s export controls restricting all non-US citizen access to Anthropic’s Fable model¹¹. The EU’s resolve to enforce its own standards for safety and security is a crucial and necessary complement, **for the meaningful protection of European citizens** and simultaneously **progression of EU sovereignty**.

Consistent enforcement is also essential for **legal certainty and EU competitiveness**. Whilst increased urgency in the US around frontier AI risks is welcomed, the uncertain policy

¹ [Anthropic’s Mythos found thousands of zero-day high-severity vulnerabilities across every major operating systems and browser](#)

² [Emerging frontier AI governance frameworks in the US](#)

³ [Signatories of the Code of Practice](#) (including leading providers Anthropic, OpenAI and Google)

⁴ [International AI Safety Report 2026](#)

⁵ [NY Times reporting on cybersecurity implications of Anthropic’s Mythos](#)

⁶ [European Commission accredited external assessor SecureBio’s evaluation of GPT-5.5](#)

⁷ [European Commission accredited external assessor METR’s evaluation of GPT-5.6](#)

⁸ [Code of Practice](#) (Appendix 1.4)

⁹ [EU Tech Sovereignty Package](#)

¹⁰ [Politico reporting on the EU being sidelined by Anthropic’s first stage release of Mythos](#)

¹¹ [Politico reporting on US imposition of export control on Anthropic’s Fable](#)

response has led to legal uncertainty for downstream industry¹²; a consequence easily avoided by the EU's clear, predictable and forward-looking rules. With **75% of EU AI investment in the downstream application layer**¹³, and **80% of Europeans supporting careful regulation of AI**¹⁴, both EU citizens and businesses stand to benefit from enforcement.

Independent resources^{15,16} already point out gaps between the Code of Practice and current practices of GPAISR providers.

We therefore ask the Commission to:

1. **Make full use of the enforcement powers the AI Act provides**, drawing on them actively, as soon as concerns arise:
 - **Article 91: the power to request documentation and information** - This keeps Europe's understanding of frontier models, and their risks, in European hands.
 - **Article 92: the power to conduct evaluations** - This provides more precise information when necessary, whilst building the EU's evaluation capacity, underpinning sovereignty and a trustworthy market.
 - **Article 93: the power to request measures** - The readiness to use these measures gives the GPAISR rules credibility, signalling that steps of escalation will be taken when compliance obligations are not met and EU citizens and businesses are put at risk.
 - **Article 101: the power to impose fines** - The credibility of the entire regime depends on penalties being used in appropriate circumstances, deterring future breaches and reinforcing incentives to comply.
2. **Empower the AI Office's Network of Evaluators**¹⁷ with the time, access and resources, to appropriately carry out external assessments under the Code of Practice.
3. **Provide the AI Office with appropriate resources and political backing** for effective enforcement, echoing numerous calls of the European Parliament¹⁸ and experts^{19,20}.

¹² [US Legal tech firm sues US Government over Executive Order limiting foreign access to Anthropic's Fable models](#)

¹³ [Prosus State of AI in Europe Report 2026](#)

¹⁴ [2026 State of the Digital Decade Report](#)

¹⁵ [Future of Life Institute's AI Safety Index](#)

¹⁶ [SaferAI's Risk Management Ratings](#)

¹⁷ [AI Office Network of Evaluators](#)

¹⁸ [European AI Act Working Group Co-chair MEPs call for appropriate AI Office resourcing](#)

¹⁹ [Statement from the Chairs and Vice-Chairs of the GPAI CoP chapter on Safety & Security](#)

²⁰ [Open joint letter on AI Office resourcing](#)

On August 2nd the Commission will hold the tools and mandate to invest in Europe's safety, security, sovereignty and competitiveness by enforcing its world-leading frontier AI rulebook. The signatories of this letter call for these powers to be used with confidence and resolve.

Sincerely,



- **Michael McNamara** – Member of the European Parliament
- **Reinier Van Lanschot** – Member of the European Parliament
- **Halldóra Mogensen** – Former Member of Parliament, Parliament of Iceland
- **Yoshua Bengio** – Full Professor of Computer Science; Scientific Director, Mila Université de Montréal
- **Stuart Russell** – Distinguished Professor of Computer Science, University of California, Berkeley
- **José Hernández-Orallo** – Professor, Leverhulme Centre for the Future, University of Cambridge
- **Lorenzo Pacchiardi** – Assistant Research Professor, Leverhulme Centre for the Future, of Intelligence, University of Cambridge
- **Vincent Corruble** – Professor, Sorbonne Université
- **Raja Chatila** – Professor, Sorbonne Université
- **Daniel Mügge** – Professor, University of Amsterdam
- **Sören Mindermann** – Scientific Lead, International AI Safety Report
- **Federico L.G. Faroldi** – Associate Professor of Ethics, Law and AI, University of Pavia
- **Kristinn Thórisson** – Professor, Co-Director, Center for Analysis of Intelligent Agents, Reykjavik University
- **Jakub Growiec** – Professor, SGH Warsaw School of Economics; Poland CEPR Research and Policy Network on AI

- **David Krueger** – Assistant Professor (Robust, Reasoning, and Responsible AI), University of Montreal
- **Tegan Maharaj** – Assistant Professor of Decision Science, HEC Montréal
- **Ivana Budinská** – Senior Researcher, Institute of Informatics, Slovak Academy of Sciences
- **Zhijing Jin** – Assistant Professor, University of Toronto
- **Florian Mai** – Postdoctoral Researcher, University of Bonn
- **Leonard Dung** – Postdoctoral Researcher, Ruhr-Universität Bochum
- **Nada Madkour** – Director, Center for Long-Term Cybersecurity
- **Deepika Raman** – Non-Resident Research Fellow, Center for Long-Term Cybersecurity
- **David Evan Harris** – Chancellor's Public Scholar, University of California, Berkeley
- **Anna Katariina Wisakanto** – Senior Researcher, Center for AI Risk Management & Alignment
- **Alexandra Abbas** – Research Engineer, Meridian Research Labs
- **Marcus Paleti** – Director, World Technology Congress