



Commission starts enforcing AI Act rules and new transparency requirements on 2 August

Brussels, 31 July 2026

From 2 August 2026, the European Commission's AI Office, together with national authorities, will begin enforcing [the Artificial Intelligence \(AI\) Act](#). On the same date, new transparency rules will start to apply, requiring certain AI systems to tell users when they are interacting with AI and when content has been generated or altered by it.

Under the new rules, chatbots and other interactive AI systems will have to tell users they are dealing with AI, not a human. Deepfakes (images, videos, or audio that have been edited or generated using AI) will have to be labelled. AI-generated or altered content will also have to carry machine-readable marks so it can be detected more easily.

The measures are intended to reduce deception and manipulation and help people make informed choices. They also give businesses clearer obligations and a practical way to show compliance. The Commission published today a [first list](#) of more than 180 organisations that have signed the [Code of Practice on transparency of AI-generated content](#) that operationalises the rules on transparency of AI-generated content.

As AI grows increasingly capable and integrated into everyday life, the AI Act helps ensure that AI is developed, deployed, and used safely, giving people and businesses across the EU greater confidence in the technology.

Enforcement of the AI Act

The [AI Office](#) can now enforce the AI Act's rules for providers of [general-purpose AI \(GPAI\) models](#). These models can perform many different tasks and can be used in a wide range of tools and services, including AI agents.

The rules also cover the most advanced GPAI models that may pose systemic risks. Their providers must meet additional obligations to address risks of large-scale harm, such as risks linked to chemical, biological, radiological and nuclear incidents, loss of control, cyber offence, harmful manipulation and threats to fundamental rights. They also address risks that have recently drawn public attention, including risks to European [cybersecurity](#) and to AI acting outside human control.

All providers of GPAI models must document certain information and provide it to competent authorities or downstream providers. They must also put in place a copyright policy and publish a sufficiently detailed summary of the content used to train their models.

Enforcement also begins for transparency obligations and [prohibited AI practices](#). These ban particularly harmful systems, including systems that manipulate people, exploit vulnerabilities in harmful ways, or unfairly score people in ways that threaten their rights.

Responsibility for enforcing the transparency rules and prohibited practices is shared across three bodies. The AI Office enforces the rules for AI systems offered by the same provider as the underlying general-purpose AI model. It also covers systems integrated into very large online platforms or very large online search engines designated under the Digital Services Act.

[National competent authorities](#) enforce the rules for other AI systems. The European Data Protection Supervisor enforces the rules for AI systems used by European Union institutions, bodies and agencies.

Effective enforcement will also depend on Member States ensuring that national competent authorities are properly designated and adequately resourced.

Scientific support

The AI Office and national competent authorities will be supported in their enforcement work by the [Scientific Panel](#), an expert advisory body made up of 60 independent AI experts. The panel recently held its first meeting.

The AI Office has also appointed Professor Alessandro Abate of the University of Oxford's Department of Computer Science as Lead Scientific Adviser. He will support the Office's scientific work on general-purpose AI models, including innovation and adoption, as well as model testing and evaluation.

Monitoring tools

To support enforcement, the AI Office has launched several tools for individuals and businesses. Natural and legal persons can use the [Complaint Tool](#) to report alleged infringements of the AI Act by providers of AI systems supervised by the AI Office. People working with providers of AI systems or general-purpose AI models can use the [Whistleblower Tool](#) to report possible violations of the Act securely. A [dedicated channel](#) is also available for downstream providers that build AI systems on general-purpose models and want to report alleged infringements by the providers of those models.

The AI Office will treat information received through these tools confidentially.

Background

Professor Abate joins the AI Office from the University of Oxford's Department of Computer Science. He previously conducted research at Stanford University and SRI International, a leading US research institute, and served as an Assistant Professor at TU Delft. He is internationally recognised for his work on the safety, verification and control of AI-enabled systems, areas that lie at the heart of the AI Office's mandate. Professor Abate holds an MS (2004) and a PhD (2007) in Electrical Engineering and Computer Sciences from the University of California, Berkeley.

Next steps

The [AI Omnibus](#) postponed the application of the rules on high-risk AI systems to 2 December 2027. It also postponed the rules for high-risk AI systems integrated into regulated products to 2 August 2028.

The AI Omnibus introduces new prohibitions on AI systems that generate non-consensual sexually explicit content and child sexual abuse material. These rules will apply from 2 December 2026.

For more information

[Enforcement under the AI Act](#)

[AI Act Service Desk](#)

[Guidelines on prohibited AI practices](#)

[General-Purpose AI Code of Practice](#)

[Template for the Public Summary of Training Content for GPAI models](#)

[Guidelines for providers of GPAI models](#)

[Code of Practice on Transparency of AI-Generated Content](#)

[Draft guidelines on the transparency obligations for certain AI systems](#)

[Draft guidelines on the classification of high-risk AI systems](#)

IP/26/1714

Quote(s):

--

"AI is a transformative technology that can bring extraordinary benefits to our people and businesses. But we are also seeing that harms can occur if AI is not properly designed and used and the most advanced models create risks on an entirely new scale. Europe anticipated this development. With the AI Act, we established a clear, risk-based and durable framework for trustworthy AI —one that gives innovators legal certainty while protecting the public interest. As enforcement begins, we are taking an important step towards AI that people and businesses can understand and trust, and whose benefits are shared widely across our society."

Henna Virkkunen, Executive Vice-President for Tech Sovereignty, Security and Democracy

Press contacts:

[Thomas REGNIER](#) (+32 2 29 91099)

[Nika BLAZEVIC](#) (+32 2 29 92717)

General public inquiries: [Europe Direct](#) by phone [00 800 67 89 10 11](#) or by [email](#)