

The EU AI Act: Challenges and questions for AI Providers as the Act starts to bite

The EU AI Act (the **Act**) has been taking effect in stages since 2025, with further obligations scheduled to take effect right through to 2030. Against that backdrop, 2 August 2026 marked an important staging post, bringing into force the Article 50 transparency obligations and activating the Act's enforcement framework.

While AI providers have already had to grapple with the rules on prohibited AI practices, AI literacy and general-purpose AI (**GPAI**) models, the application of the broader compliance and enforcement framework, together with the Act's transparency rules, marks a major milestone for the EU's AI regulatory regime.

At the same time, the rules continue to evolve. The AI Omnibus (the EU's "simplification" package – see our [blog](#)), multiple consultations from the European AI Office, and new guidelines and codes of practice have all continued to shape the compliance landscape in recent months. Despite this, many of the issues which are most important to the providers of AI systems and GPAI models remain unclear.

In this briefing, we highlight five key challenges facing AI providers today and consider how they can manage the legal and regulatory risks around those challenges. Taking a step back, we also consider whether the Act is likely to have a "Brussels effect" by influencing the emergence of AI regulation, or the behaviour of AI providers, outside Europe.

Five challenges for AI providers

Challenge 1: EU law, global risk?

As with many EU Regulations, the Act has broad extra-territorial effect. As well as covering deployers (users) based in the EU it covers:

1. providers (wherever based) who place AI systems or GPAI models on the market, or put them into service, in the EU - these concepts are reasonably well-established in EU law; and
2. non-EU providers (and deployers) of AI systems "*where the output produced by the AI system is used in the [EU]*".

This second limb is difficult. The recitals¹ confirm that it was included, at least in part, to stop organisations circumventing the Act by sending EU data to be processed by AI systems outside the EU only to receive the AI's output back in the EU. But the provision has much broader potential application, because there are many ways in which the output of an AI system may end up in the EU. For example, a single piece of AI-generated content, produced in the US, India or China, but shared on social media in the EU, could bring the original provider in scope of the Act.

So far, the most helpful guidance comes from the final version of the Commission's guidelines on transparency (the **Transparency Guidelines**) (see [blog](#)). Here, the Commission notes that "... *incidental, unforeseeable or unauthorised downstream use should not alone trigger the application of the obligations...*". But it's not clear what a provider needs to do to benefit from this clarification, particularly as other AI Act guidance suggests that relying on restrictions in usage terms may be insufficient.² What does a provider need to do to show that downstream use was unauthorised (beyond referring to its usage terms), or unforeseeable?

¹ See Recital 22.

² See the High-Risk Guidelines, para 12, which discuss the limitations of simply relying on terms of use.

It's easy to see how this might concern AI system providers outside the EU. In circumstances where a given AI output may be impossible to control once exported from the platform and where the boundary between intended EU use and incidental downstream use remains uncertain, the AI Act could represent an open-ended risk for ex-EU providers (including smaller providers without the resources to apply to regulatory compliance). The extra-territorial burden is compounded by the Act's requirement that non-EU providers appoint an EU-based authorised representative in certain circumstances, whose obligations go beyond simply acting as a mailbox (adding additional internal governance, resourcing and liability considerations for providers).

What can we expect? Possibly a combination of strict usage terms, persistent user-facing language on how outputs can be used and distributed – and, where available, technical measures limiting dissemination into the EU. As we've seen in other areas of EU regulation, this could ultimately prevent or delay some AI providers from releasing features (or possibly whole products) to European customers, particularly as some compliance solutions will need to be implemented at the model level (and geography-specific model serving is often prohibitively compute-intensive). And – where tolerable for providers – we could see model or system-level changes implemented globally to comply with EU rules (as we've already seen with Anthropic – see further below). As much as this would delight Brussels, it will be a difficult decision for providers and may not win support among users.

All told, we think more guidance will be needed here to illuminate the "safe harbours" that providers can rely on.

Challenge 2: AI System - yes or no?

There is still a fundamental lack of clarity around many of the key definitions in the Act, which presents real challenges when it comes to applying the Act in practice.

This is a particular issue with the definition of "AI system", given how central it is to most of the Act's obligations. The European Commission did publish [guidelines](#) on this definition in July 2025 (the **AI System Definition Guidelines**), although they offer little beyond what is already in the Act and its Recitals and have failed to resolve several key issues that we see many providers struggling with.

What is AI? AI v traditional software

The first of these issues is: what is "AI" as opposed to "traditional software"? This largely comes down to whether a system has the necessary degree of inference to qualify as "AI". In many cases this will be obvious – for example, any system making use of a Large Language Model will almost certainly have the necessary degree of inference to be caught by the Act. However, providers of more traditional data analytics or decision support systems (and their lawyers) quickly find themselves in a grey area. For example, how do you draw a sensible line between simple multi-dimensional statistical optimisation techniques and "AI" techniques?

The AI System Definition Guidelines do touch on the inference requirement, offering some basic examples of what sits outside the definition of an AI system and therefore the scope of the Act. These include *"systems for improving mathematical optimisation"*, *"basic data processing"*, *"systems based on classical heuristics"* and *"simple prediction systems"*. However, at the more advanced end, it again becomes unclear.

Technical teams may be inclined to take a practical approach, assuming that it is generally obvious when a system constitutes AI, particularly where it is described as "AI" or "machine learning" in technical (rather than marketing) documentation. However, where significant regulatory obligations hinge on whether a system falls within the definition of an "AI system", providers and their legal advisers would benefit from greater certainty than a "know it when you see it" approach can provide.

What is an AI "system"?

The second issue we see is how to draw the line around any particular AI "system". At one extreme, every device and piece of software that connects to the internet could be part of a single system. At the other, every individual function within any piece of hardware or software could be an individual system.

The Act must intend something in-between these two extremes, but the lack of good guidance creates uncertainty. The issue is not addressed in the AI System Definition Guidelines, which are more focused on the meaning of AI (versus other types of predictive technology). Providers must instead look to the Commission's draft guidelines on high-risk AI systems (the **High-Risk Guidelines**), which note that: *"...[w]here several AI systems form part of a more complex AI system... the combined configuration is treated as a single AI system for the purpose of high-risk classification..."*³ But that neither

³ See the [High-Risk Guidelines](#), para 75-76.

answers the question of where a single system begins and ends, nor clarifies if this analysis will apply in all circumstances (as it was discussed in the context of whether an exemption to the rules can apply). The guidance goes on to confirm that: "[to] avoid circumvention of the high-risk classification rules by system design, split architectures are assessed as a whole."

This is a particular issue for internal teams grappling with where the compliance boundary lies. Can they legitimately draw a narrow boundary even where an obviously AI-based system plugs into other non-AI systems in a larger architecture? Or can they legitimately draw a wide boundary where that architecture comprises AI and non-AI elements but is used for a single purpose (or for a particular corporate function)? The compliance impact may be significant between these decisions, because – particularly for high-risk AI Systems – the burden of the obligations will often correspond to the size and scope of the AI System (particularly the requirements around risk management systems, technical documentation, record keeping and event logs, human oversight and transparency), with obvious resource and organisational implications. Without more guidance, providers are having to reason from first principles (and with little sense of what stance the regulators will take).

Challenge 3: Just how much transparency?

Under the Act, system providers have two transparency obligations – users must know: (i) they are interacting with AI; and (ii) when content is generated or manipulated by AI.

While there are some questions around when, how and how often users should be told about their AI interactions, it is the second, marking, obligation which is proving more challenging (and controversial). This requires providers to add machine-readable marks which allow AI-generated or manipulated content to be identified, subject to limited exceptions⁴ (the so-called watermarking rules - Article 50(2)). In most formats, the related Code of Practice on transparency (the **Transparency Code**) has concluded that, in practice and given the current state of the art, this requires both digitally signed metadata and imperceptible watermarking. However, there are limitations as to what the technical solutions can currently do to satisfy these requirements, particularly where plain text is involved (see further below).

It is a challenge impinging on both AI system providers and on GPAI model providers upstream. Despite the rules not directly applying to model providers (an odd quirk of Article 50(2)), model providers are nonetheless facing demands to support such marking tools from downstream system providers who have embedded a model into their AI system – and the Transparency Guidelines encourage model providers to implement appropriate transparency measures at the model level.⁵ Of course, many model providers are also system providers under the Act, and so are separately incentivised to facilitate compliance.

The practical difficulties are also raising questions among some AI system providers over whether they should sign the Transparency Code (the Codes of Practice are discussed further in Challenge 4 below).

Watermarking plain text

Marking plain text is proving particularly challenging. As plain text is composed of simple strings of text characters (as opposed to text containerised in files such as PDFs or Word documents which offer additional possibilities for watermarking, as well as embedding metadata or applying other markings), watermarking is proving intrinsically more difficult.

While some methods do exist for this, they have their limitations. For example, visible or invisible/zero width characters can be added to the text, but these can impact the user experience, make it difficult to use the text (e.g. impacting search or copying functions) and may not be very robust. Advanced statistical watermarking techniques such as Google's SynthID can also be used, but these generally work by subtly impacting the output tokens chosen by an AI model to embed a detectable watermark signal in generated text, which can affect the output (for example, by biasing in favour of certain token types). In other words, the watermarking changes the text itself.⁶ Anthropic has also **announced** its own technique for embedding watermarks into generated text (as well as attaching signed provenance metadata to files), but this also seems to affect token choice (although in a way Anthropic says is imperceptible to humans).

As these rules have broken into developer and engineering circles, they have proven incredibly controversial. While watermarking images and video has not attracted much concern (and providers were supporting those techniques before the AI Act), watermarking text has touched a nerve.

⁴ See Article 50(2). E.g. the rules do not apply to the extent the AI systems perform an assistive function for standard editing, and the guidelines include translation as another example.

⁵ See the **Transparency Guidelines**, para 26-27.

⁶ There are also concerns around their reliability and robustness (they can be unreliable for short text segments, and the watermark can be easily damaged or lost if the output text is subject to significant editing) and the potential risks around them producing false positives.

Anthropic's announcement, for example, generated significant debate online, attracting anger that the watermarking technology can affect the actual output (even if only marginally changing quality or utility). There are also concerns that even lightly edited or proofread text could be marked as an AI product (robbing the human author of true credit, or suggesting AI was used where it wasn't), although Anthropic have sought to dispel those. Similar concerns have been raised with over-inclusive AI tags on lightly edited photography, raising the question of whether the Act is driving behaviours which ultimately defeat its purpose of reliably identifying AI versus human-made content (for example, where providers take an inclusive approach to marking to ensure compliance).

The Commission may know it is creating more problems than solutions here. It recently outlined several exceptions to the text marking requirements which are not obviously based on the text of the Act itself (including for very short text of 200 tokens or fewer).⁷ Further lobbying may, however, be needed if it becomes clear that genuine compliance for plain text is hampered by technological restrictions, and that the solutions being developed to seek to resolve that gap are doing more harm than good. More generally, it raises the question of whether the Act is too ambitious to suppose that AI-generated text can be robustly marked and traced – and perhaps naïve in supposing that this would be without controversy among users.

Challenge 4: Navigating the GPAI rules

The GPAI rules have been controversial from their inception. A late addition to the Act, they were introduced following criticisms that the (then draft) Act failed to cater for GPT-style generative AI models – and they nearly derailed the whole process over concerns they could harm the development of EU models like Mistral.

They raise a variety of issues, from the way in which systemic risk is assessed to concerns around creating extraterritorial reach for copyright law and requirements to disclose commercially sensitive material. A number of these were debated at length as part of the process to agree the Code of Practice on GPAI models (the **GPAI Code**) (see [blog](#)).

Assessing "systemic risk"

One goal of the AI Act is to closely regulate GPAI models that pose "systemic risks", including any actual or reasonably foreseeable negative effects in relation to major accidents, disruptions of critical sectors and serious consequences to public health and safety; on democratic processes, public and economic security; and the dissemination of illegal, false, or discriminatory content.⁸

"... it is quite conceivable that a non-GPAI model, or a GPAI model falling below the current systemic risk thresholds, could be launched in Europe with significant disruptive capabilities – and be completely outside the systemic risk rules under the Act. That would be an odd outcome for Europe's flagship AI safety regulation."

Here the Act's approach is curious – and potentially misguided.

First, only GPAI models can be designated as posing "systemic risks". This follows the logic that the highest impact, most broadly capable models are most likely to pose systemic risks – but there's no reason why a non-GPAI model couldn't pose those risks (and, in fact, more focused models could prove highly risky, e.g. models that identify security vulnerabilities in critical software, or models used for biological research).

Some of those riskier non-GPAI models could be caught by the list of prohibited AI practices,⁹ but that list is quite narrow, and the Commission would need to bring new primary legislation to expand it.¹⁰ And that would not, of course, serve the same purpose as the GPAI systemic risk

⁷ See the Transparency Code, Sub-measure 1.1.2 and Glossary definition of "very short text"; and the Transparency Guidelines, paras 64-68 and 76.

⁸ See Recital 110.

⁹ See Article 5.

¹⁰ For example, so-called "nudification" apps were added during the Omnibus negotiations as a new prohibited AI practice.

rules (of applying higher ongoing safety requirements on riskier models). It would simply ban them outright.

Second, the current guidance for assessing which GPAI models are presumed to have "systemic risks" under Article 51(1)(a) (the **GPAI Guidelines**) focuses on cumulative training compute (and whether GPAI models pass a defined FLOP threshold).¹¹ This is consistent with the focus on GPAI models (and the general idea that "big means risky") – but is an approach that could prove both too wide and too narrow. On the one hand, ever more powerful models could mean the Commission is inundated with models that surpass the FLOP thresholds, making enforcement difficult (given the system was designed to only catch the most powerful models). On the other hand, models falling below the FLOP thresholds could be missed. One example could be so-called "student" models that can assimilate many of the capabilities of "teacher" models via distillation, but do not use the same levels of compute (and the GPAI Guidelines say that compute used to train the teacher model shouldn't be counted in assessing the cumulative training compute of the student model).

The Commission itself acknowledges this approach is "*an imperfect proxy for generalities and capabilities*",¹² and the Act does contain additional criteria. For example, a GPAI model will be classified as having systemic risk if it has "*high impact capabilities*", and these can be evaluated based on certain indicators and benchmarks. The Commission can also designate a GPAI model as having systemic risk on alternative grounds,¹³ and adopt delegated acts to introduce different benchmarks and amend the FLOP thresholds. The Commission may need to explore these options with some urgency if the system is to be robust. And as things stand, it is quite conceivable that a non-GPAI model, or a GPAI model falling below the current systemic risk thresholds, could be launched in Europe with significant disruptive capabilities – and be completely outside the systemic risk rules under the Act. That would be an odd outcome for Europe's flagship AI safety regulation.

This therefore seems an area that calls for much more agility from the Commission, and a flexible understanding of which market participants pose genuine risk. The current approach, which has hallmarks of the regulatory

hierarchies set out in the Digital Markets Act and Digital Services Act,¹⁴ may prove misguided in a landscape that is full of change and dynamic risks.

Should I sign the codes?

The Act¹⁵ envisages that approved codes of practice will help with compliance – and two have been agreed and published: the GPAI Code and the Transparency Code (together, the **Codes of Practice**).

These create a difficult question for providers, for two reasons:

- first, in many areas they appear to go further than the Act's requirements (for example, the Transparency Code contains substantially more "rules" around copyright issues than the simple requirement in the Act to implement a compliant policy); and
- second, while the Commission says that providers may "*...rely on the code to demonstrate compliance...*" both the GPAI and Transparency Codes state that "*...adherence to the Code does not constitute conclusive evidence of compliance with these obligations*".

The FAQs do go on to confirm that signatories will benefit from "*increased trust*" from the Commission and others, but it's unclear what this will mean in practice – and time will show if the Commission would ultimately refrain from enforcement action against a Code-compliant GPAI provider in the event of a major issue. At the same time, we see providers coming under pressure from the AI Office (and from the fact competitors have signed) to sign up to the Codes of Practice without a clear view on what this will mean for them in terms of their compliance (including how the AI Office and Member State authorities will enforce the underlying obligations in practice) – and without significant confidence that code compliance will help them defend against future enforcement.

A rational end state for the Codes of Practice is that they become an expanded framework for discussions about compliance. In that context, signatories may benefit from a higher baseline of goodwill (and some "*ready-made*" factual defences). That said, non-signatories may have good arguments that their approaches are still compliant with the Act's requirements. That leaves non-signatories

¹¹ The GPAI Guidelines describe FLOPs (or floating-point operations) as a measure used to quantify the training compute (which is the number of computational resources used to train the model). See our blogs [here](#) and [here](#) on the evolution of the guidelines.

¹² GPAI Guidelines, para 16.

¹³ See Article 51(1)(b).

¹⁴ Which, respectively, identify "gatekeepers" considered to have entrenched market power, and "Very Large Online Platforms" and "Very Large Online Search Engines" considered to pose intrinsically higher risks of online harms.

¹⁵ See Article 56.

without those good arguments clearly in a position of higher risk.

Challenge 5: Not ready for high-risk

Since the inception of the Act, a key focus for many providers has been whether their systems qualify as "high-risk", given the compliance burden under the Act is significantly greater for these systems.

The rules create several challenges for providers (and potential providers) and there are big open questions even as these obligations finally take effect after the Omnibus negotiations. For example:

Guidance needed

The Act envisages that providers will have standards and guidance to help them comply with the high-risk rules – and one of the reasons the delay to the high-risk rules was first floated (in the context of the Omnibus) was because those resources didn't yet exist.

Today we are in a similar place. A first draft of the High-Risk Guidelines¹⁶ has been published (see [blog](#)) which looks at what is and is not in scope. They do provide some assistance – clarifying, for example, that the headings in each of the points in Annex III form part of the test as to whether a system is high-risk (so, for example, credit scoring is only in scope if it is for essential private services and essential public services and benefits), and providing some useful practical advice on the application of the "filter" exception under Article 6(3). But these guidelines arrived (in draft) only weeks before the rules would originally have come into force – and the technical standards (together with further Commission guidance) intended to help providers comply with the obligations themselves are still nowhere to be seen.

A key request from providers is therefore that the necessary standards and guidance are available in good time ahead of the new high-risk deadlines – particularly as other key pieces of guidance have been published late or extremely close to compliance deadlines.¹⁷ The Omnibus negotiations ultimately dispensed with the proposal of tying the high-risk application date to the availability of the necessary standards and guidance – but the worst case would be approaching the new (fixed) deadline still without good guidance.

The accidental provider

A user/deployer of an AI system effectively takes over the provider role from the original provider if they: (a) put their name or trademark on a high-risk AI system (without otherwise contractually allocating provider responsibilities); (b) substantially modify a high-risk AI system; or (c) modify its purpose such that it becomes high-risk.¹⁸ Where this happens, the original provider must cooperate with and assist the new provider.

"... there is a degree of legal fiction in assuming that a user can step into the shoes of a provider when many of the provider obligations presuppose knowledge or control of systems to which a deployer may have little or no access."

With little guidance to help, there are many outstanding questions around how this will work in practice. For example, what will the contractual allocation of responsibilities look like? What does an organisation need to do to modify a system's purpose (is simply using it for a new high-risk use case sufficient)? And what level of cooperation and assistance is required from the original provider?

Whilst the regime is a creative tool for plugging a gap in the value chain, there is a degree of legal fiction in assuming that a user can step into the shoes of a provider when many of the provider obligations presuppose knowledge or control of systems to which a deployer may have little or no access.

Overlapping regimes

There has been much discussion as part of the AI Omnibus process about how the AI Act works with the Annex I product safety legislation, resulting in changes to the conformity assessment rules, the removal of the Machinery Regulations from Annex I Section A (meaning a more limited set of rules applies under the Act)¹⁹ and

¹⁶ See Article 6(5).

¹⁷ E.g. the GPAI Code and GPAI Guidelines, and the draft High-Risk Guidelines themselves.

¹⁸ See Article 25.

¹⁹ It is now in Annex I Section B and the Machinery Regulation is being amended through delegated acts to incorporate some AI Act requirements

questions around whether other legislation (particularly Medical Device rules) should follow suit.

However, some of the Annex III high-risk categories raise similar questions around how the Act's rules work with general legislative regimes. For example, AI used in an employment context (including recruitment, where AI is increasingly used) can be an Annex III high-risk use case – but is also an issue very much on the radar of many EU privacy legislators. Likewise, the Annex III high-risk categories pertaining to insurance and credit-scoring provide further examples of use cases that sit within mature and heavily supervised regulatory frameworks. These overlaps raise a broader question: whether the incremental protections delivered by the AI Act justify the additional compliance burden in areas that are already subject to significant regulatory oversight.

Bad branding?

Already we see a distinction between the Commission's view of "high-risk" and that of private organisations.

The Commission appears to see the "high-risk" rules as a mechanism for providers to demonstrate their commitment to safety and increase market confidence in their AI systems – so increasing uptake. But many private organisations struggle with the idea that their product being described as "high-risk" will help their go-to-market (particularly where they have invested significantly in the security and reliability of a product independently of the AI Act requirements).

Over time, we may see the idea of "high-risk" converted into positive technical standards (when those are eventually developed) such that it becomes a benchmark of quality and trustworthiness, but today the idea is causing more alarm than enthusiasm among organisations who fall in scope.

Will the AI Act have a "Brussels effect"?

There are, as can be seen, a range of practical compliance challenges for providers of AI models and systems caught by the Act. A compounding issue is that many will be both providers and deployers in practice, meaning they are also caught by the deployer rules (which raise their own issues) – and some will also have GPAI model provider obligations thrown into the mix. Add to that a complex enforcement environment, with the AI Office notionally supervising the

larger providers (but subject to untested exceptions), and multiple market surveillance authorities in many Member States, and providers are facing a complex landscape despite the Act being a directly applicable Regulation.

Practical issues aside, a bigger question is whether the Act is likely to have a "Brussels effect" and influence the development of AI regulation outside Europe (similar to the GDPR and, to a lesser extent, the DMA (Digital Markets Act)). Do we see AI Act-inspired regulation emerging elsewhere, or evidence that providers are adjusting their activities outside Europe to mirror their AI Act compliance?

Today, the answer is mixed. There have been signs of similar rules emerging in some jurisdictions. For example, in South Korea and Brazil there are new laws (enacted or proposed) which share commonalities with the Act, particularly in the establishment of a tiered risk approach for AI use cases, and transparency obligations.²⁰ At the US state level (but not at the federal level), laws have been introduced (or proposed) around comparable issues, in particular the transparency of AI content, and in New York, rules on using AI in recruitment (one of the AI Act's high-risk activities). Looking at providers, perhaps most significant has been Anthropic's decision to implement its text watermarking at the model level globally, which it explicitly linked to AI Act compliance (although this may be more about managing compute restraints, as geographic model serving can be extremely compute-intensive).

On the other hand, many jurisdictions are taking a lighter touch approach, particularly the US (at federal level), the UK and China (unsurprisingly, as all these countries have ambitions for AI leadership). More broadly, much of the global debate around AI regulation has moved to discussing the cyber risks posed by frontier models and agentic AI systems, and how to regulate Artificial General Intelligence (**AGI**) if and when it arrives – neither of which are a direct focus of the AI Act.

Can the Act have a voice here? It's unclear. We are already seeing different models emerging, including calls for more independent pre-release testing of frontier models.²¹ The US, for example, has recently established a voluntary framework under which providers give the Federal Government pre-release access to "covered frontier models", and major industry voices have called for an international standards-setting body that conducts

²⁰ Namely the South Korea AI Basic Act (Framework Act on Artificial Intelligence Development and Establishment of a Foundation of Trust) (2024), and the Brazil AI Bill (PL 2338/2023).

²¹ It should be noted that the Act's high-risk rules do include conformity assessments requirements, albeit this assessment can often be undertaken internally.

rigorous (third-party) pre-release safety testing.²² The UK's AI Security Institute is established to perform a similar role.

These types of third-party model testing could have greater force than the AI Act's requirements (which generally put the onus on providers to carry out testing internally) – particularly in circumstances of intense rivalry where providers are incentivised to rapidly release models. This could be a key factor driving calls from providers for these types of rules.²³

In this context there is a risk that the AI Act ends up "downstream" of these types of frameworks. The ultimate test of the AI Act may not be whether providers can comply with it, but whether it maintains relevance as one of the principal frameworks through which the most capable AI systems are governed, or whether its requirements begin to look pedestrian as attentions shift towards frontier models, autonomous capabilities and AGI. Put another way, will the Act be an important part of a wider international framework, or subsidiary to rules which are ultimately set and imposed elsewhere?

Digital Regulation at Slaughter and May

Our Digital Regulation practice covers the full spectrum of digital rules, from AI regulation, competition-focused platform regulation, to online harms and content regulation, and data access and portability. Clients look to us to offer practical, risk-adjusted advice that helps them navigate this ever more complex landscape – and to innovate at speed.

For more information, see our [Digital Regulation Practice](#) page, our digital blog [The Lens](#) or contact one of the team (below) or your usual Slaughter and May contact.

²² See the [proposals of Demis Hassabis](#), who proposes a new Standards Body modelled on a federally overseen public-private partnership or self-regulatory organisation, comparable to the Financial Industry Regulatory Authority (FINRA), with a board that includes independent leading technical experts and open-source representatives.

²³ See, for example, [Dario Amodei's comments](#) that frontier models could be "required to go through technical testing and auditing, and their release should be blocked or reversed as a threat to public safety if they do not meet high standards of safety."

For more insights on the AI Act itself, see our article [The EU AI Act enters into force](#), our [AI Act Essentials flyer](#) and our [AI Blogs on The Lens](#).

Contact



Will Manley
Head of Digital Regulation
T: 0032493406939
E: Will.Manley@slaughterandmay.com



Laura Houston
Partner
T: (0)20 7090 4230
E: Laura.Houston@slaughterandmay.com



Natalie Donovan
Head of Knowledge, Tech & Digital
T: (0)20 7090 4058
E: Natalie.Donovan@slaughterandmay.com

We would also like to thank Ian Ranson, Lily Latimer Smith and Jack Eames for their help in writing this publication.

London

T +44 (0)20 7600 1200
F +44 (0)20 7090 5000

Brussels

T +32 (0)2 737 94 00
F +32 (0)2 737 94 01

Hong Kong

T +852 2521 0551
F +852 2845 2125

Beijing

T +86 10 5965 0600
F +86 10 5965 0650

Published to provide general information and not as legal advice. © Slaughter and May, 2026.
For further information, please speak to your usual Slaughter and May contact.

www.slaughterandmay.com