

# A Practical Guide to the EU AI Act

Building Governance Before the Deadline Forces It



# Executive Summary

The Digital Omnibus has postponed the EU AI Act's most demanding requirements for high-risk AI systems until December 2027 and August 2028. Many organizations will view this as a welcome reprieve. In reality, it presents organizations with an opportunity to move ahead of the market, build compliance maturity early, and gain ground while competitors remain on the sidelines.

GDPR offers a useful preview of how this tends to unfold, with early enforcement targeting the easiest, best-evidenced gaps, not complex judgment calls, and organizations that wait for a deadline to force the work typically arrive underprepared and competing for the same scarce expertise as everyone else who also waited.

This paper sets out what should change inside an organization because of the Act. How classification, human oversight, monitoring, documentation and role determination need to work in practice. It argues that the real challenge isn't any single requirement, but building these into a standing way of working rather than a one-off compliance project. Role reclassification is the clearest example: ordinary product and delivery decisions can quietly shift an organization from deployer to provider, picking up a materially heavier set of legal obligations, often without anyone realizing it happened.

Getting this right depends on bringing an organization's AI lifecycle, its technology and its governance into one place, backed by genuine regulatory expertise rather than just generic risk tooling. One of the more compelling examples we have worked with directly is ServiceNow's AI Control Tower. However, technology alone is not enough, the real value comes from combining the right platform and tools with regulatory expertise and governance processes designed around the Act's actual requirements.

That's the combination we've spent the past year building for clients and it's the basis for the argument in the pages that follow. The next 12 to 18 months is an opportunity to grasp a head start, not a distant deadline you don't have to consider yet.

“

The organizations best placed when the 2027/2028 deadlines land won't be the ones who waited for certainty, they'll be the ones who built experience while everyone else was waiting.

”

## Authors



### Oliver Vistisen

ServiceNow Architect, Synechron



## About Synechron:

Synechron is a global technology consulting firm that helps leading organizations accelerate digital transformation through innovation, expertise and agility. With more than 17,000 plus professionals across around 58 offices in 22 countries, we combine deep industry knowledge with advanced capabilities in such areas as AI, cloud, cybersecurity, and data engineering.

Our regional teams, supported by strategic delivery centers, provide scalable, cost-efficient solutions tailored to local markets. Through our award-winning SynechronLabs accelerators and strategic partnerships with AWS, Microsoft, Databricks, Salesforce, and ServiceNow, we enable clients to innovate fast and lead with confidence. For more information on the company, please visit our website or LinkedIn community.



## 1.

## Why Now, Not Later

The recent adoption of the Digital Omnibus formalized the previously suggested 12 to 18-month delay to some of the EU AI Act's most demanding obligations. Since its publication in November 2025, many organizations have adopted a wait-and-watch approach, hoping it would do exactly that, and most compliance teams will take the European Parliament's approval as relief.

This may well be a mistake. Or at the very least, it's the wrong place to stop reading.

## The Risk

This isn't the reprieve it looks like. From 2 August 2026, enforcement and penalty powers materially ramp up, with fines running up to €15 million or 3% of global turnover, and transparency obligations become enforceable alongside them. However, it is worth being realistic about how this tends to play out. Using GDPR as a parallel: the headline fines that proved that regulation had teeth took years to arrive, and years after that, remain unsettled. Meta's €1.2 billion penalty<sup>1</sup> (still the largest GDPR fine to date) took five years to land and remains under active appeal in the Irish High Court. Amazon's \$886 million penalty<sup>2</sup> took three years to arrive, was upheld at first instance two years after that, and was then annulled on appeal in March 2026 on procedural grounds, with the case remanded for a fresh penalty assessment<sup>3</sup>. Even the fines held up as proof of serious enforcement are, in practice, live litigation for the better part of a decade rather than done deals.

But when GDPR enforcement did start, it went after the easy cases first: a missing legal basis (Google €50m<sup>4</sup>), undocumented security practices (BA £20m<sup>5</sup>, Marriott £18.4m<sup>6</sup>), and blatant surveillance (H&M €35.3m<sup>7</sup>). None of these were subtle interpretive disputes. The organizations caught in that first wave weren't unlucky. They had visible, easily evidenced gaps. If the Act's enforcement follows a similar pattern, the early risk isn't a headline-scale fine; rather, it's being one of the obvious, undocumented cases a regulator picks off first simply because they're easiest to prove.

Alongside that, from the same 2 August 2026 date, transparency obligations become enforceable. AI-generated content needs to be clearly marked, deepfakes disclosed and chatbots identified as such. These duties apply broadly, including to organizations that are deploying AI in the EU, not just those building or selling it.

## The Cost

Waiting doesn't just risk a fine. It changes how prepared you'll be when it matters. The standards and detailed guidance that will eventually fill in the Act's remaining specifics aren't written yet; the first harmonized standards for high-risk systems aren't expected until H2 2026, with the ones that matter most for Annex III conformity realistically landing in 2027<sup>8</sup>. GDPR ran this experiment already. Its own interpretive guidance was still being finalized right up against and past the 2018 deadline, and some of the hardest questions weren't settled until years afterward, through case law nobody could have anticipated in advance.

An organization that spends this runway building a working governance layer has somewhere to put that detail as it lands. A new standard becomes one more thing an existing process absorbs. An organization that spends the runway waiting must stand up an entire governance approach from scratch, against the full weight of a finished regulation, with no practical experience of applying any of it, on a compressed timeline. The first version is incremental. The second is theoretical right up until it's a scramble. And a scramble creates its own problem. GDPR's deadline produced a market flooded with support built to meet sudden demand, most of it without the track record that only comes from working the problem before it was urgent. Genuine expertise takes time to build and can't be compressed into the same few months everyone else is also trying to compress it into.

There was a full two-year gap between GDPR's adoption and its application date, and a 2018 survey found 60%<sup>9</sup> of compliance professionals still expected to miss the deadline. A year after it landed, only 28% had actually achieved compliance despite 78% having expected to<sup>10</sup>. Two years of runway produced almost the same readiness curve as no runway at all, because most of the work only happened once the deadline was unavoidable. The AI Act has just handed organizations an even longer window on its hardest obligations. Extending that same pattern isn't a safe bet. It's a well-documented risk and, for anyone who moves differently, a genuine head start.

“

GDPR's headline fines took three to five years to arrive. The first wave went after the easiest cases to prove, not complex legal disputes. That's the more useful precedent for what's coming.

”

## 2.

### Operational Ramifications of the Act

Being ready sounds straightforward until you ask what “ready” actually requires.

Most of what’s demanded by the Act isn’t complicated in isolation. In the simplest terms, it’s a series of concrete, achievable changes to how an organization designs, builds and runs AI systems. The difficulty is less about any single requirement and more about making all of them ordinary practice, rather than a project someone completes once and files away.

### Classification Before Development Starts, Not After

Risk classification determines almost everything else that follows: what oversight a system needs, what gets logged, what should be documented and whether a conformity assessment is required at all (Article 6, Annex III). Treating it as a sign-off at the end of a build gets this backwards. By the time a system is finished, most of the decisions that classification should have shaped have already been made, usually without anyone deliberately thinking about them as compliance decisions at the time.

A team that knows early on that a system might fall into a high-risk category can design logging and human-oversight touchpoints into the architecture from day one, at a fraction of the cost of adding them later. A team that finds this out after the system is live is choosing between an expensive retrofit and an uncomfortable conversation about why it wasn’t caught sooner.

The greater risk isn’t just a costlier retrofit. It risks discovering, only after months of development and investment, that the system falls into a category the Act prohibits outright. Article 5’s prohibited practices aren’t a compliance tier that can be managed; they are a hard stop. AI systems built without an early classification check can end up as a sunk cost in their entirety because they can’t legally be deployed at all.

### Human Oversight Has to Be a Person, Not a Policy

It’s common for organizations to describe human oversight in general terms. A human is “in the loop,” or a team “monitors” the system. It should instead be a named individual, trained on what the system does and where it can fail, with actual authority to intervene or halt it. That is a significant difference in commitment. It means someone owns the role, someone trains them for it and someone is held accountable if oversight didn’t happen the way it was supposed to.

In practice, this is often the requirement organizations are least prepared for. Mainly because it’s a resourcing decision nobody has explicitly made. “Someone will be watching it” is not the same as naming who, and organizations that leave this vague tend to discover the gap only when something has already gone wrong.

“Someone will be watching it”  
isn’t the same as naming who.”



### Monitoring Needs to be Active, and Incident Response Should be Pre-Built

Once a system is live, it needs to be watched against how it’s actually meant to be used, not just logged in case a question ever comes up. Logs need to be retained for a minimum of six months (Article 26 §6) and be genuinely accessible, not buried in a system nobody remembers how to query. And if a deployer has reason to believe a system is posing a risk, they’re expected to suspend their use and notify the right people without undue delay. For avoidance of doubt: the provider, importer/distributor, then the market surveillance authority (Article 26 §5).

“Without undue delay” is significant here. It assumes the organization already knows who gets notified, in what order and by whom. It’s not compatible with figuring out the escalation path after something’s already gone wrong. Building that route before it’s needed is a small amount of upfront planning that avoids a much larger amount of scrambling under pressure later.

## 2.

### Documentation Needs to be Built as the System is Built, not Reconstructed for an Audit

Technical documentation (architecture, training data provenance, testing and validation procedures, accuracy and robustness metrics) holds up far better when it's written alongside development than when it's assembled after the fact to satisfy a review. This is also where GDPR offers a useful preview of how early AI Act enforcement is likely to unfold. The first wave of major GDPR fines didn't target subtle interpretive disputes; they targeted clean, well-evidenced failures. Regulators picked the easy cases first, because they were the easiest to prove. If AI Act enforcement follows a similar pattern, and there are credible reasons to suspect it may, organizations may be most exposed when they can't produce evidence against the basic requirements of the regulation.

“The first wave of major GDPR fines didn't target subtle interpretive disputes; they targeted clean, well-evidenced failures.”

### Determining Your Role for Each System

Most organizations answer “Are we a provider or a deployer?” once, early on, and then treat it as settled. That's a reasonable starting point, but that is not where the risk sits. Role determination is generally AI system- and use-case-specific, and it can shift as responsibilities change over time.

Article 25 sets out the conditions under which a deployer can become a provider: putting your own name or trademark on a system; making a substantial modification to how it works; or fine-tuning a general-purpose model in a way that changes its intended purpose. These triggers don't depend on an organization expressly intending to take on the provider role. Organizations can find they've crossed the line from deployer to provider through common product decisions, inheriting all the additional legal obligations that arise with it.

For most organizations, role determination should be treated as an operational checkpoint per use case, not a one-off assessment. Most processes ask the question once at the start and never again. Whenever a system is materially developed, modified, rebranded, integrated or extended, it's prudent to re-check whether responsibilities have shifted. Skipping that check doesn't remove the obligations that come with a heavier role; it just means an organization discovers them later, usually at a worse moment than if it had checked when the change was made.

“Are we a provider or a deployer?”

## It's Worth Distinguishing Which Deadline Governs What

The prohibition of certain practices has been in force since February 2025 (Article 5 and 113). That's not a future risk but a current one, and AI systems falling into a prohibited category are already a live exposure rather than something to plan for later.

Article 50's transparency obligations become enforceable from 2 August 2026, alongside the AI Office's full penalty powers (Article 50 and 113).

The high-risk system obligations that most compliance content focuses on don't land until December 2027 and August 2028. Conflating these dates tends to produce one of two mistakes: either spending energy worrying about something that is still well over a year away, or missing something that's a matter of weeks.

Handled with the right sequencing, none of the above needs to be a scramble. It comes down to a small number of decisions, made early, owned by named people, and documented while they're actually happening rather than reconstructed afterwards.

“

The high-risk system obligations that most compliance content focuses on don't land until December 2027 and August 2028.

”



## 3.

### A Question of Culture, Not Checklists

Everything in the previous section is achievable. Assign an owner, classify before you build and keep documentation current as you go. What's harder is doing all of it consistently, across every system, as an organization keeps building new AIs and changing existing ones. That's where compliance programs typically fail. The challenge is establishing a standing way of working rather than a one-off project.

A checklist gets completed once, filed and revisited under pressure, usually after something's already gone wrong. A team that's built the habit into how it works doesn't need reminding. The person building the system should know what oversight it needs, because that question got asked at the start, not appended at the end.

The clearest test of whether an organization has actually made that shift is what happens when an existing system changes, rather than when a new one gets built. This is one of the most common sources of exposure because it is rarely treated as a compliance moment.

Take a reasonably ordinary scenario such as when an organization licenses a general-purpose AI model from a third-party provider and uses it, unmodified, to power a feature in its own product: a document-summarization tool that condenses incoming job applications for a recruiter to skim. At this point, the organization is simply a deployer. The obligations are fairly light: follow instructions, assign appropriate human oversight, monitor performance and retain logs where required.

Then the product team decides to fine-tune that model on the organization's own data, so instead of just summarizing applications for a human to read, it filters and shortlists candidates itself, flagging which applicants should move to interview. Reasonable, common and from a product perspective, barely worth more than a routine sign-off.

But recruitment and candidate selection is an explicit high-risk use case under the Act, and fine-tuning a general-purpose model in a way that changes its intended purpose to a high-risk one is one of the conditions under which a deployer becomes a provider (Article 25 §1c). Once that line is crossed, the organization has picked up an entirely different set of obligations. A quality management system (Article 17), technical documentation built to a specific standard (Article 18), a conformity assessment before the system can legally be used this way (Article 43), a declaration of conformity (Article 47) and registration in an EU database (Article 49). None of that existed as a requirement the day before the fine-tuning shipped. All of it exists as a requirement the day after.

The organization didn't set out to become a provider. Nobody in the product review meeting was thinking about Article 25; they were thinking about shortlisting accuracy. It is not what the team intended, but what is important for the Act is what the system now does and who's responsible for it. And because most organizations' change-management process fails to ask, "Has our role just shifted?" at the point a modification like this ships, the gap sits there as an unidentified risk until an audit, an incident or a market surveillance request forces the question.

## 4.

This pattern shows up in other forms too. A fintech rebrands a vendor's embedded credit-scoring widget (already an explicit high-risk use case) under its own name for an affordability check at checkout: the moment its own brand sits on someone else's, it picks up the same set of provider obligations described above (Article 25 §1a).

An insurer's product team migrates to a new claims platform and streamlines the approval flow for a vendor's high-risk life-insurance risk-assessment system. Every decision used to route to a human underwriter before a policy was issued, but now standard applications go straight to the customer. Nobody on the migration thought of that step as a compliance control, just friction to cut. As soon as the human-oversight safeguard the original conformity assessment relied on was removed, they became the provider (Article 25 §1b).

In neither case is the mechanism a typical governance decision. It's an ordinary product or delivery decision, made without it being flagged as a compliance moment, that results in significant compliance consequences beyond the scope of the teams that made it.

Governance built into the culture matters more than governance bolted onto a project plan. A checklist completed at launch won't catch a role reclassification that happens 18 months later, as part of routine product work nobody thought to route through compliance.

A team that is used to asking these questions each time something changes can catch it. In practice, it is rarely a formal step but rather a well-fostered habit: Have we classified this before building it? Do we know who, by name, has oversight? Has anything about this system (what it does, how it's built or whose name is on it) changed since we last checked our role? Nothing here gets checked off. It's just asked, every time, because asking it has become how the team works rather than a stage in a process.

“

Governance built into the culture matters more than governance bolted onto a project plan.

”



## Bringing Lifecycle, Technology and Governance together

If culture is the right frame, the next question is practical. What does an organization need in place to run that culture day to day? Something concrete enough that a team can point to it and use it, not just agree with it in principle?

Our view, from doing this work directly, is that it comes down to one principle: that an organization's AI lifecycle, the technology it builds and uses and the governance that sits over both, need to live in one place. In practice, that usually means moving away from three separate functions that occasionally sync up, toward a single connected view where a system's classification, its documentation, its oversight assignments and its ongoing monitoring are visible together. They're all part of the same decision, not three different ones made by three different teams at three different times.

That principle isn't theoretical. Assessing how organizations approach AI governance in practice, against the Act's actual text rather than a general sense of “AI risk,” surfaces the same handful of repeated gaps. These are not failures of any particular tool, but gaps in how governance gets designed and run.

The most consistent one is a confusion between an operational risk score and a legal risk classification. Most organizations build some form of internal triage, ranking systems by how risky they seem, usually on a numeric scale, so attention goes where it's needed most. That's a genuinely useful instinct. But the Act doesn't ask “roughly how risky is this?” It asks a binary legal question. “Does this system fall into a high-risk category under Article 6, or doesn't it?” An internal severity score and a statutory classification can look similar enough that organizations treat them as the same thing, when getting that classification wrong has real legal consequences a risk score was never designed for.

“

The AI Act doesn't ask “How risky is this?”  
It asks “Is it high-risk or not?”

”

## 4.

A second, related pattern is that organizations are generally good at recording that a check happened, and much less good at forcing the right check to happen at the right moment. Screening for prohibited practices under Article 5, for instance, needs to happen before a system moves into development, not as one question among many in a general assessment further down the line. Where that gate isn't built in early, nothing stops a project accumulating months of work and budget before anyone asks whether it's legally viable at all. I've seen that gap specifically, and it's a far more expensive way to find out than the alternative.

A third is role determination. Many organizations treat every AI the same way regardless of whether they're acting as a provider or a deployer for it, applying one uniform process and set of questions across all to ensure comprehensive compliance. They don't determine the operational role they are actually participating as for this specific system. Given that role can shift the moment a system is modified, rebranded or repurposed, a process that doesn't force this question at the point of change is likely to miss exactly the reclassification risk described earlier in this paper.

A fourth pattern cuts the other way from the first three: over-applying the Act can be a governance failure mode, not just a safe default. It's tempting to treat every AI system an organization touches as in scope of the Act and run the full weight of obligations against all of them. Known as a "better safe than sorry" approach, this can backfire if it obscures the territorial/market nexus that determines whether obligations from the Act even apply at all. Only AIs genuinely deployed to the EU should be scoped and bear the additional compliance burden. Governance approaches that don't address this early end up applying rigorous, resource-intensive controls to systems that were never in scope to begin with. This doesn't make an organization safer; it just spends effort on the wrong things and risks the real in-scope systems getting less attention because the net was cast too wide.

A fifth point, see's the confusion of treating transparency as if it were a rung on the risk ladder, sitting somewhere between High-Risk and Minimal-Risk. It isn't. The Act's classification structure produces a binary legal outcome: a system is High-Risk or it isn't. Transparency obligations under Article 50 are a separate, independent layer that can apply regardless of where a system sits on that classification. Where governance processes build transparency in as if it were a lighter compliance tier, two things go wrong at once. Systems correctly classified as High-Risk end up running transparency controls that don't apply to them, while systems that are Minimal-Risk but still trigger a genuine Article 50 duty (a chatbot, a deepfake generator) could be waved through as low-risk without anyone checking the transparency question at all. Treating transparency as its own determination, asked independently of classification, avoids both failure modes at once.

A sixth, and probably the least visible, issue is that having a broad set of controls isn't the same as having the right ones. Most organizations build compliance controls from established general risk and compliance frameworks, which is a sound starting point. But the Act creates a number of specific, narrow obligations, such as particular disclosure duties, formal artifacts a system needs before it can be placed on the market, and specific notification timelines once something goes wrong, that a generic set of controls wouldn't address, even when it looks, on paper, like it should. A control set can appear comprehensive and still leave an organization with nothing that maps to what a regulator might specifically ask for when they come looking. The fix isn't more controls in general. Its controls built against the Act's own text, article by article, rather than adapted from a broader risk framework and assumed to be close enough.

None of these are exotic findings, and none of them are failures of intent or effort. They're what happens when general-purpose governance and risk management meets a regulation that asks very specific legal questions. They're just not the same problem, and organizations that assume they are, tend to discover the difference later than they'd like.

This is also why the tooling question and the regulatory question can't be separated. A platform configured only for general governance would faithfully record decisions without knowing whether they're the right ones. Regulatory expertise without a working system to apply it against stays theoretical. Both are needed together, which is precisely the combination most organizations find hardest to build in-house.

It's the combination I've spent the past year building for clients working through EU AI Act compliance, pairing platform configuration with regulatory expertise applied directly against the Act's text, rather than treating either as sufficient on its own. ServiceNow's AI Control Tower is the platform I work with most closely for this. It's built as the foundation this kind of governance needs, holding classification, documentation, oversight assignment, and monitoring status together in one connected view. Getting the most out of that foundation, on any platform, means configuring it deliberately for what the Act specifically requires, rather than assuming general governance capability and regulatory compliance are the same thing. That configuration work, where regulatory expertise is applied directly to how a platform is set up and used, is where the real value sits, and it's the same work regardless of which platform an organization starts from.



# 5.

## The Window, and What to Do With It

We opened this paper with the claim that the next 12-18 months are an opportunity, not a deferral, and that the organizations who treat it that way may end up somewhere genuinely different from the ones who don't.

The deadline that matters most immediately is now, not eighteen months away. Prohibited-practice enforcement and transparency obligations are live from 2 August 2026 and if GDPR is any guide, the organizations caught early won't be the ones facing complex, contestable judgment calls. They'll be the ones with the easiest gaps to prove, the missing documentation and the absent screening steps described earlier in this paper. That's worth acting on now, on its own, regardless of anything else in this piece.

**1 August  
2024**

### Prohibited Practices

Art. 4 AI literacy duties and Art. 5 bans already took effect. No route to compliance possible for prohibited practices, with limited exceptions.

**2 February  
2025**

### GPAI Obligations

Arts. 53-55 Provider duties have already taken effect for General Purpose AIs. Organisations training (and, in some cases, fine-tuning) their own models may already carry provider duties depending on how the model is placed on the market / used / modified.

**2 August  
2026**

### Transparency Obligations + Enforcement Powers

Art. 50 transparency disclosure duties step into effect (disclosing AI to users i.e. with chatbots and generative tools). Note: Existing AI Systems can defer watermarks until 2 Dec 2026 Enforcement Powers to issue penalty fines running up to €15 million or 3% of global turnover become active.

**2 December  
2027**

### Stand-Alone AI (High Risk)

High-Risk obligations for stand-alone AI systems under Annex III come into force (recruitment, credit scoring, law enforcement, education, border control tools).

**2 August  
2028**

### Embedded AI (High Risk)

High-Risk obligations for Embedded AI in regulated products under Annex I will take effect (medical devices, machinery, vehicles, etc.).



5.

But the larger opportunity sits in the space the Digital Omnibus opened up. Most of the market is going to spend the next year and a half relieved that the harder obligations moved out to 2027 and 2028. Some of that relief will turn into genuine inaction, the same way it did the first time an EU regulation handed organizations a long runway. GDPR's own two-year gap produced almost the same readiness curve as no runway at all, because most of the work only happened once the deadline was unavoidable.

The organizations that use this time differently won't just be compliant sooner. They'll have built the habit this paper has tried to describe rather than the checklist most probably default to.

- Scoped AIs in and out of the EU, clear on whether the regulation applies at all.
- Classification asked and assessed before a system is built, not after.
- Role determination revisited whenever a system materially changes, not settled once and forgotten.
- Distinguished between an operational risk score and a legal risk classification.
- Transparency treated as its own question, not folded into a risk tier it doesn't belong to.
- Controls built against what the Act specifically requires, not borrowed from a broader framework and assumed to be close enough.

It's easier to build early, a piece at a time, than to assemble it all at once against a deadline that is rapidly approaching or has passed. It also means having real experience to draw on, rather than a theoretical reading of a regulation, by the time the harder obligations and the detailed standards eventually land.

If there's a single practical starting point, it is not to wait for a neatly completed answer before you begin. The standards that will eventually fill in the Act's remaining detail aren't published yet and won't be for some time. Organizations that wait for total clarity before acting are likely to spend this window exactly the way GDPR's did—mostly idle, then competing for the same thin layer of genuine expertise as everyone else who also waited, right as the deadline makes waiting impossible. Organizations that start now, working from the Act's intent rather than a finished product, will be better placed to capture future competitive advantage.

“

It's easier to build early, a piece at a time, than to assemble it all at once against a deadline that is rapidly approaching or has passed.

”

Key Terms

| TERM                 | DEFINITION (AS USED IN THE PAPER)                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       | DEFINING ARTICLE(S)                                                          |
|----------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------|
| DIGITAL OMNIBUS      | The legislative package (COM(2025) 836 final), adopted June 2026, that deferred the Act's high-risk obligations to Dec 2027 (Annex III) and Aug 2028 (Annex I)                                                                                                                                                                                                                                                                                                                                                                                                                                          | N/A.<br>An amending document, not part of the base Regulation (EU) 2024/1689 |
| PROHIBITED PRACTICES | AI practices banned outright regardless of risk classification: manipulation, exploitation of vulnerabilities, social scoring, certain biometric uses, etc.<br><br>• Manipulative or deceptive techniques<br>• Exploitation of vulnerabilities<br>• Social scoring<br>• Untargeted facial image scraping<br>• Emotion recognition in workplaces and schools<br>• Real-time remote biometric identification for law enforcement<br>• Profiling-based individual crime-risk prediction<br>• Biometric categorisation for sensitive attributes<br><br>In force since 2 Feb 2025. Limited exceptions apply. | Art. 5                                                                       |
| RISK CLASSIFICATION  | The binary legal determination of whether a system is High-Risk or not. Not a scale or score.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           | Art. 6                                                                       |

|                                              |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    |                                      |
|----------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------|
| <b>HIGH-RISK AI SYSTEM</b>                   | <p>A system falling within one of the Annex III use-case categories:</p> <ul style="list-style-type: none"><li>• Biometric use</li><li>• Critical Infrastructure</li><li>• Education/Training</li><li>• Workforce Management</li><li>• Access to Services</li><li>• Law / Justice</li><li>• Migration / Border-control</li></ul> <p>Or Annex I:</p> <ul style="list-style-type: none"><li>• A safety component</li><li>• Critical product type</li></ul> <p>... triggering the full compliance obligation set.</p> | Annex I, Annex III                   |
| <b>MARKET/ TERRITORIAL NEXUS</b>             | The condition that triggers the Act's applicability (placing a system on the EU market, or its output being used in the EU) rather than the operator's location alone.                                                                                                                                                                                                                                                                                                                                             | Art. 2                               |
| <b>PROVIDER</b>                              | The natural/legal person that develops an AI system (or has one developed) and places it on the market or puts it into service under its own name/trademark.                                                                                                                                                                                                                                                                                                                                                       | Art. 3(3); obligations at Art. 16    |
| <b>DEPLOYER</b>                              | The natural/legal person using an AI system under its own authority, other than for personal non-professional use.                                                                                                                                                                                                                                                                                                                                                                                                 | Art. 3(4); obligations at Art. 26    |
| <b>RECLASSIFICATION (PROVIDER/ DEPLOYER)</b> | The mechanism by which a deployer, distributor, importer, or third party becomes legally treated as a provider (by rebranding, substantially modifying, or repurposing a system).                                                                                                                                                                                                                                                                                                                                  | Art. 25                              |
| <b>SUBSTANTIAL MODIFICATION</b>              | A change to an AI system, after it's on the market, that affects its compliance with the Act or changes its intended purpose. One of the triggers for reclassification.                                                                                                                                                                                                                                                                                                                                            | Art. 3(23)                           |
| <b>GENERAL-PURPOSE AI (GPAI) MODEL</b>       | A model trained on broad data, capable of competently performing a wide range of tasks, that can be integrated into various downstream systems.                                                                                                                                                                                                                                                                                                                                                                    | Art. 3(63)                           |
| <b>HUMAN OVERSIGHT</b>                       | The requirement that a named, competent, trained individual has real authority to monitor and intervene in a high-risk system's operation.                                                                                                                                                                                                                                                                                                                                                                         | Art. 14; deployer duty at Art. 26(2) |
| <b>QUALITY MANAGEMENT SYSTEM (QMS)</b>       | The documented system a provider must maintain to ensure ongoing compliance across a high-risk system's lifecycle.                                                                                                                                                                                                                                                                                                                                                                                                 | Art. 17                              |

|                                            |                                                                                                                                                                                                      |                                                                        |
|--------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------|
| <b>TECHNICAL DOCUMENTATION</b>             | The provider's record of a system's architecture, training data, testing, and performance. Required before market placement.                                                                         | Art. 11, Art. 18                                                       |
| <b>LOGS / LOG RETENTION</b>                | System-generated records that providers and deployers must keep; deployers must retain for a minimum of six months.                                                                                  | Art. 12 (providers), Art. 19, Art. 26(6) (deployers)                   |
| <b>CONFORMITY ASSESSMENT</b>               | The formal procedure a high-risk system must undergo before being placed on the market or put into service.                                                                                          | Art. 43                                                                |
| <b>DECLARATION OF CONFORMITY (DOC)</b>     | The provider's formal statement that a system meets the Act's requirements.                                                                                                                          | Art. 47                                                                |
| <b>CE MARKING</b>                          | The physical/documentary mark affixed to a compliant high-risk system, in accordance with the DoC.                                                                                                   | Art. 48                                                                |
| <b>EU DATABASE REGISTRATION</b>            | The registration of a high-risk system (and, where applicable, the deployer) in the EU-wide public database.                                                                                         | Art. 49                                                                |
| <b>TRANSPARENCY OBLIGATIONS</b>            | The distinct duty layer requiring disclosure for AI-generated content, deepfakes, chatbot interactions, and emotion-recognition/biometric-categorisation systems. Independent of risk classification | Art. 50                                                                |
| <b>SERIOUS INCIDENT</b>                    | An incident or malfunction causing death, serious harm, critical infrastructure disruption, or fundamental-rights infringement, triggering mandatory notification.                                   | Definition at Art. 3(49); notification duty at Art. 73                 |
| <b>MARKET SURVEILLANCE AUTHORITY (MSA)</b> | The national authority responsible for enforcement and to whom risks/incidents must be reported.                                                                                                     | Art. 3(26); role detailed across Arts. 74–79                           |
| <b>AI OFFICE</b>                           | The Commission body with centralised competence, primarily over GPAI models and systemic-risk obligations.                                                                                           | Established under the Regulation; powers referenced throughout Ch. VII |

## References:

Meta inquiry concerning data transfers – Data Protection Commission Ireland - <https://www.dataprotection.ie/en/dpc-guidance/decisions/inquiry-concerning-data-transfers-eueea-us-meta-platforms-ireland-limited-its-facebook-service>

Amazon hit with \$886m fine – British Broadcasting Corporation - <https://www.bbc.co.uk/news/business-58024116>

Administrative Court ruling in the case between Amazon and the CNPD – CNIL (COMMISSION NATIONALE DE L'INFORMATIQUE ET DES LIBERTES) - <https://cnpd.public.lu/en/actualites/national/2026/03/arret-ca-amazon-cnpd.html>

France fines Google €50m for 'lack of transparency & valid consent' – Business and Human Rights Centre - <https://www.business-humanrights.org/en/latest-news/france-fines-google-50-million-for-lack-of-transparency-valid-consent-regarding-advert-personalization-using-gdpr-rules/>

British Airways fined £20m over data breach – British Broadcasting Corporation - <https://www.bbc.co.uk/news/technology-54568784>

Marriott Hotels fined £18.4m for data breach that hit millions – British Broadcasting Corporation - <https://www.bbc.co.uk/news/technology-54748843>

Hamburg Commissioner Fines H&M – EFDPO (European Federation of data protection officers) - <https://www.efdpo.eu/hamburg-commissioner-fines-hm-35-3-million-euro-for-data-protection-violations-in-service-centre/>

Standardisation of the AI Act – European Commission - <https://digital-strategy.ec.europa.eu/en/policies/ai-act-standardisation>

2018 GDPR Compliance Report – Crowd Research Partners - <https://cybersecureforum.co.uk/briefing/only-40-of-uk-businesses-ready-for-gdpr/>

2019 GDPR readiness survey – Capgemini Research Institute - <https://www.helpnetsecurity.com/2019/09/30/companies-gdpr-readiness/>



## Contact us:

Synechron Limited, 95 Gresham Street,  
London - EC2V 7NA, United Kingdom

[enquiries@synechron.com](mailto:enquiries@synechron.com)

[synechron.com](https://synechron.com)

