

EU AI Act Compliance Kit

Risk classification, live deadlines and the four documents you will be asked for

Dates in this document were verified on 14 August 2026, after the Digital Omnibus (Regulation (EU) 2026/1744, in force 27 July 2026) moved several deadlines. Much of the compliance material circulating online still quotes the pre-Omnibus calendar. Check the date on anything else you read.

1. Where the calendar actually stands

Obligation	Applies from	Status
Prohibited practices (Art. 5)	2 February 2025	In force
General-purpose AI model obligations	2 August 2025	In force
Transparency obligations (Art. 50)	2 August 2026	In force now
Machine-readable marking of synthetic content	2 December 2026	Upcoming
High-risk systems, Annex III	2 December 2027	Postponed by the Omnibus (was 2 Aug 2026)
High-risk systems, Annex I	2 August 2028	Postponed by the Omnibus

The one that bites today. Article 50 was deliberately left out of the Omnibus deferral. If your product has a chatbot, generates images, audio or video, does emotion recognition, or produces anything a person could mistake for human-made, the transparency duties apply *now* and national market surveillance authorities can enforce them. The high-risk postponement bought time for hiring, credit and healthcare systems. It bought none for chatbots.

2. Classify your system in four questions

1. **Does it make or materially inform decisions about people?** Hiring, credit, insurance, education, healthcare, law enforcement. Yes puts you in Annex III high-risk.
2. **Does it interact with people, or generate synthetic content?** Chatbot, assistant, image or voice generation. Yes puts you in Art. 50 transparency, which is live.
3. **Does it run in critical infrastructure?** Water, electricity, transport, digital infrastructure. Yes is high-risk.
4. **Does it process personal data?** Yes means GDPR applies alongside, not instead.

None of the above and you are in minimal risk: no specific AI Act duties, GDPR unchanged. Say so in writing anyway. A documented "we assessed this and it is minimal risk" is worth considerably more in an audit than silence.

The four levels

- **Unacceptable** — banned outright. Social scoring, untargeted facial-image scraping, real-time remote biometric identification in public spaces (with narrow law-enforcement exceptions).
- **High** — permitted with a heavy compliance load. Annex III covers the decisions-about-people categories.
- **Limited** — transparency only. Chatbots, deepfakes, synthetic media.
- **Minimal** — no specific obligations. Spam filters, recommendation engines with no consequential effect, game AI.

3. What each level actually requires

Limited risk (live today)

- Users are told they are interacting with an AI system, before or at first interaction, unless it is obvious to a reasonably observant person.
- Synthetic image, audio and video content is disclosed as artificially generated.
- Deepfakes depicting real people or events are labelled as such.
- Machine-readable marking of that content becomes mandatory 2 December 2026. If you generate media, that is your next build task, not a 2027 problem.

High risk (build now, obligatory 2 December 2027)

- Risk management system, maintained across the lifecycle rather than written once.
- Data governance: quality criteria, provenance, bias examination of training and validation sets.
- Technical documentation to Annex IV.
- Automatic logging of events, retained so a decision can be reconstructed.
- Human oversight designed in, with a real ability to intervene or override.
- Accuracy, robustness and cybersecurity appropriate to the purpose.
- Registration in the EU database before the system is placed on the market.

Penalties

Prohibited practices: up to €35 million or 7% of worldwide annual turnover, whichever is higher. Most other breaches: up to €15 million or 3%. Supplying incorrect or misleading information to authorities: up to €7.5 million or 1%. Smaller caps apply to SMEs and start-ups.

4. The four documents

These are the artefacts an auditor, an enterprise customer's procurement team, or a market surveillance authority will ask to see. Blank templates for all four ship with the open-source kit.

1. **Risk assessment** — what the system is, who it affects, which level it falls in and the reasoning that got you there. This is the document everything else hangs from.
2. **Technical documentation** — architecture, data sources, model provenance, performance measures, security controls. Required for high-risk under Annex IV; worth keeping for limited-risk as audit readiness.
3. **Transparency notice** — the user-facing disclosure. Short, in the interface, not buried in a policy page.
4. **Incident response record** — for a harmful or biased decision, a data leak, a successful prompt injection, or sustained wrong output. Providers of high-risk systems must report serious incidents to market surveillance authorities.

5. A five-phase run, if you are starting from nothing

1. **Inventory and classify.** List every AI system and feature you ship, including third-party APIs behind your own branding. Classify each. Write down the reasoning.
2. **Close the transparency gaps first.** They are live, they are cheap, and they are the ones visible from outside your company.
3. **Risk management and logging** for anything that landed in high risk.
4. **Write the technical documentation** while the build is fresh, not at audit time.
5. **Monitoring and review.** Quarterly is a reasonable cadence. The regulation is still being interpreted, and guidance moves.

What this kit is not. It is an engineering-side working document, written by people who build and secure AI systems. It is not legal advice, it does not create a lawyer-client relationship, and it does not substitute for counsel on your specific product. Classification at the boundaries, particularly around Annex III, is genuinely contested and worth paying a lawyer to decide.

What we can tell you accurately is the technical shape of compliance: what to log, how to make oversight real rather than decorative, what belongs in Annex IV documentation, and how to mark synthetic content in a machine-readable way.

DeviDevs Technologies S.R.L. · AI agents, automation and AI security

Open-source kit, including the risk calculator and all four templates: github.com/DeviDevs-Technologies/eu-ai-act-compliance-kit (MIT)

Primary sources: artificialintelligenceact.eu/article/50 · Regulation (EU) 2024/1689 · Regulation (EU) 2026/1744
devidevs.com

This document is yours to keep, use and share, whether or not you ever contact us.