

EU AI ACT: HIGH-RISK SELF-ASSESSMENT CHECKLIST

The Digital Omnibus pushed the Annex III high-risk compliance deadline to **2 December 2027**. That is runway to get your classification right — not a reason to leave it for next year's roadmap. These are the ten questions we ask first when we run AI Act classification for a client. Answering them honestly, against your actual architecture and use case rather than your product description, is what makes a classification defensible later — including in Series A diligence.

1. Does the system perform any function listed in Annex III — biometrics, critical infrastructure, education or vocational training, employment or worker management, access to essential public or private services, law enforcement, migration/asylum/border control, or administration of justice and democratic processes?
 - Why it matters: this is the primary trigger under Article 6(2). If none of these apply, the Annex III route doesn't apply on its own — though the product may still fall under the separate Annex I route (Q2).
2. Is the system a safety component of, or itself, a product regulated under EU harmonisation legislation (e.g. medical devices, machinery, toys) that requires third-party conformity assessment?
 - Why it matters: this is the independent Article 6(1) trigger. It runs on its own timeline — deferred to 2 August 2028 under the Digital Omnibus — separate from the Annex III deadline your team may already be tracking.
3. Does the system profile natural persons — automatically process personal data to evaluate aspects like performance, behaviour, preferences, or location?
 - Why it matters: under Article 6(3), any Annex III system that profiles individuals is always high-risk. None of the four narrow-task exceptions below apply once profiling is present.
4. Is the system limited to a narrow procedural task, such as converting unstructured data into a structured format or categorising incoming documents?
 - Why it matters: one of four Article 6(3) conditions that can take an Annex III system out of high-risk — but only where profiling isn't involved and the task genuinely doesn't influence the substantive decision.
5. Does the system only improve the result of a previously completed human activity, without materially changing the outcome?
 - Why it matters: a second Article 6(3) exception condition. The answer depends on whether the AI output could plausibly shift the human decision — a defensible answer requires looking at the actual workflow, not the product description.

6. Does the system detect decision-making patterns or deviations from them, without replacing or influencing the prior human assessment absent proper human review?
 - Why it matters: a third Article 6(3) exception condition. “Proper human review” is doing a lot of work here — the exception fails if review is nominal rather than substantive.
7. Does the system perform a preparatory task ahead of an Annex III-relevant assessment, rather than the assessment itself?
 - Why it matters: the fourth Article 6(3) exception condition — relevant where your system feeds into, but doesn't make, a high-risk decision, such as screening applications before a human evaluates them.
8. If you believe the system isn't high-risk, has that assessment been documented before the system is placed on the market or put into service?
 - Why it matters: Article 6(4) requires this documentation regardless of which way the classification lands, and triggers a registration obligation under Article 49(2). Undocumented self-assessments are exactly what surfaces as a gap in diligence.
9. Which compliance deadline actually applies to your system — 2 December 2027 (Annex III use-case route) or 2 August 2028 (Annex I product-safety route)?
 - Why it matters: the two routes now sit on different timelines under the Digital Omnibus. Building your roadmap against the wrong one means either scrambling late or over-investing a year early.
10. Does any part of the system's functionality touch a prohibited practice under Article 5 — for example subliminal manipulation, social scoring, untargeted facial recognition scraping, or emotion inference in workplaces or schools?
 - Why it matters: prohibited practices have been enforceable since February 2025, unaffected by the Digital Omnibus delay. This is a separate, already-live compliance track from the high-risk questions above.

What this checklist doesn't do

This is a starting point for spotting where your classification needs a defensible answer — it isn't a substitute for one. If more than a couple of these questions don't have a clear, documented answer, that's the gap worth closing before it surfaces somewhere less convenient.