

DOCUMENT 1 OF 2 · 10 PAGES

EU AI Act Classification Questionnaire

Use-case-by-use-case classification against the EU AI Act risk tiers (prohibited · high-risk · limited-risk · minimal-risk) plus GPAI obligations — with named sign-off and decision rendering.

Audience · AI governance, Legal, DPO, programme owner, business unit, board / audit committee

Companion documents · Conformity Gap-Assessment + NIST AI RMF Cross-Mapping

Use as · Audit-grade artefact requiring named human sign-off

Version · v1.0 · 2026-05

§ 1

EU AI Act overview & risk tiers

The EU AI Act applies a risk-based framework. Providers and deployers of AI systems must classify each system, and obligations scale with risk. Most enterprise AI sits in either "limited risk" or "high risk"; misclassification is a recurring audit finding.

Tier	Examples	Core obligations
Prohibited (Title II)	Subliminal manipulation, social scoring by public authority, untargeted facial-image scraping, exploitation of vulnerabilities, real-time biometric ID in public spaces (with narrow exceptions)	Cannot be placed on the market or used
High-risk (Annex III + sector products)	Biometric ID, critical infrastructure, education, employment, essential services (credit, insurance), law enforcement, migration, justice, democratic processes; safety-component AI in regulated products (medical devices, machinery, toys, etc.)	Risk management (Art. 9), data governance (Art. 10), technical doc (Art. 11), record-keeping (Art. 12), transparency (Art. 13), human oversight (Art. 14), accuracy/robustness/cybersecurity (Art. 15), conformity assessment, CE marking, post-market monitoring
Limited risk	Chatbots, emotion recognition, biometric categorisation, deep-fake content	Transparency / disclosure obligations
Minimal / no risk	Most general productivity AI	Voluntary codes of conduct
GPAI (general-purpose AI)	Foundation models & downstream systems	Technical documentation, training-data summaries, copyright policy; systemic-risk models add evaluations, incident reporting, cybersecurity

§ 2

Step 1 · Prohibited-practice screen

#	Question	Y / N
P-1	Does the system use subliminal techniques beyond a person's consciousness to materially distort behaviour?	
P-2	Does it exploit vulnerabilities of specific groups (age, disability, socio-economic) to materially distort behaviour?	
P-3	Is it used by a public authority for social scoring leading to detrimental treatment outside the original collection context?	
P-4	Does it perform untargeted scraping of facial images from the internet or CCTV for biometric databases?	
P-5	Does it infer emotions in workplace or educational contexts (outside medical / safety)?	
P-6	Does it categorise individuals based on biometric data to infer race, political opinions, trade-union membership, religion, sex life or sexual orientation?	
P-7	Does it use real-time remote biometric identification in publicly accessible spaces by law enforcement (outside the narrow exceptions)?	

Any "Yes" on this page means the use case is prohibited under Title II and must not proceed. Re-scope or terminate.

§ 3

Step 2 · High-risk classification screen (Annex III)

Annex III lists categories of high-risk AI. A "Yes" to any of the questions below places the system in high-risk unless the narrow Article 6(3) carve-out applies (purely accessory to the decision, no material influence on the outcome, no profiling).

#	Question	Y / N
H-1	Biometric identification, categorisation, or emotion recognition (outside prohibited list)?	
H-2	Safety component of critical infrastructure (water, gas, electricity, road, rail, digital infrastructure)?	
H-3	Education & vocational training (access, evaluation, exam-cheating detection)?	
H-4	Employment, workers' management or access to self-employment (recruitment, allocation of tasks, performance, termination)?	
H-5	Access to and enjoyment of essential private services (credit scoring, insurance pricing, emergency dispatch)?	
H-6	Access to essential public services and benefits (social security, eligibility)?	
H-7	Law enforcement (profiling, risk assessment, polygraph, evidence reliability)?	
H-8	Migration, asylum or border control management?	
H-9	Administration of justice and democratic processes (judicial decision research, election influence)?	

§ 4

Step 3 · Safety-component classification (regulated products)

Use this screen when the AI is a safety component of, or itself, a product covered by Union harmonisation legislation listed in Annex I.

Product area	Examples	Y / N
Medical devices (MDR / IVDR)	SaMD, diagnostic imaging, IVD	
Machinery	Industrial machinery with AI safety	
Lifts & safety components		
Toys		
Radio equipment / RED		
Pressure equipment, gas appliances		
Civil aviation / rail / vessels / vehicles		

Yes to any → high-risk under the AI Act, with sector-conformity interaction (e.g. MDR + AI Act). Refer to product-specific notified-body procedure.

Step 4 · Article 6(3) carve-out test (high-risk only)

If Annex III applies (§ 3) but the system meets all of the following, it may avoid high-risk classification under Article 6(3). Carve-out does not apply to profiling of natural persons.

1. The system is intended to perform a narrow procedural task (e.g. structuring information, classifying documents).
2. It is intended to improve the result of a previously completed human activity.
3. It is intended to detect decision-making patterns or deviations without replacing or influencing the human assessment.
4. It performs a preparatory task to an assessment that is materially carried out by humans.

Decision

Field	Detail
Carve-out applies?	Yes / No
Rationale	Specific reference to (1)–(4)
Profiling involved?	If yes → carve-out unavailable; remains high-risk
Sign-off	Legal + AI governance

Step 5 · Limited-risk transparency triggers

#	Trigger	Obligation	Applies?
L-1	System interacts directly with natural persons (chatbot)	Inform the person they are interacting with AI	
L-2	Emotion recognition / biometric categorisation	Inform exposed persons	
L-3	Generates synthetic content (deep-fake)	Mark output as artificially generated	
L-4	Generates / manipulates text on matters of public interest	Mark as artificially generated, unless human review	

§ 7

Step 6 · GPAI / foundation-model classification

Question	Y / N
Is the use case based on a general-purpose AI model (GPAI)?	
Are we acting as a deployer of GPAI, or are we the provider?	Deployer / Provider
Has the provider disclosed: technical documentation, training-data summary, copyright policy?	
Does the GPAI present systemic risk (model meets training-compute / capability thresholds defined by Commission)?	
If systemic-risk: are evaluations, incident reporting and cybersecurity obligations evidenced by the provider?	

A deployer of a GPAI does not automatically become a provider of high-risk AI; classification depends on use. However, the deployer's downstream obligations (transparency, GPAI provider acknowledgment) must be respected.

Decision rendering

Field	Detail
Use case (one sentence)	
Steps 1–6 outcome	P · H · A6(3) · L · GPAI flags
Final classification	Prohibited / High-risk / Limited / Minimal & GPAI obligations Y/N
Rationale (3–5 sentences)	
If high-risk: conformity-assessment route	Internal control · Notified body (when required)
Reviewer	AI governance + Legal
Approver	DPO + Legal + sponsor

Re-classification triggers. Material change to intended use, target population, or feature set. Vendor model update that changes capability. Regulator clarification. Annual review.

Worked examples

Use case	Classification	Why
Retail recommender on website	Limited risk (if no profiling-based price discrimination) / Minimal	Annex III not triggered; transparency on AI-generated content not required for product recommendations
Credit scoring for consumer loans (EU)	High-risk (Annex III H-5)	Essential private service access
Triage copilot in EU hospital	High-risk (safety component of MDR-classified device)	Sector + AI Act
Foundation model used internally for code drafting	Minimal (deployer); provider obligations on the vendor	No Annex III trigger; transparency for end users
Predictive maintenance on plant assets	Limited risk (or minimal)	Operational; not a safety-component of a regulated product in this configuration
Resume screening for hiring	High-risk (Annex III H-4)	Employment access
Chatbot answering customer questions	Limited risk	Transparency: tell users they are interacting with AI

Examples are illustrative; final classification requires Legal sign-off on the specific configuration and population.

Sign-off & attestation

Sign-off block

Role	Name	Date	Signature
AI governance lead			
Legal			
DPO			
Programme sponsor			
Business owner			

The questionnaire is completed for every new initiative and re-completed on material change. Outputs feed into the gating workbook (G0 evidence) and the audit-committee briefing pack.