

EU AI Act: Implementation Guide

FOR DEPLOYERS OF AI SYSTEMS



EU AI Act: Implementation Guide

For deployers of AI systems

Introduction

As AI continues to transform industries and reshape business operations, organizations operating within or doing business in the EU must prepare for the sweeping changes introduced by the EU AI Act.

In this guide, we provide a practical overview of the new regulatory landscape **for ‘deployers’ of AI systems**, highlighting the importance of understanding your AI systems, ensuring compliance, and safeguarding data and privacy. With significant penalties for non-compliance and evolving requirements for high-risk AI systems in particular, now is the time to take a proactive approach to AI governance.

Most organizations will be ‘deployers’ of AI systems. The EU AI Act defines a deployer as a person or entity (business, non-profit organization, government body, etc) using an AI system under its authority. Organizations can have overlapping obligations as ‘providers’ of AI systems, which are outside the scope of this guide.

Why the EU AI Act matters

If your organization uses AI and operates within the EU — or is an external company conducting business in the EU — it’s important to understand the new regulations and ensure adherence.

Non-compliance with certain provisions in the EU AI Act could result in fines of up to **€35 million EUR (\$38 million USD)** or up to **7% of your gross revenue**.

The new Act aims to ensure the safety and reliability of AI systems, positioning the EU as a leader in AI governance. It outlines obligations that include transparency, risk management, self-assessment, mitigation of systemic risks, and more.

In late 2025, the European Commission introduced the Digital Omnibus on AI, a set of targeted amendments then adopted by Parliament and Council that pushed back the compliance deadline for high-risk obligations while leaving most transparency requirements on their original schedule. The Digital Omnibus on AI (Regulation (EU) 2026/1744) became law on July 27, 2026. The dates below reflect the current, post-Omnibus timeline.

Dates to be aware of

February 2, 2025	AI literacy obligations under Article 4 took effect, as well as the original eight prohibitions under Article 5.
August 2, 2026	Article 50 transparency obligations generally apply (biometric-categorization/emotion-recognition disclosures, deep-fake labelling, public-interest text labelling). This date was not delayed.
December 2, 2026	Two further Article 5 prohibitions took effect covering AI-generated non-consensual intimate imagery and CSAM.
December 2, 2027	High-risk obligations apply to standalone Annex III AI systems (deferred 16 months from August 2, 2026).
August 2, 2028	High-risk obligations apply to AI embedded in regulated products under Annex I (deferred 12 months from August 2, 2027).

DNA of the EU AI Act

RECITALS

The recitals set the context and intent of the document. They explain the rationale behind the legislation or agreement, outlining the guiding principles and objectives.

Although not legally binding, they provide interpretative guidance. Courts or stakeholders often refer to recitals to understand the intent or spirit of the law.

ARTICLES & ANNEXES

There are **113 articles** within the EU AI Act, which are the primary operative provisions of the document. They define rights, obligations, procedures, and substantive rules. Articles form the core legal framework. They are binding and enforceable, and each article usually addresses a specific topic or rule. Articles are integral to the document and apply universally within its scope. They are mandatory unless explicitly stated otherwise.

The EU AI Act features **13 annexes**, which provide supplementary details, technical specifications, or additional information that support the articles. Annexes form an integral part of an EU legal act and have the same binding force as the articles. The distinction between articles and annexes is one of drafting convenience: annexes typically contain lists, technical specifications, tables or forms that would be unwieldy in the article text. Articles typically cross-refer to the relevant annex to give it operative effect.

SUPPORTING INSTRUMENTS & GUIDANCE

Beyond the articles, annexes, and recitals, the AI Act works through a layer of supporting instruments that do much of the practical compliance work.

Codes of Practice: Voluntary codes developed with the Commission and AI Office; signing up is a recognized way to demonstrate compliance. The Transparency Code (covering Article 50 transparency obligations) was finalized in June 2026.

Commission and AI Office guidance: Non-binding guidance interpreting how the Act applies in practice. Highly influential, and in the early years the main source of clarity before courts start ruling on the Act. Final Guidelines have been published on prohibited practices, the definition of an AI system, and Article 50 transparency obligations.

Harmonised standards: Technical standards developed by European standards bodies (CEN, CENELEC, ETSI) at the Commission's request. Using them is voluntary, but doing so creates a presumption that you meet the corresponding AI Act requirements. None have been finalized yet — the first likely to be published is prEN 18286 on quality management systems, expected in late 2026.

Classifications of AI systems

The EU AI Act categorizes systems according to the risk they can pose to users and consumers, with each having its own set of requirements and standards. The Act's primary regulatory focus for deployers covers **three main categories**:

1. Prohibited AI practices
2. High-risk AI systems
3. Systems requiring transparency

These categories are not mutually exclusive: the same system can fall into more than one depending on how and where it is used. Let's take a closer look.

1. PROHIBITED AI PRACTICES

Prohibited AI practices refer to the uses of AI that are considered to present an unacceptable risk to fundamental rights, safety, or EU values. These are explicitly banned under Article 5 and include, for example, AI systems that manipulate human behavior through subliminal techniques, exploit vulnerabilities of specific groups, enable social scoring by public authorities, or create facial recognition databases through untargeted scraping. Narrow exceptions exist, for example for medical or law enforcement purposes.

Certain uses of emotion recognition and real-time remote biometric identification in public spaces for law enforcement are also prohibited, except in narrowly defined circumstances.

These prohibitions took effect on February 2, 2025.

Important update: A further prohibition was **added to Article 5 by the 2026 Digital Omnibus on AI**: AI systems intended to generate non-consensual intimate imagery (“nudifiers”) or child sexual abuse material (“CSAM”), or that lack reasonable safeguards against producing such content, will be banned outright. This prohibition takes effect on December 2, 2026.

2. HIGH-RISK AI SYSTEMS

High-risk AI systems are defined in Article 6. These are systems that, due to their intended purpose or the sector in which they are deployed, pose a significant risk to health, safety, or fundamental rights if abused or deployed improperly.

High-risk AI systems are subject to strict requirements, including conformity assessments, risk management, data governance, transparency, human oversight, and post-market monitoring. AI systems regarded as high risk negatively affect safety or fundamental rights and are divided into two categories:

- **Annex I:** AI systems used in products falling under the EU’s product safety legislation, including toys, aviation, cars, medical devices, and lifts.
- **Annex III:** AI systems falling into specific areas that will have to be registered in the relevant database, including biometric identification and categorization, critical infrastructure, education and vocational training, employment and workers’ management, essential public and private services, law enforcement, and migration, asylum, and border control management. Some AI systems may be exempt if they do not pose a significant risk of harm or materially influence decision-making.

Important update: The 2026 Digital Omnibus on AI fixed these as unconditional deadlines: December 2, 2027 for standalone Annex III systems, August 2, 2028 for AI embedded in regulated products under Annex I. Unlike the Commission’s original proposal, these dates do not depend on technical standards being finalized, they apply regardless. Given the lead time needed for conformity assessments, human oversight processes, and logging infrastructure, treat this as implementation time, not a reprieve.

3. SYSTEMS REQUIRING TRANSPARENCY

For deployers, this category is defined by whether the system can expose an individual to AI without their knowing it. Article 50 identifies three such features:

- Inference from the individual's biometric data: emotion recognition and biometric categorization.
- Generation or manipulation of image, audio, or video content resembling real persons, objects, places, or events that would falsely appear authentic ('deep fakes').
- Generation or manipulation of text published to inform the public on matters of public interest.

Requirements for deployers of AI systems

Most organizations will be 'deployer'(s) of AI systems. Article 3 defines a deployer as a person or entity (business, non-profit organization, government body, etc.) using an AI system under its authority.

Organizations may have overlapping obligations as 'provider(s)' of AI systems, which are outside the scope of this guide.

DATA PROTECTION AND PRIVACY

Deployers must comply with all applicable Union data protection laws (such as GDPR) when using AI systems that process personal data (see Recital 10, Article 2(7)).

TRANSPARENCY OBLIGATIONS (Article 50) - Due Aug 2026

Unlike the high-risk deadlines, these obligations were not delayed by the 2026 Digital Omnibus and apply on the original schedule:

- Deployers of emotion recognition or biometric categorization systems must inform affected individuals that the system is in use.
- Deployers of AI systems that generate or manipulate AI-generated or manipulated image, audio, or video content ("deep fakes") must disclose that it was artificially generated or manipulated.
- Deployers publishing AI-generated or manipulated text to inform the public on matters of public interest must disclose that it was artificially generated. Limited exceptions apply, including for law enforcement, artistic works, and editorially reviewed content.

COPYRIGHT AND INTELLECTUAL PROPERTY

If a deployer uses outputs from an AI system for commercial or public purposes, they must ensure that such use does not infringe on copyright or related rights.

Use in high-risk contexts

'High-risk AI systems' (as defined above) are subject to strict regulatory requirements under the EU AI Act to ensure safety, transparency, and the protection of fundamental rights. This section outlines the key obligations that organizations must meet when deploying high-risk AI systems.

Important update: The 2026 Digital Omnibus on AI fixed unconditional deadlines for when the obligations below become mandatory: December 2, 2027 for standalone Annex III systems, and August 2, 2028 for high-risk AI embedded in regulated products under Annex I. Unlike the Commission's original proposal, these dates do not depend on technical standards being finalized, they apply regardless. Given the lead time needed for conformity assessments, human oversight processes, and logging infrastructure, treat the intervening period as implementation time, not a reprieve.

1. USE IN ACCORDANCE WITH INSTRUCTIONS

Deployers must take appropriate technical and organizational measures to ensure they use the system in accordance with the instructions for use supplied by the system provider.

2. HUMAN OVERSIGHT

Deployers must assign human oversight to natural persons who have the necessary competence, training, and authority, as well as the necessary support.

3. INPUT DATA CONTROL

If the deployer controls the input data, they must ensure it is relevant and sufficiently representative for the intended purpose of the high-risk AI system.

4. MONITORING AND REPORTING

- Deployers must monitor the operation of the high-risk AI system based on the instructions for use.
- If they believe the system presents a risk, they must inform the provider/distributor and the relevant market surveillance authority, and suspend use.
- Serious incidents must be reported immediately.

5. LOG RETENTION

Deployers must keep the logs automatically generated by the high-risk AI system (to the extent under their control) for at least six months, unless otherwise specified by law.

6. WORKPLACE NOTIFICATION

Before putting a high-risk AI system into service at the workplace, deployers who are employers must inform workers' representatives and affected workers.

7. REGISTRATION

Deployers that are public authorities or similar bodies must register themselves and the system in the relevant database before using a high-risk AI system.

8. DATA PROTECTION IMPACT ASSESSMENT

Where applicable, deployers must use the information provided by the provider to comply with their obligation to carry out a data protection impact assessment.

9. SPECIAL RULES FOR BIOMETRIC IDENTIFICATION

Deployers using high-risk AI systems for post-remote biometric identification for certain cases must request prior authorization from a judicial or administrative authority, except for initial identification based on objective facts directly linked to the offence.

10. TRANSPARENCY TO INDIVIDUALS

Deployers of Annex III systems must inform natural persons when they are subject to a high-risk AI system's decision-making, except for certain law enforcement uses.

Meeting EU AI Act requirements

As a deployer, you first need to understand what AI is deployed in your environment, how it is being used, what data it is accessing, and what data it is creating. You would want to have a grasp of these variables before the AI system is deployed, along with how they change over time.

An organization may also have a mix of high-risk AI systems and systems requiring transparency, which is why consistent processes may be more easily applied to both types of systems.

UNDERSTAND WHAT AI IS DEPLOYED AND WHERE

There needs to be a clear understanding for users and leaders within the organization on how to approve and onboard new AI systems.

Users should follow the process to register the solution for initial assessment, and then security teams need an AI governance platform or AI security platform for long-term log retention and compliance management.

The platform needs to include an inventory of all AI systems used across the organization and detect 'Shadow AI.' Without visibility into which AI systems are running, it's impossible to make informed decisions about which should or shouldn't be used in the organization.

Additionally, AI monitoring is critical for understanding how AI systems are used, and AI-SPM ensures that AI systems are secure and in compliance with the EU AI Act.

UNDERSTAND AI AND DATA USE

Nearly all AI systems process data, including personal data. GDPR continues to apply in full wherever personal data is involved, and the EU AI Act adds AI-specific obligations on top of it (see Recital 10, Article 2(7)). The EU AI Act also requires protections for copyright and intellectual property (IP), input data controls, and data protection impact assessments for deployers of high-risk AI systems.

In short, compliance with the EU AI Act is closely linked to data security and data governance.

For an end-to-end understanding of how data is being used within an AI system, you need a combination of both a **Data Security Platform** and an **AI Security Platform**.

DSP: A full DSP can see what users have access to data (with and without AI agency), along with what agents can access without human interaction.

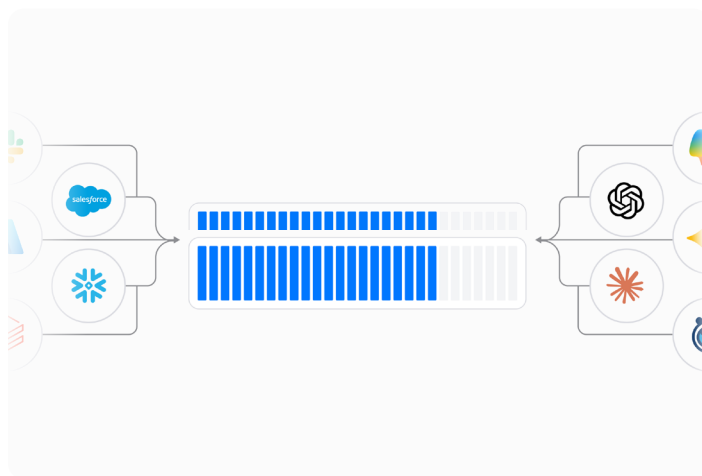
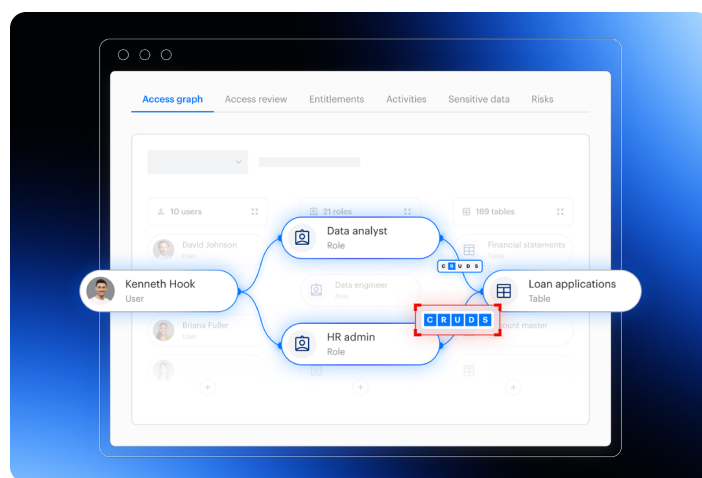
The leading AI systems follow the permissions of the identity using the AI application. Therefore, it's imperative to understand what access controls exist for each data set/resource and apply least privilege across them all.

To effectively report on incidents related to AI misuse or exploited vulnerabilities, you also need a way to alert and investigate on the day-to-day AI activities.

AI-SP: AI security and governance capabilities within these platforms play a crucial role in helping organizations comply with the EU AI Act's requirements for both high-risk AI systems and systems requiring transparency. Other similar solutions, like AI Security Posture Management (AI-SPM) solutions, fall into this category.

For all AI systems, Data Security Platforms and AI Security Platforms work together to maintain detailed records about what AI systems and agents exist and how they are connected to data. That includes training data and intended-use data, which are essential to demonstrating conformity with the Act's transparency and accountability standards.

For high-risk AI systems in particular, AI Security Platforms help organizations establish and maintain input-data relevance, monitoring, and log retention, all of which are mandated for deployers by the EU AI Act.



By automating documentation, monitoring, and reporting, platforms make it easier for providers and deployers to fulfill obligations such as incident reporting, log retention, and cooperation with authorities.

BUILDING COMPLIANCE AT EVERY LAYER OF THE AI AND DATA LIFECYCLE

Varonis provides the two most important platforms necessary to meet EU AI Act requirements along with other regulatory demands as they evolve:

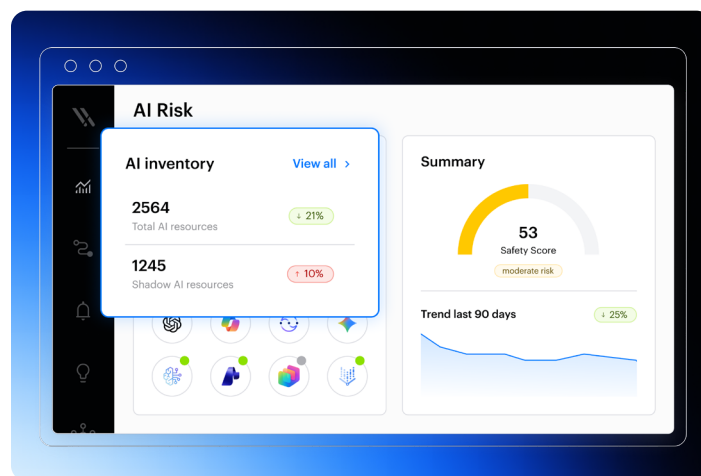
Data Security and AI Security.

The **Varonis Data Security Platform** covers structured, unstructured, and semi-structured data across domains. The platform was also named a **Leader in The Forrester Wave™: Data Security Platforms Q1 2025** and recognized across all three dimensions — with the strongest current Offering, strongest strategy, and the most recognition as a Customer Favorite.

When paired with **Varonis Atlas, an end-to-end AI Security Platform**, organizations have full visibility and control of regulated AI systems and the underlying data.

- **Visibility:** A complete inventory of the AI systems and tools running in the organization, including Shadow AI, mapped to the underlying data and applications that those AI systems can access, using agentless scanning that integrates with existing tools to minimize disruption.
- **Least Privilege:** A mapping of identities — human and agent — to data and AI systems across the organization with behavior analysis and data access governance. Automated remediation ensures least privilege is enforced, from removing stale users to right-sizing access.
- **Posture and Data-Centric Protection:** AI Security Posture Management (AI-SPM) scans for vulnerabilities and misconfigurations and enforces posture policies aligned to threat frameworks like MITRE ATLAS, and tracks every finding through to remediation.
- **AI and Data Monitoring:** Continuous monitoring of AI and data usage with real-time threat and intent detection to alert on non-compliant behavior, whether by humans or agents.
- **AI Red Teaming:** Automated LLM penetration testing and red-team validation probe your AI systems for prompt-injection, jailbreak, and data-leakage weaknesses before attackers find them.
- **Third-Party Compliance:** AI Third-Party Risk Management (TPRM) automates and manages supply chain AI risks. It replaces manual document reviews with automated surveys, analyzes AI vendor-inherited risks, and continuously monitors third-party AI.

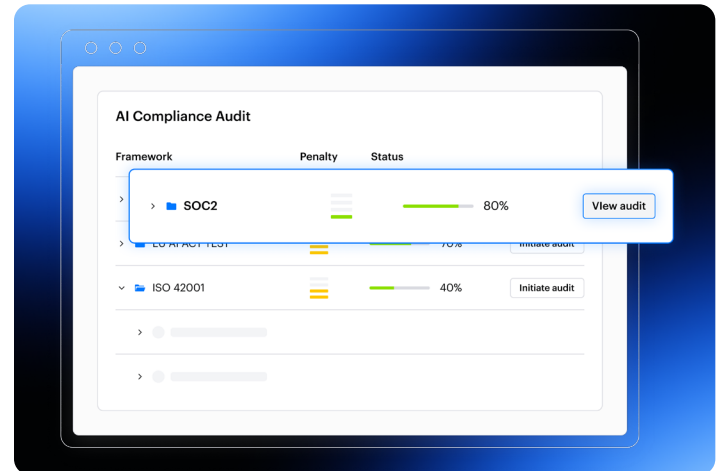
Beyond discovery and posture management, Varonis adds runtime protection and observability. It monitors all AI interactions, captures requests and responses under policy-based controls, and stores logs for compliance reporting.



Threat detection capabilities go beyond prompt injection, using contextual and historical analysis to identify suspicious activity. Organizations can also leverage built-in dashboards for governance and risk management, ensuring safe adoption of AI technologies while meeting EU AI Act requirements.

AUTOMATE YOUR EU AI ACT AUDIT

This is where security becomes compliance. Varonis Atlas includes an automated compliance module for the EU AI Act: it assesses each of your AI systems against the Act's requirements and positions you to compile and present compliance evidence based on what the platform already holds (AI inventory, access permissions, guardrail configuration and enforcement, monitoring activity, red-team results). This includes assistance with drafting compliance policies such as the quality management system required by Article 17.



The same evidence powers assessments for other frameworks such as ISO 42001 and the NIST AI RMF, so you gather the evidence once for overlapping requirements and reuse it.



See what others have to say

Varonis is the leader in data security and is trusted by more than 8,000 organizations to classify and protect their data. Here's what some of our customers said:

"Varonis allowed us to actually deploy AI. Without it, I don't think I would have been able to safely greenlight and recommend that we could. The only reason I'm prepared and feel safe in recommending our AI products through AI Governance Committee is because we have Varonis."

Jim Bowie, CISO
Tampa General Hospital

[Watch testimonial >](#)

"Varonis offers configuration options for global standards, including Europe and PCI compliance. Varonis' automation saves so many hours, days, and probably weeks over the year. I just cannot sing the praise is enough for the Varonis, and not just the product, but the service and the people in the business."

Infrastructure and Security Manager
Leading charity in the U.K.

[Read case study >](#)

Your partner on the frontlines of the agentic revolution

To ensure your organization is fully prepared for the evolving EU landscape of AI regulation, now is the time to take proactive steps toward compliance and security.

Engage with the Varonis team to access expert guidance, comprehensive risk assessments, and powerful platforms that deliver visibility, protection, automated compliance and actionable insights across your AI environment. Start your journey toward AI readiness today and empower your team to confidently navigate the requirements of the EU AI Act and beyond.

See Atlas in action

Below are the fast and easy steps to success the Varonis team will take to implement Atlas so your organization can trust autonomous systems to act safely, reliably, and in line with policy.

LEARN MORE

In a live demo, you'll see how to:

- Discover AI usage across models, agents + shadow AI
- Identify risky data exposure and permissions
- Enforce AI guardrails with real-time runtime protection
- Monitor AI activity with full data context

[Request a demo](#)

FREE TRIAL

Following the demo, take advantage of a full proof of value where you and your teams can access a tailored environment to match your current AI system architecture.

