



AI Act Service Desk

Frequently Asked Questions



This list of FAQs has been compiled based on queries received during the AI Pact webinars as well as submissions from stakeholders. This list will be updated regularly and as needed.

AI Act: General questions

What is the AI Act, and what are its objectives?

The EU AI Act is the world's first comprehensive AI law. It aims to promote innovation and uptake of AI, while ensuring a high level of protection of health, safety and fundamental rights, including democracy and the rule of law.

The uptake of AI systems has a strong potential to bring societal benefits, economic growth and enhance EU innovation and global competitiveness. However, in certain cases, the specific characteristics of certain AI systems may create risks related to user safety, including physical safety, and fundamental rights. Some powerful AI models that are widely used could also pose systemic risks.

This leads to legal uncertainty and potentially slower uptake of AI technologies by public authorities, businesses and citizens, due to the lack of trust. Disparate regulatory responses by national authorities could risk fragmenting the internal market.

Responding to these challenges, legislative action was needed to ensure a well-functioning internal market for AI systems and models where both benefits and risks are adequately addressed.

Related resources

- [Artificial Intelligence – Questions and Answers](#)

The AI Act applies progressively, with a full roll-out by 2 August 2027.

- The prohibitions, definitions and provisions related to AI literacy became applicable on 2 February 2025;
- The rules on governance and the obligations for general-purpose AI models became applicable on 2 August 2025;
- The obligations regarding high-risk systems listed in Annex III, the transparency requirements (Article 50) as well as the measures in support of innovation will apply as of 2 August 2026. This is also the date when the enforcement of AI Act will start;
- The obligations for high-risk AI systems that classify as high-risk because they are embedded in regulated products, listed in Annex I ([list of Union harmonisation legislation](#)), will enter into force on 2 August 2027.

Related resources

- [Timeline for the Implementation of the EU AI Act](#)

The AI Act is designed as a flexible and future-proof regulation that allows to adapt the rules to the rapid pace of technological development, as well as the potential changes in the use of AI systems and emerging risks.

While in general the AI Act can only be amended through the legislative procedure, in certain cases, the Commission is empowered to amend certain parts of the AI Act. For instance, the following parts of the AI Act can be adapted by the Commission:

- The list of high-risk use-cases in Annex III. The Commission is obliged to carry out a yearly review to assess if changes to the list are needed.
- The threshold above which general-purpose AI models are presumed to have high impact capabilities and are classified as presenting systemic risks.

The Commission also regularly assesses if other changes to the AI Act are needed and reports to the European Parliament and the Council. Such regular evaluation is foreseen directly in the AI Act.

What type of systems are regulated under the AI Act?

The AI Act does not apply to all AI solutions, but only to those that fulfil the definition of an 'AI system' within the meaning of Article 3(1) AI Act.

The AI Act follows a risk-based approach and introduces rules for AI systems based on the level of risk they can pose. Any AI practices with an unacceptable risk to health, safety or fundamental rights enshrined in the Charter of Fundamental Rights are prohibited (e.g. AI systems used to detect emotions of employees at work, except medical and safety reasons; certain social scoring practices). AI systems with high risk for health, safety or fundamental rights need to meet certain requirements to make sure they are safe and trustworthy (e.g. AI systems used at border control management; law enforcement, or autonomous vehicles could be examples of high-risk AI systems). Certain AI systems need to meet transparency requirements (e.g., deep fakes will have to be labelled as AI-generated; chatbots should inform that a person is not communicating with a human). All other AI systems remain unregulated and can be placed on the market, put into service or used in the EU without any requirements – at the time of preparation of the proposal of AI Act, it was estimated that these would be 85%.

Related resources

- [Artificial Intelligence – Questions and Answers](#)
- [The Commission publishes guidelines on AI system definition to facilitate the f...](#)

The AI Act does not automatically apply to all AI systems placed on the market before its application date. Instead, compliance obligations are phased in depending on the category and whether the system undergoes significant modifications.

The AI Act will apply from 31 December 2030 to AI systems which are components of the large-scale IT systems established by the legal acts listed in Annex X that have been placed on the market or put into service before 2 August 2027. In case such large-scale IT systems are evaluated or the legal acts in Annex X are replaced or amended before 31 December 2030, the AI Act shall be taken into account.

The AI Act will also apply to high-risk AI systems that have been placed on the market or put into service before 2 August 2026 only in case those systems are significantly modified. The AI Act will also apply from 31 December 2030 to high-risk AI systems that have been placed on the market or put into service before 2 August 2026 and are (intended to be) used by public authorities.

Lastly, the AI Act will apply from 2 August 2027 to general-purpose AI models that have been placed on the market before 2 August 2025.

However, the rules on prohibited AI practices (Article 5) apply to all AI systems, without taking into account the date of their placement on the market.

Focusing specifically on the building phase of the AI system, Recital 12 of the AI Act further clarifies that '[t]he techniques that enable inference while building an AI system include machine learning approaches that learn from data how to achieve certain objectives, and logic- and knowledge-based approaches that infer from encoded knowledge or symbolic representation of the task to be solved.' Those techniques should be understood as 'AI techniques'.

In addition to various machine learning approaches, that are often understood as an AI technique, the second category of AI techniques mentioned in Recital 12 of the AI Act are 'logic- and knowledge-based approaches that infer from encoded knowledge or symbolic representation of the task to be solved'. Instead of learning from data, these AI systems learn from knowledge including rules, facts and relationships encoded by human experts.

Based on the knowledge encoded by human experts, these systems can 'reason' via deductive or inductive engines or using operations such as sorting, searching, matching and chaining. By using logical inference to draw conclusions, such systems apply formal logic, predefined rules or ontologies to new situations. Logic and knowledge based approaches include, for instance, knowledge representation, inductive (logic) programming, knowledge bases, inference and deductive engines, (symbolic) reasoning, expert systems and search and optimisation methods.

The AI Act distinguishes between an AI system, defined in Article 3(1) and a general-purpose AI model defined in Article 3(63).

The AI Act refers to AI models as those that 'are essential components of AI systems. They do not constitute AI systems on their own. AI models require the addition of further components, such as for example a user interface, to become AI systems. AI models are typically integrated into and form part of AI systems' (Recital 97). 'Large generative AI models are a typical example for a general-purpose AI model, given that they allow for flexible generation of content, such as in the form of text, audio, images or video, that can readily accommodate a wide range of distinctive tasks' (Recital 99).

The AI Act defines an AI system as a 'machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments' (Article 3(1), Recital 12).

For more information, please refer to the [Guidelines on the definition of an artificial intelligence system](#).

Related resources

- [Artificial Intelligence – Questions and Answers](#)

What does the 'risk-based approach' mean?

The AI Act follows a risk-based approach, classifying AI systems into four different risk categories: unacceptable risk, high risk, transparency risk, and minimal to no risk.

For the unacceptable risk category, the AI Act lists specific practices that are prohibited (Article 5 of the AI Act). The high-risk AI systems are defined in accordance with Article 6 of the AI Act in conjunction with Annex I ([list of Union harmonisation legislation](#)) and Annex III of the AI Act. Annex III comprises eight areas in which the use of AI can be particularly sensitive and lists concrete use cases for each area that are assessed by the co-legislator as posing significant risks to health, safety and fundamental rights. For certain AI systems where transparency is particularly important, rules on transparency are foreseen. All the other AI systems are considered as "minimal to no risk" AI systems and the AI act does not foresee any obligations for them. Voluntarily, providers of those systems may choose to apply the requirements for trustworthy AI and adhere to voluntary codes of conduct.

What are the obligations for high-risk, transparency-risk, and minimal-risk AI systems?

Depending on the level of risk, obligations and requirements related to specific AI systems will vary.

Providers of high-risk AI systems need to ensure that their high-risk AI system complies with the requirements of the AI Act before the system is placed on the market or put into service. Examples of such requirements include having a risk management system in place, recording logs, and ensuring qualitative data and data governance, as well as human oversight. In addition, providers will need to register their system in the EU database and accompany the system with instructions for its use for deployers.

Providers of transparency-risk systems, such as chatbots, AI social companions or deep fakes, need to comply with certain transparency obligations. This includes, for example, informing natural persons that they are interacting with an AI system or ensuring that AI generated content can be detected as such. Deployers of systems that can generate or manipulate content need to disclose that the content is generated by AI.

The AI Act does not prescribe obligations for minimal-risk AI systems, but Member States can facilitate the drawing up of voluntary codes of conduct for AI systems that are not high-risk.

AI Act: General questions | High-Risk AI Systems
How should evolving AI systems be handled?

In principle, the evolution of an AI system does not have legal implications. However, in the case of high-risk AI systems, such developments should be anticipated and addressed within the risk management framework and related compliance obligations. Where an AI system undergoes substantial modifications, this may give rise to legal consequences, including the need for a renewed conformity assessment under the AI Act.

AI Act: General questions | High-Risk AI Systems
Is an "AI Officer" or governance board recommended within companies?

The AI Act does not require any particular internal governance within the company.

However in line with Article 17(1)(m) the **providers of high-risk AI systems** should put a quality management system in place, that should include an accountability framework, setting up the responsibilities of the management and other staff with regard to all the aspects related to the quality management system listed in Article 17 of the AI Act.

AI Act: General questions
How do I access AI-related funding programs?

The [EU Funding & Tenders Portal](#) of the European Commission is the central place to find any funding opportunities from the EU, including AI-relevant ones. Of particular interest among the funding programmes are Horizon Europe and Digital Europe Programme. Within Horizon Europe, there are different areas of interest, namely Cluster 4 and the European Innovation Council. Other areas such as the bottom-up European Research Council may be also relevant.

The GenAI4EU flagship initiative offers ample opportunities to develop and deploy trustworthy, generative AI in Europe's strategic sectors. With nearly 700 million euros committed, there are plenty of opportunities to help Europe become more competitive and innovative. A dedicated [website](#) gives an overview of the different call for proposals available under Horizon Europe and the Digital Europe Programme.

Related resources

- [EU Funding & Tenders Portal](#)
- [GenAI4EU: Funding opportunities to boost Generative AI "made in Europe"](#)

AI literacy (Article 4)

What is AI literacy as described in Article 4 of the AI Act?

The concept of AI literacy mentioned in Article 4 of the AI Act relies on the definition of the term given in Article 3(56) of the AI Act, according to which: 'AI literacy' means skills, knowledge and understanding that allow providers, deployers and affected persons, taking into account their respective rights and obligations in the context of this Regulation, to make an informed deployment of AI systems, as well as to gain awareness about the opportunities and risks of AI and possible harm it can cause.

Read the complete [AI Literacy FAQ here](#).

Related resources

- [AI Literacy - Questions & Answers](#)
- [Living repository to foster learning and exchange on AI literacy](#)
- [Third AI Pact webinar on AI literacy](#)

AI literacy (Article 4)

How can companies ensure AI competency (e.g., employee training)?

There is no single approach that works for all. It depends on the type of AI system that the company deals with and the knowledge of the staff in question. There is, however, no obligation for external training or external certification.

Please see further [AI Literacy - Questions & Answers](#).

Related resources

- [AI Literacy - Questions & Answers](#)

AI literacy (Article 4)

Will the Commission provide training materials or courses for AI Officers?

The AI Act does not require that companies designate AI Officers. However, the AI Act requires that providers and deployers of AI systems ensure, to their best extent, a sufficient level of AI literacy of their staff and other people dealing with the operation and use of AI systems on their behalf.

The Commission published a [living repository to foster learning and exchange on AI literacy](#), [Questions and Answers on AI literacy](#), and a [recorded webinar on AI literacy](#). The repository will be expanded. There are many cost-free online training courses. The choice depends on the individual skills of the staff in question.

Related resources

- [Living repository to foster learning and exchange on AI literacy](#)
- [AI Literacy - Questions & Answers](#)
- [Third AI Pact webinar on AI literacy](#)

Prohibited AI practices

What AI systems are prohibited?

The prohibitions of Article 5 AI Act include certain manipulative practices, social scoring, predictive policing based solely on profiling, scraping the internet and CCTV material to build-up or expand facial recognition databases, emotion recognition at education and workplace, biometric categorisation to infer sensitive attributes such as political opinion or sexual orientation and real-time remote biometric identification for law enforcement purposes at publicly accessible spaces (with some limited exceptions).

For more information, please refer to the [Guidelines on prohibited artificial intelligence practices](#).

Prohibited AI practices

What systems are prohibited under Article 5 of the AI Act (e.g., social scoring, emotion recognition)?

The prohibitions of Article 5 include certain manipulative practices, social scoring, predictive policing based solely on profiling, scraping the internet and CCTV material to build-up or expand facial recognition databases, emotion recognition at education and workplace, biometric categorisation to infer sensitive attributes such as political opinion or sexual orientation and real-time remote biometric identification for law enforcement purposes at publicly accessible spaces (with some limited exceptions).

Prohibited AI practices

Will guidelines be issued for prohibited AI systems?

[Guidelines were issued](#) in February 2025.

Related resources

- [Commission publishes the Guidelines on prohibited artificial intelligence \(AI\) ...](#)

High-Risk AI Systems

When will some guidelines on the classification of high-risk AI systems be published?

The Commission is currently preparing guidelines on the classification of high-risk AI systems, which should be published in February 2026.

High-Risk AI Systems | Prohibited AI practices

Is emotion recognition prohibited in workplace/education settings or for these purposes (e.g. employee monitoring)?

The placing on the market, the putting into service for this specific purpose, or the use of AI systems to infer emotions of a natural person in the areas of workplace and education institutions, is prohibited. However, this prohibition does not apply to the use of an AI system when the AI system is intended to be put in place or into the market for medical or safety reasons.

High-Risk AI Systems

What is the difference between the list of use cases and the list of areas of Annex III?

Annex III comprises eight areas and lists specific use-cases for each, which the EU co-legislator has assessed as posing a high risk of harm to health and safety or to fundamental rights.

In simple terms, the areas are enumerated as points in Annex III (e.g. point 1 – Biometrics, point 2 – Critical infrastructure, point 3 – Education and vocational training, etc.). The use-cases are listed under each area. For example, point 6 of Annex III covers the area of law enforcement, which includes five use-cases ((a)–(e)).

To be classified as high-risk, an AI system must fall under one of the areas of Annex III *and* correspond to one of the use-cases listed therein. In other words, not all AI systems used in a particular area (such as law enforcement) are automatically considered high-risk; only those falling within the specific use-cases identified in Annex III qualify.

High-Risk AI Systems

What documentation/assessments are needed for providers of high-risk AI systems to comply with the AI Act?

Article 16 of the AI Act gives an overview of the obligations that providers of high-risk AI systems need to follow in order to comply with the AI Act.

Prior to placing their high-risk AI system or putting it into services, providers of such systems need to ensure that their high-risk AI system is compliant with Articles 8-15 of the AI Act and has undergone a conformity assessment. This conformity should be demonstrated upon a reasonable request by a competent national authority.

Following this conformity assessment, providers need to draw up an EU declaration of conformity and affix the CE marking to the system (or on its packaging or accompanying documentation), as well as indicate their contact information on the high-risk AI system so they can be contacted. Providers also need to register their high-risk AI system in the EU database.

In addition, providers need to comply with Articles 17-20 AI Act by setting up a quality management system, keep documentation for 10 years, keep automatically generated logs by the high-risk AI system, and take necessary corrective actions and provide information related to these actions.

Finally, providers of high-risk AI systems need to ensure that their system complies with relevant accessibility requirements.

High-Risk AI Systems

Does location data qualify as biometric data under the AI Act?

Location data is normally not biometric. Accordingly, it does not fall within the high-risk use cases listed in point 1 of Annex III, nor within the prohibited practices set out in Article 5(1)(g).

High-Risk AI Systems

Are biometric systems (e.g., facial recognition for bank security) permitted?

Most biometric AI systems are not prohibited under the AI Act. The prohibited AI practices regarding biometric systems are limited and include in particular: emotion recognition in the workplace, certain biometric categorisation systems, and real-time remote biometric identification in publicly accessible spaces for law enforcement purposes. These are subject to narrowly defined exceptions.

There are, however, some additional rules that apply to permitted biometric systems. In particular, certain biometric applications are considered high-risk (remote biometric identification systems, certain AI systems intended for biometric categorisation, and AI systems intended for emotion recognition). The AI Act sets out specific requirements for high-risk AI systems. These requirements relate to, among others, data and data governance, documentation and record-keeping, transparency and the provision of information to users, human oversight, robustness, accuracy, and security.

The use of AI systems for biometric verification - that is, confirming that an individual is who they claim to be - is not prohibited under the AI Act and does not fall within the category of high-risk systems.

For more information, please refer to [Guidelines on prohibited artificial intelligence practices](#).

High-Risk AI Systems

What does the requirement for human oversight mean?

The EU AI Act emphasises the importance of human oversight for high-risk AI systems. This means that these systems must be designed and developed so that humans can monitor them while they are in use. Appropriate tools and measures should be in place to help people supervise AI effectively, ensuring that any risks to health, safety, or fundamental rights are prevented or minimised. When the system is in use, human oversight also includes appropriate measures to be implemented by the deployer. Those measures should be described in the instructions of use and be effective and commensurate with the risk, the level of autonomy of the system and the context of use of the high-risk AI system.

Deployers should implement human oversight measures and assign this task to natural persons who have the necessary competence, training and authority, as well as the necessary support.

High-Risk AI Systems

Are there simplified compliance paths for SME providers of high-risk AI systems?

The AI Act introduces a number of support measures tailored to small and medium-sized enterprises (SMEs). These include the possibility to prepare simplified technical documentation and to implement simplified quality management systems for high-risk AI systems. SMEs are also granted free-of-charge access to AI regulatory sandboxes, enabling them to test and develop AI systems in a controlled environment. In addition, the European Commission and national authorities are required under Article 62 to provide specific support measures to SMEs. Their interests are further represented through a special membership category in the advisory forum, ensuring that their needs are taken into account in the regulatory process. Finally, SMEs benefit from special consideration when penalties are imposed, reflecting their particular position in the market and the potential impact of sanctions.

What is a general-purpose AI (GPAI) system?

When a general-purpose AI model is integrated into or forms part of an AI system, the resulting system should be considered a general-purpose AI system if it gains the capability to serve multiple purposes through this integration. General-purpose AI systems can be used directly or embedded in other AI systems.

Given their versatility, general-purpose AI systems may be used as high-risk AI systems in their own right or as components of other high-risk AI systems. To ensure a fair distribution of responsibilities throughout the AI value chain, providers of general-purpose AI systems should cooperate closely with providers of relevant high-risk AI systems, unless otherwise specified in this Regulation. This cooperation should enable high-risk AI system providers to comply with their obligations under this Regulation and facilitate interactions with the competent authorities established by the Regulation.

Related resources

- [Questions and Answers on the Code of Practice for General-Purpose AI](#)
- [General-Purpose AI Models in the AI Act – Questions & Answers](#)

General-purpose AI models

How do AI agents fit within the GPAI framework under the AI Act?

While AI agents are not a separate category under the AI Act, they may have to comply with the requirements for AI systems and/or the obligations for providers of general-purpose AI models. Particularly relevant are the AI Act's prohibitions of harmful manipulation or exploitation of vulnerabilities, which can require technical safeguards in the design of an agentic AI system. If the agentic AI system classifies as high-risk, it is subject to additional requirements that ensure its safety and trustworthiness.

Furthermore, if the agentic AI system is intended to interact with natural persons or generate content, transparency rules apply (for further information, see [these FAQs](#) and [this consultation](#)). As regards general-purpose AI models, factors like the level of autonomy or tool use of a model can be decisive in its designation as having systemic risk. Providers of general-purpose AI models with systemic risk are subject to risk management obligations, which include considerations regarding the model's autonomous capabilities and its agentic use.

Given that the developments are recent and fast evolving, the European Commission's regulatory considerations are only preliminary at this stage, also in light of the fact that the understanding of what constitutes an AI agent varies greatly. The European Commission continues to closely monitor the development of AI agents and, if need be, will consider actions in that regard. For example, the AI Office's recent [call for tenders](#) on technical assistance for AI safety includes a dedicated lot to evaluating safety and security of AI agents.

General-purpose AI models

Which AI models does the AI Act apply to?

The AI Act regulates amongst others general-purpose AI models that are placed on the Union market, including their development and use. For further information under which conditions an AI model is considered 'general-purpose', see the question 'When does a model qualify as a general-purpose AI model?' and Section 2.1 of the [Guidelines on the scope of the obligations for general-purpose AI models](#).

'Placing on the market' means the first supply of an AI system or a general-purpose AI model for distribution or use on the Union market during a commercial activity, whether in return for payment or free of charge (Article 3(9)(10) AI Act). Some examples of what constitutes a placing on the market can be found in Section 3.1.2 of the [Guidelines on the scope of the obligations for general-purpose AI models](#).

An open-source release can constitute a placing on the market. Notably, the commercial supply of the model can be free of charge (Article 3(10) AI Act). However, providers of an open-source general-purpose AI model that does not present systemic risk are exempt from certain obligations. For further information, see the question 'How does the AI Act apply to general-purpose AI models released as open-source?' and Section 4 of the [Guidelines on the scope of the obligations for general-purpose AI models](#).

The internal use of a general-purpose AI model constitutes a placing on the market at least if it is essential for providing a product or service to third parties on the Union market or if it affects the rights of natural persons in the Union. Further examples are provided in Section 3.1.2 of the [Guidelines on the scope of the obligations for general-purpose AI models](#).

AI models that are specifically developed and put into service for the sole purpose of scientific research and development are out of scope of the AI Act (Article 2(6) AI Act).

Article 3(63) AI Act defines a 'general-purpose AI model' as 'an AI model, including where such an AI model is trained with a large amount of data using self-supervision at scale, that displays significant generality and is capable of competently performing a wide range of distinct tasks regardless of the way the model is placed on the market and that can be integrated into a variety of downstream systems or applications, except AI models that are used for research, development or prototyping activities before they are placed on the market'. This definition lists in a general manner factors that determine whether a model is a general-purpose AI model. Nevertheless, it does not set out specific criteria that potential providers can use to assess whether their model is a general-purpose AI model.

Section 2.1 of the [Guidelines on the scope of obligations for general-purpose AI models](#) provide an indicative criterion: **a model is considered a general-purpose AI model if the computational resources used for its training exceeds 10^{23} floating point operations and it can generate language (text or audio), text-to-image, or text-to-video**. This compute threshold corresponds to the typical amount of compute used to train models with one billion parameters on large datasets. Moreover, the chosen modalities enable flexible content generation capable of accommodating a wide range of distinct tasks. However, this is not an absolute rule — models meeting this criterion may exceptionally not qualify as general-purpose AI models if they lack significant generality, while models below this threshold may still be general-purpose AI models if they display significant generality and can competently perform a wide range of tasks.

A general-purpose AI model is classified as a general-purpose AI model with systemic risk if it meets either of the following two conditions.

First, if it has capabilities that match or exceed those of the most advanced models.

The AI Act presumes models trained with a cumulative amount of computational resources that exceeds 10^{25} floating point operations to have such capabilities (Article 51(2) AI Act). The threshold is set on the basis that the use of this amount of compute indicates that the model could significantly impact the Union market due to its reach or potential negative effects on public health, safety, security, fundamental rights, or society. The European Commission shall adjust this threshold when necessary to account for technological developments.

Second, the European Commission can designate a model as a general-purpose AI model with systemic risk, either through its own initiative or following a qualified alert from the scientific panel, if the model has capabilities or an impact equivalent to those of the most advanced models (Article 51(1)(b) AI Act). This accounts for the fact that there are models that may present systemic risks despite not meeting the compute threshold.

For a model meeting either condition, its provider must comply with additional obligations, including assessing and mitigating systemic risks (see the question 'What are the additional requirements for general-purpose AI models with systemic risk?').

Providers of models trained with less than 10^{25} floating point operations of cumulative compute are not expected to notify the European Commission of their model and comply with the safety and security obligations, unless their model was designated as presenting systemic risks by the European Commission.

General-purpose AI models

When should providers assess whether their model will be a general-purpose AI model and/or a general-purpose AI model with systemic risk?

Training a general-purpose AI model (with or without systemic risk) requires considerable upfront allocation of compute resources. Therefore, providers are able to know whether their model will meet the indicative criterion for constituting a general-purpose AI model, given in Section 2.1 of the [Guidelines on the scope of the obligations for general-purpose AI models](#), or the cumulative training compute threshold in Article 51(2) of the AI Act, before the training is complete (recital 112 AI Act). In particular, providers are expected to estimate whether their model will be a general-purpose AI model (trained with more than 10^{23} floating point operations of cumulative compute) and a general-purpose AI model with systemic risk (trained with more than 10^{25} floating point operations of cumulative compute) before the start of a large pre-training run.

Once a model exceeds the cumulative training compute threshold of 10^{25} floating point operations, set in Article 51(2) of the AI Act, or the provider knows that the model will exceed this threshold, the provider has two weeks to notify the European Commission (Article 52 AI Act). For further information on the notification see the question 'If a general-purpose AI model is above the compute threshold for such models with systemic risk, is it always considered to have systemic risk?.'

General-purpose AI models

What is the methodology for the cumulative training compute estimation (including synthetic data)?

Providers may choose their method to calculate/estimate the cumulative training compute in floating point operations. The [Guidelines on the scope of the obligations for general-purpose AI models](#), in their Annex, offer two possible approaches: architecture-based and hardware-based; as well as general principles for what should be included/excluded in the calculation. The compute used to generate synthetic data — artificially created training data produced by algorithms (including other models) rather than collected from real-world sources — must be included in the cumulative training compute calculation, additionally to the compute used to train on the synthetic data, if the synthetic data is not publicly accessible. However, the cumulative training compute of the parent model (i.e., the model generating the synthetic data) does not need to be included in the calculation/estimation.

General-purpose AI models

If a general-purpose AI model is above the compute threshold for such models with systemic risk, is it always considered to have systemic risk?

Providers may contest the automatic classification as a general-purpose AI model with systemic risk, when meeting the compute threshold in Article 51(2) of the AI Act. For instance, **they may rebut the presumption that their model has high-impact capabilities by providing evidence that their model's capabilities do not match or exceed those recorded in the most advanced models.** Section 2.3 of the [Guidelines on the scope of the obligations for general-purpose AI models](#) provides further guidance on how to do this. If the European Commission accepts the provider's arguments, the model will cease to be classified as systemic risk, thereby reducing the obligations for the provider with regard to this specific model from the moment of acceptance onwards.

In line with the collaborative and proportionate application of the AI Act, and the provider's right to good administration, the European Commission may take into account new information submitted by the provider after the notification, once it becomes available. Furthermore, the European Commission's decision to accept or reject the provider's arguments may take into account uncertainties in the information provided and outline conditions for potential future reassessment accordingly. However, the obligations for providers of general-purpose AI models with systemic risk apply from the moment the model meets the condition of Article 51(1)(a) of the AI Act (i.e., high-impact capabilities, presumed by meeting the compute threshold).

General-purpose AI models

Do built-in safety measures, such as output filters, affect the classification as a general-purpose AI model with systemic risk?

Generally, no. Safety measures may be relied upon to mitigate systemic risks stemming from the provider's general-purpose AI model with systemic risk. However, mitigations do not change the fact that the model presents (acceptable) systemic risk. If the model does not present systemic risks, even without implemented safety and security measures, providers may contest the automatic classification pursuant to Article 52(2) of the AI Act (see the question 'If a general-purpose AI model is above the compute threshold for such models with systemic risk, is it always considered to have systemic risk?').

The European Commission considers any subsequent development of the model downstream of its large pre-training run as part of the same model lifecycle and therefore part of the same model, if performed by the same provider or on behalf of the same provider (see Section 2.2 of the [Guidelines on the scope of the obligations for general-purpose AI models](#)). **Different considerations apply if another actor modifies the model** (see the question 'If a general-purpose AI model is modified by another actor than the original provider, does the modifying entity become the provider?' and Section 3.2 of the [Guidelines on the scope of the obligations for general-purpose AI models](#)).

This understanding of a model was chosen by the European Commission after careful consideration. In modern AI development, it is common to make small iterative changes to AI models after market placement. Providers also often perform more significant post-training to their model already placed on the market and the resulting altered model gets marketed under a new name, and colloquially being referred to as a 'new model'. As our understanding of the technology is still evolving, currently it is hard to define a good boundary in what counts as a relevant or irrelevant modification for the purpose of delineating subsequent model versions into distinct models. Therefore, for now, the European Commission considers every model version descended from the same large pre-training run by the same provider to be the same model to provide legal certainty. As a result, not every 'release' that is branded as such by providers is necessarily a new model for the purposes of the AI Act, following the interpretation in the guidelines. Non-public technical details about the training process may be needed to determine whether a 'release' constitutes a new model.

As models placed on the market before 2 August 2025 only need to be compliant with the AI Act starting from 2 August 2027 (Article 111(3) AI Act), it is possible that this transitional period still applies to a 'release' happening now.

General-purpose AI models

What entities or roles does the AI Act apply to when it comes to models?

Different to the provisions on AI systems, which cover multiple different roles such as the provider, deployer, importer, or distributor, the provisions on general-purpose AI models only regulate the 'provider'. A 'provider' in this context is whoever develops a general-purpose AI model or that has a general-purpose AI model developed, and places it on the market under its own name or trademark, whether for payment or free of charge (Article 3(3) AI Act). Importantly, 'places it on the market' refers to the **first** commercial supply of the model for distribution or use on the Union market (Article 3(9)(10) AI Act).

For example, an employer that provides its employees with licenses for a model already available on the Union market would neither be considered the provider of that model nor be required to verify its compliance under the AI Act. However, if a downstream actor integrates a model that was already made available on the Union market into its AI system, it may face obligations under the AI Act's provisions on AI systems.

General-purpose AI models

If a general-purpose AI model is modified by a different actor than the original provider, does the modifying entity become the provider?

The European Commission considers a downstream modifier to become the provider of the modified general-purpose AI model only if the modification leads to a significant change in the model's generality, capabilities, or systemic risk. An indicative criterion for this is that the training compute used for the modification is greater than a third of the training compute of the original model (see the Annex of the [Guidelines on the scope of the obligations for general-purpose AI models](#) for how 'training compute' should be understood). However, if the downstream modifier cannot be expected to know this value and cannot estimate it then the threshold should be replaced as follows.

If the original model is a general-purpose AI model with systemic risk, the threshold should be replaced with a third of the threshold for a model being presumed to have high-impact capabilities (i.e., currently 10^{25} floating point operations). Otherwise, it should be replaced with a third of the threshold for a model being presumed to be a general-purpose AI model (i.e., currently 10^{23} floating point operations, see Section 2.1 of the [Guidelines on the scope of the obligations for general-purpose AI models](#)).

This threshold is based on the expectation that a model that is modified with this amount of compute will display a significant change which warrants the downstream modifier being subject to the obligations for providers or general-purpose AI models and potentially of general-purpose AI models with systemic risk. While currently few modifications may meet this criterion, the number of downstream modifiers that become providers of general-purpose AI models may increase over time as the compute used to modify models increases. The criterion is thus primarily forward-looking, and in line with the risk-based approach of the AI Act. Therefore, the European Commission's approach may change in the future as technology and the market evolve.

For further information see Section 3.2 of the [Guidelines on the scope of the obligations for general-purpose AI models](#).

General-purpose AI models

Can entities contractually determine who the provider is?

The AI Act recognises that the entity that trains the model may not be the one that is contractually, economically, and/or technically in charge of the development (see the provider definition stating 'has developed' in Article 3(3) AI Act). The question of under whose authority a model is trained requires a case-by-case assessment, which can take account of contractual relations amongst other factors. Some examples are provided in Section 3.1.1 of the [Guidelines on the scope of the obligations for general-purpose AI models](#).

General-purpose AI models

What legal obligations do providers of general-purpose AI models face under the AI Act as of 2 August 2025?

Providers of general-purpose AI models must (Article 53 AI Act):

- draw up and maintain technical documentation about the model, including information about the development process, to provide to the AI Office upon request; national competent authorities can also ask the AI Office to request information on their behalf when this information is necessary for them to exercise their supervisory tasks;
- provide information and documentation to downstream AI system providers to help them understand the model's capabilities and limitations and comply with their own obligations;
- implement a policy to comply with Union copyright law and related rights, including identifying and respecting rights reservations through state-of-the-art technologies; publish a sufficiently detailed summary of the content used for training the model; if they are established outside the EU, appoint an authorised representative in the Union before placing their model on the market.

Providers can demonstrate compliance through the [General-Purpose AI Code of Practice](#) which was [assessed as adequate](#), or via alternative adequate means.

General-purpose AI models

What are the additional requirements for general-purpose AI models with systemic risk?

In addition to the standard obligations for providers of all general-purpose AI models, providers of general-purpose AI models with systemic risk must (Article 55 AI Act):

- perform model evaluation using standardised protocols and state-of-the-art tools, including conducting and documenting adversarial testing to identify and mitigate systemic risks;
- assess and mitigate possible systemic risks at Union level, including their sources, that may stem from the development, placing on market, or use of these models;
- track, document, and report relevant information about serious incidents and possible corrective measures to the AI Office and, as appropriate, national authorities without undue delay;
- ensure adequate cybersecurity protection for both the model and its physical infrastructure, to prevent unauthorised access, theft, or leakage.

Providers of such models can demonstrate compliance through adhering to the [General-Purpose AI Code of Practice](#) or show alternative adequate means of compliance. If providers choose to comply via alternative means, they must present arguments for why such means are adequate, for assessment by the European Commission.

General-purpose AI models

How does the AI Act apply to general-purpose AI models released as open-source?

Providers of general-purpose AI models released as open-source may be exempt from certain obligations (Articles 53(2), 54(6) AI Act), specifically:

- the requirement to maintain technical documentation for authorities;
- the requirement to provide documentation to downstream AI system providers;
- the requirement to appoint an authorised representative (for non-EU providers).

These exemptions require that the general-purpose AI model:

- is released under a free and open-source license, allowing access, use, modification, and distribution without monetisation;
- has its parameters, including weights, architecture, and usage information publicly available;
- is not a general-purpose AI model with systemic risk — providers of general-purpose AI models with systemic risk must comply with all the obligations for providers of general-purpose AI models, regardless of whether the model is released as open-source.

These exemptions recognise that open-source models contribute to research and innovation while already providing transparency through their open nature.

Nevertheless, providers whose models meet the above open-source requirements are not exempt from the copyright policy obligation or the requirement to publish a training data summary (Article 53(1)(c)(d) AI Act), since their open-source nature does not necessarily make available information on the data used for training or modifying the model, nor on how compliance with copyright law was ensured.

Section 4 of the European Commission [Guidelines on the scope of the obligations for general-purpose AI models](#) provides further guidance on the open-source exemption.

The [General-Purpose AI Code of Practice](#) is a voluntary tool, prepared by independent experts in a multi-stakeholder process, designed to help industry comply with the AI Act's obligations for providers of general-purpose AI models, ensuring that general-purpose AI models placed on the European market are safe and transparent, including the most powerful ones.

It outlines a way for providers of general-purpose AI models and of general-purpose AI models with systemic risk to demonstrate compliance with the AI Act's obligations laid down in Articles 53 and 55 AI Act.

The Code has three chapters: **Transparency** and **Copyright**, both addressing all providers of general-purpose AI models, and **Safety and Security**, relevant only to the limited number of providers of the most advanced models that are subject to the AI Act's obligations for providers of general-purpose AI models with systemic risk. Any provider (or potential future provider) of a general-purpose AI model can sign the code by completing the [Signatory Form](#) and sending it to EU-AIOFFICE-CODE-SIGNATURES@ec.europa.eu. The form should be signed by a person with sufficient authority to bind the provider to the Code (e.g. a senior executive). Further information about the signing process can be found [in these Q&A](#).

For further information on which AI models are 'general-purpose' and under which conditions they are classified as 'systemic risk', see the other questions and Section 2 of the [Guidelines on the scope of the obligations for general-purpose AI models](#). To learn more, read about the General-Purpose AI Code of Practice [in these Q&A](#).

The European Commission will regularly monitor and evaluate the achievement of the objectives of the [General-Purpose AI Code of Practice](#) (Code). The updating mechanism is further elaborated in the European Commission's [adequacy assessment](#) of the Code.

In particular, the European Commission will consider facilitating formal updates to the Code at least every two years, for instance based on the emergence of standards, relevant technological developments, or changes in the risk landscape.

To monitor the achievement of the objectives of the Code, the European Commission will remain in exchange with the signatories to understand where implementation support is necessary and may cooperate with national competent authorities, downstream providers, rightsholders, and other actors. Further, the Code leaves the European Commission's responsibility unaffected to issue guidance on the application of the AI Act, which may be of relevance for concepts in the Code. In particular, in the case of an imminent threat of large-scale irreversible harm or to address its negative effects, the European Commission will consider whether rapid guidance on the AI Act's application, or a rapid update to the Code agreed upon by the signatories, are appropriate, in addition to adequate enforcement actions. The [Guidelines on the scope of the obligations for general-purpose AI models](#) will be regularly reviewed and updated as appropriate.

Pursuant to Article 53(1)(a) and Annex XI, Section 1, point 2(e), of the AI Act, providers of general-purpose AI models must document the known or estimated energy consumption of their model. If it is unknown, this estimation may be based on the computational resources used. Moreover, the European Commission is empowered to adopt a delegated act to detail measurement and calculation methodologies that providers should use to measure or estimate the energy consumption of their models, to allow for comparable and verifiable documentation. For the period until such a delegated act is adopted, the Model Documentation Form in conjunction with the Transparency Chapter of the [General-Purpose AI Code of Practice](#) (Code) provides guidance on how providers may demonstrate compliance with this requirement, under 'Energy consumption (during training and inference)'.

In particular, if a Signatory to the Code does not know the energy consumption of training their model, they commit to reporting an estimated amount unless they are lacking critical information about compute or hardware which prevents them from being able to make an estimate. In this case, the Signatory commits to documenting the information they lack. For the purpose of estimating energy consumption, the AI Office will use the information available about the computational resources used for training to derive an estimate. To do so, the AI Office will rely on knowledge of what critical information the provider is lacking, available scientific resources, as well as on preliminary results from the ongoing European Commission study on energy-efficient and low-emission AI and draw on expertise from the scientific panel.

What are the obligations of the provider if a serious incident involving its model occurs?

Providers of general-purpose AI models with systemic risk need to keep track of, document, and report, without undue delay, to the AI Office and, as appropriate, to national competent authorities, relevant information about serious incidents and possible corrective measures to address them (Article 55(1)(c) AI Act). The European Commission considers that this obligation covers serious cybersecurity breaches related to the model or its physical infrastructure, including the (self-)exfiltration of model parameters and cyberattacks, due to their possible implications for the obligations provided for in Article 55(1)(b) and (d) AI Act. Apart from this, the European Commission considers a 'serious incident' in the context of Chapter V AI Act as any incident or malfunctioning of a general-purpose AI model that directly or indirectly leads to any of the events listed in the corresponding definition for AI systems in Article 3(49)(a) to (d) AI Act. The Safety and Security Chapter of the [General-Purpose AI Code of Practice](#), by way of its Commitment 9, provides a means to demonstrate compliance with this obligation.

General-purpose AI models

What does the European Commission expect from providers in the first year, before enforcement begins in August 2026?

The European Commission expects providers to comply with the obligations for providers of general-purpose AI models in the AI Act, which entered into application in August 2025. Some of these are straightforward — for instance, the notification to the AI Office when a provider is training a general-purpose AI model with systemic risk in Article 52 AI Act. Providers are expected to send these notifications without any further delay. Other obligations are more complex. For those, providers may demonstrate compliance through the [General-Purpose AI Code of Practice](#) (Code) or by demonstrating alternative adequate means of compliance.

The European Commission understands that, especially regarding the more complex obligations, it may not be possible for some providers to fully implement the measures in the Code, or similarly adequate means of demonstrating compliance, by August 2025. In cases where a provider is not fully compliant with all the Code's commitments right away, the European Commission will consider them to act in good faith and will be ready to collaborate to find ways to ensure full compliance. In particular, the European Commission encourages close informal cooperation with providers during the training of their general-purpose AI models with systemic risk to facilitate compliance and ensure timely market placement. Furthermore, the European Commission expects proactive reporting by providers of general-purpose AI models with systemic risk, whether as part of commitments under the Code or as part of alternative means of demonstrating compliance. This is particularly important during the first year, where the European Commission will focus on technical exchanges on working level to facilitate providers' full compliance as soon as possible.

Innovation measures

What are AI regulatory sandboxes, and how can providers/deployers participate?

AI regulatory sandboxes are safe and controlled environments for the experimentation, development, training and testing of innovative AI systems. Competent authorities supervising the sandboxes can provide guidance and advice on the interpretation of provisions of the AI Act and other relevant Union or national legislation. Any provider with an AI system that falls within the scope of the AI Act can apply for participation in an AI regulatory sandbox.

Member States should have their AI regulatory sandboxes in place by 2 August 2026, after which (prospective) providers can apply for participation in the sandboxes. At least one AI regulatory sandbox will be available in each Member State.

Innovation measures

When do Member States need to establish an AI regulatory sandbox?

Member States should have in place an AI regulatory sandbox with national coverage by 2 August 2026.

Innovation measures

What type of AI systems can be tested with an AI regulatory sandbox?

Any provider with an AI system that falls within the scope of the AI Act can apply for participation in an AI regulatory sandbox. This means that AI regulatory sandboxes can (among others) supervise specific AI projects to fulfil the requirements and obligations in the AI Act, including for conformity assessment of high-risk AI systems, compliance of interactive and generative AI systems with the transparency requirements in Article 50 of the AI Act, or to enable identification, testing and implementation of effective preventive and mitigating measures to ensure specific borderline cases of AI systems do not constitute prohibited practices.

Innovation measures

How will the AI regulatory sandboxes and the AI Act Service Desk interact with each other?

The AI regulatory sandboxes and the AI Act Service Desk are expected to complement each other's aim. The AI regulatory sandboxes are designed for more specific and complex regulatory questions and challenges, and the aim of the AI Act Service Desk is to provide answers to more general questions concerning the AI Act. The regulatory challenges dealt with, and questions answered in the AI regulatory sandboxes can also support the aim of the AI Act Service Desk by making use of the lessons learnt, the guidance and answers provided in the AI regulatory sandboxes.