



Guide AI Regulation

Version 1.0 - October 2024

Reading guide and disclaimer

You are developing AI systems or may want to implement them in your organization. If so, you may have to deal with the AI Regulation. This guide has been prepared for you, a tool to help you understand in an accessible way the main aspects of the AI Regulation to be obtained. **No rights can be derived from the contents of this guide.** The legal text of the AI regulation always remains leading.

Do you have feedback on this guide? If so, e-mail ai-verordening@minezk.nl. Your feedback will be used to improve this guide.

Are you reading a printed version of this guide? You can always find the latest version at ondernemersplein.nl/AIA.

The AI Regulation

The AI Regulation is a comprehensive law on artificial intelligence (AI) for the entire European Union (EU). The AI Regulation sets out rules for the responsible development and use of AI by companies, governments and other organizations with the aim of protecting people's safety, health and fundamental rights. This allows organizations to be confident that the AI they use is responsible and to take full advantage of the opportunities of AI.

The regulation goes into effect in stages and will be mostly applicable by mid-2027. Some AI systems are already prohibited as of February 2025. Therefore, it is smart to prepare now. To help you do this, this guide lists the most important provisions of the AI regulation. No rights can be derived from this document, which is purely for support. The full legal text can be found [here](#).¹

What does the AI regulation mean for your organization?

Depending on the AI system and what an organization does with that system, requirements are going to apply to its development and use. Whether requirements are going to apply depends, among other things, on the risk the AI system poses to safety, health and fundamental rights. There will also be different requirements for organizations that develop (or have AI developed) than for organizations that use AI. To find out what the AI regulation means for your organization, it is important to go through the following four steps. These steps are explained in the rest of the guide:

Step 1 (Risk): *Does our (AI) system fall within one of the risk categories?*

Step 2 (AI): *Is our system "AI" according to the AI regulation?*

Step 3 (Role): *Are we the provider or usage manager of the AI system?*

Step 4 (Obligations): *What obligations must we adhere to?*

¹ <https://eur-lex.europa.eu/legal-content/NL/TXT/?uri=CELEX:32024R1689>

Step 1 (Risk) Does our (AI) system fall within one of the risk categories?

All AI systems are covered by the AI regulation, but there are different requirements for different categories based on risk. This risk is determined by the intended **application or product** for which the AI system is developed, sold and used:

- **Prohibited AI:** these AI systems may not be sold, used or put into operation.
- **High-risk AI:** These AI systems must meet several requirements to mitigate risk before they can be sold or used.

In addition, there are going to be requirements for AI models and AI systems that **can do** specific things:

- **General purpose AI models and systems:** specific information requirements will apply to these. In some cases, risk mitigation requirements must also be met.
- **Generative AI and chatbots:** here, specific transparency requirements are going to apply regardless of whether the system is high-risk or not.²

Sometimes the same AI system can fall under multiple categories. For example, a chatbot may be used for a high-risk application. AI systems that do not fall under any of the above categories do not have to meet the requirements of the AI Regulation. However, these systems may still have to meet requirements from other regulations, such as the General Data Protection Regulation (AVG).

To know whether you need to comply with requirements of the AI regulation, it is important to first determine under which category your AI system falls. The various risk categories are explained in detail below.

1.1. Prohibited AI systems

Certain AI practices pose an unacceptable risk to people and society and are therefore banned from February 2025. This means that these systems may not be marketed, used or put into service. These bans apply to both **providers** and **those responsible for use** (explained further under [Step 3. Are we the provider or use manager of the AI system? on page 10](#)).

² Specific transparency requirements will also apply to emotion recognition and biometric categorization systems. Because these are also high-risk AI systems, these requirements are included under the high-risk obligations in step 4.2.

Prohibited AI systems

1. Systems designed to **manipulate human behavior** to restrict the free choice of individuals and which may result in significant harm to those individuals.
2. Systems that **take advantage of the vulnerabilities** of individuals due to their age, disability or specific social or economic situation and can lead to significant harm to those individuals.
3. Systems for establishing rewarding and punitive point systems based on social behavior or personal characteristics, known as **social scoring**, resulting in adverse and unfair treatment.
4. Ban on **risk assessment** systems **for committing crimes** based solely on profiling or (personality) characteristics.
5. Systems that create or expand **facial recognition databases** through the indiscriminate **scrapping** of facial images from the Internet or CCTV images.
6. **Emotion recognition** systems in the workplace and education, unless done for medical or safety reasons.
7. Systems used to **categorize people based on biometrics** in certain sensitive categories.
8. The **use of remote real-time biometric identification in public places for law enforcement purposes**. Some exceptions apply in cases where the use is strictly necessary, such as for the search for specific victims of kidnapping and human trafficking or missing persons. Additional safeguards do apply to these uses.

1.2. High-risk AI systems

High-risk AI systems may have risks to the health, safety or fundamental rights of individuals, such as the right to privacy and not to be discriminated against. At the same time, these systems can also have applications with positive effects on people and organizations, if they are reliable and the risks are mitigated. Therefore, starting August 2026, high-risk AI systems must meet several requirements before they can be marketed, used or commissioned. This means that **providers** must ensure during the development of the system that it meets the requirements before it is first sold or used. A professional party who uses the AI system under their own responsibility is a **use representative** (explained further under [Step 3. Are we the provider or use manager of the AI system? on page 10](#)).

For use representatives

obligations also apply, to mitigate risks resulting from the specific deployment of the system. There

are two types of high-risk AI systems:

- **High-risk products:** AI systems that are also directly or indirectly covered by a selection of **existing product regulations** (see below). For example, an AI system as a safety component of an elevator or an AI system that is a medical device.
- **High-risk applications:** AI systems developed and deployed for particular applications in "high-risk application areas. These are eight application areas ranging from AI for law enforcement to AI in education. Included within those eight areas are about 30 different specific applications that may have high risks, such as AI systems to help dispatch emergency services.

The product groups and application areas in which AI systems are classified as high-risk are shown in the figures below.

The obligations for this category are described under [4.2. High-risk AI on page 11](#).

High-risk AI as a (safety component of) existing products

These are products that are already regulated in the EU. A product is considered high-risk when existing product regulations require third-party approval before it can be placed on the market (conformity assessment). If AI is a safety component of the risk product or if the risk product itself is an AI system, then it is considered high-risk AI. This applies to products covered by the following product legislation:

- **Machinery** (Directive 2006/42/EC)
- **Toys** (Directive 2009/48/EC)
- **Pleasure boating** (Directive 2013/53/EU).
- **Elevators** (Directive 2014/33/EU)
- **Equipment and protective systems for use in potentially explosive atmospheres** (Directive 2014/34/EU)
- **Radio equipment** (Directive 2014/53/EU)
- **Pressure equipment** (Directive 2014/68/EU)
- **Cableway installations** (Regulation (EU) 2016/424)
- **Personal protective equipment** (Regulation (EU) 2016/425)
- **Gas burning appliances** (Regulation (EU) 2016/425)
- **Medical devices** (Regulation (EU) 2017/745)
- **In vitro diagnostic medical devices** (Regulation (EU) 2017/746)

In addition, the AI Regulation mentions another list of products that are also considered high-risk AI, but for which **no direct requirements will apply from the AI Regulation**. However, the requirements of the AI Regulation will be used at a later date to implement the specific product legislation that applies to these products. When this will happen is not yet known and will vary per product. It concerns the products covered by the following product legislation:

- **(Security of) civil aviation** (Regulation (EC) 300/2008 and Regulation (EU) 2018/1139)
- **Two or three-wheeled vehicles and quadricycles** (Regulation (EU) 168/2013)
- **Agricultural and forestry vehicles** (Regulation (EU) 167/2013)
- **Marine equipment** (Directive 2014/90/EU).
- **Rail system interoperability in the EU** (Directive (EU) 2016/797)
- **Motor vehicles and trailers** (Regulation (EU) 2018/858 and Regulation (EU) 2019/2144)

High-risk application areas

An AI system falls under one of the application areas if the provider has intended the use of the AI system for an application within one of these areas. The provider must explicitly include the purpose in the AI system's documents, including the instructions for use, advertising materials and further technical documentation.

There are eight high-risk application areas. This does not mean that all AI systems that fall within that, often abstractly defined, application area are high-risk. Some specific applications are listed for each area.

Tip: First check if your AI system falls under one of the eight areas, and then see if your AI system is one of the described AI systems in that category. Only in that case are you dealing with a high-risk AI system that needs to start meeting requirements.

1. Biometrics

- Remote biometric identification systems, unless the system is used for verification *only*.
- Systems for biometric categorization based on sensitive traits or characteristics.
- Emotion recognition systems.

2. Critical infrastructure

- Systems as security components for management and operation of critical digital infrastructure, for road traffic or for supply of water, gas, heating and electricity.

3. Education and vocational training

- Educational admission or assignment systems.
- Systems for evaluating learning outcomes.
- Systems for assessing educational attainment.
- Systems for monitoring students during exams.

4. Employment, workforce management and access to self-employment

- Systems for recruiting or selecting candidates.
- Systems for making decisions affecting the contract of employment and working conditions, for allocating duties, and for monitoring and evaluating employees.

5. Essential private and public services and benefits

- Systems for assessing (the degree of) access to essential government benefits and services.
- Systems for assessing the creditworthiness or score of individuals, unless done to detect financial fraud.
- Risk assessment and pricing systems for life and health insurance.
- Systems for prioritizing the deployment of emergency services or patient triage in healthcare.

6. Law Enforcement

- Law enforcement systems to determine the likelihood of a person becoming a victim of crime.
- Law enforcement systems deployed as lie detectors.
- Law enforcement systems to assess reliability of evidence.
- Law enforcement systems to assess or predict the likelihood that a person will commit a criminal offense or to assess previous criminal behavior of (groups of) individuals.
- Law enforcement systems for profiling individuals during investigation or prosecution.

7. Migration, asylum and border management

- Systems for government agencies to use as lie detectors.
- Systems for government agencies to assess risks to security, to illegal migration or a health risk upon entry to a country.
- Systems to assist government agencies in dealing with asylum, visa or residency issues, including related complaints.
- Government agency systems for detecting, recognizing or identifying individuals, excluding travel document verification.

8. Justice and democratic processes

- Systems to support the administration of justice and litigation to investigate and explain facts or to explain or apply the law to a set of facts.
- Systems for influencing elections or referendums or the voting behavior of individuals, with the exception of supporting political campaigns from an administrative or logistical point of view.

Exceptions to the high-risk application areas

There are some specific **exceptions** where AI systems that fall under one of the application areas are nevertheless not considered high-risk AI. This is the case when the AI system **does not materially affect a decision** because the system is intended to:

- Perform a **limited procedural task**;
- **Improve** the result of a **previously completed human activity**;
- To **check whether human decisions deviate** from a previous pattern without replacing or influencing the decision;
- Perform a **preparatory task** for an assessment relevant to one of the high-risk application areas.

If you have determined that your AI system falls under one of the exceptions, you must document this and register the AI system in the EU database for high-risk AI systems. The European Commission will provide a list of examples at a later date to clarify what does and does not fall under the exceptions.

1.3. AI models and general purpose AI systems

A **general purpose AI model**, also known as "General Purpose AI," can properly perform a wide range of different tasks and can therefore be integrated into a variety of different AI systems. Often, these models are trained on large amounts of data and using *self-supervision* techniques.

The broad applicability of these models, through specific AI systems, allows them to be used for all kinds of applications. These can include high-risk applications. Because of the potentially high impact of these models, they must meet several requirements as of August 2025.

If an AI system is based on a general purpose AI model and can actually serve different purposes itself, then it is a **general purpose AI system**.

The obligations for this category are described under [4.3. General purpose AI models and systems on page 16](#).

1.4. Generative AI and chatbots

Transparency requirements are imposed on some AI systems. These are systems that people often interact with directly. Therefore, it must be clear to people that AI is interacting with them or that the content has been manipulated or generated.

- Systems used to generate audio, image, video or text (**generative AI**);
- Systems made for interaction (**chatbots**).

The obligations for this category are described under [4.4. Generative AI and Chatbots on page 17](#).

1.5. Other AI

See [4.5. Other AI on page 17](#) for more information on AI systems that do not fall into any of the risk categories described above.

Step 2. Is our system "AI" according to the AI regulation?

The AI regulation sets requirements for AI systems. There are different images about what "AI" is or is not. The AI regulation contains the following definition that is intended to delineate what AI is as a product in the marketplace:

*"An AI system is a machine-based system that is designed to operate **with different levels of autonomy** and can exhibit **adaptability** after deployment, and that, for **explicit or implicit objectives**, derives from the received inputs how **to generate outputs** such as predictions, content, recommendations or decisions."*

What does this include?

- Systems that use machine learning (**machine learning**) where data is used to learn how to achieve certain objectives;
- Systems using **knowledge and logic-based** approaches that enable learning, reasoning or modeling.

What is **not included** here?

- Systems based on rules established solely by humans to perform actions automatically.

If your system is not considered "AI" under the AI Regulation but does fall into one of the risk categories, it is important to have the conversation within your organization as to what extent the system does not still pose risks and whether these risks can be mitigated by meeting (certain) requirements from the

AI regulation to comply. In addition, systems that fall outside the AI regulation may have to comply with requirements from other regulations.

Step 3. Are we the provider or user manager of the AI system?

After you have determined which risk category your AI system falls under and whether your AI system is also covered by the AI regulation, you must determine whether you are a provider or a user agent.

- **Provider (provider)** means a person or organization that develops or commissions the development of an AI system or model and markets or self-deploys that system.
- **User responsible (deployer)**: a person or organization that uses an AI system under its own responsibility. It does not include non-professional use.

The description of the requirements in Step 4 describes, for each risk category, the obligations that apply to providers and those responsible for use. They must each meet different obligations, with the most onerous obligations applying to providers.

Note: As a use case manager, in some cases you may also become a provider of a high-risk AI system and you must meet the high-risk obligations for providers. This is explained further under steps [4.2. High-risk AI on page 11](#) and [4.3. AI models and systems for general purposes on page 16](#).

Note that there are also other roles under the AI regulation, for example, agents, importers and distributors. Obligations for these actors are not addressed in this guide.

Step 4. What obligations must we abide by?

4.1. Forbidden AI

These AI systems pose an unacceptable risk and are therefore banned from February 2, 2025. This means that these systems may not be marketed or used. These bans apply to both **providers** and **users**.

The prohibition of real-time remote biometric identification in public places for law enforcement purposes is subject to sharply defined exceptions, and their deployment must have a national legal basis. Additional safeguards also apply around the deployment of these systems.

4.2. High-risk AI

The most onerous obligations of the AI regulation will apply to high-risk AI systems. **Providers** must meet (various) obligations:

- Risk management system;
- Data and data governance;
- Technical Documentation;
- Logging;
- Transparency and information;
- Human supervision;
- Accuracy, robustness and cybersecurity;
- Quality Management System;
- Monitoring.

As a provider, if you meet or think you meet all of these obligations, you will need to conduct a **conformity assessment**. In some cases you may do this yourself and in some cases a third party must do it for you. A later version of this guide will address when you must perform which procedure.

Users must also meet various obligations, with additional obligations applying to government organizations using AI systems.

The figure below explains each obligation. These obligations will be further developed into European standards in the coming years, in which everyone can participate through the standardization organizations in the European member states. In the Netherlands, this is the NEN.³

Note: In two cases, as the user manager of a high-risk AI system, you can become the provider of that system yourself:

- When you, as a user manager, put your own name or brand on the high-risk system;
- When you, as the person responsible for use, make a significant change to the high-risk AI system that was not anticipated by the provider and causes the system to no longer meet the requirements or changes the purpose of the system as intended by the provider.

³ <https://www.nen.nl/ict/digitale-ehetiek-en-veiligheid/ai>

Rules for providers of high-risk AI systems

1. Risk management system

Several steps must be taken for this system:

- Identify and analyze foreseeable risks to health, safety or fundamental rights, among others.
- Taking appropriate risk management measures that ensure that risks remaining after the measures *are acceptable*.

The following points should be taken into account:

- The identification and addressing of risks should take place once before the AI system is marketed or used and then continuously during the use of the AI system.
- Allowance must be made for foreseeable misuse of the system.
- Consideration should be given to the context of use, including the use manager's knowledge and experience with such AI systems or whether children or vulnerable groups are affected by the AI system. For example, it may be necessary to offer training to those who will be working with the AI system.
- Risk management measures should be tested to verify that they actually work. This should be done using benchmarks appropriate to the purpose for which the AI system is being deployed.
- If a risk management system must also be established under existing product legislation, it may be combined into a single risk management system.

2. Data and data governance

There are different requirements for the datasets used to train, validate and test high-risk AI systems.

- Data management appropriate to the purpose of the AI system, including:
 - Recording processes, including those of data collection and processing;
 - Establishing assumptions about the datasets;
 - An assessment of the availability, quantity and adequacy of the datasets, including possible biases that may affect people;
 - Measures to detect, prevent and mitigate biases;
 - Addressing deficiencies in the datasets that may impede compliance with the AI regulation (e.g., consider mitigating risks under the risk management system).
- Datasets should be sufficiently representative and as error-free as possible for the purpose for which they will be used. This should also take into account the context in which the AI system will be used; for example, the geographical or social context.
- Under a number of strict conditions, special categories of personal data (a concept from the General Data Protection Regulation) may be processed to counteract biases in datasets.

3. Technical Documentation

The technical documentation must demonstrate that the high-risk AI system meets the requirements of the AI Regulation. The technical documentation must include, among other things:

- A general description of the AI system, including the intended purpose of the system, the name of the provider, and instructions for use;
- A detailed description of the elements of the AI system and the process for its development, including the steps of development, the design choices, the expected output of the system, the risk management system and the datasets used.

- Detailed information on the monitoring, operation and control of the AI system, including the level of accuracy at individual and overall levels, risks, the system for in-service evaluation, and measures for monitoring and human supervision.
- An overview of the standards used.
- The EU declaration of conformity (the CE mark).

SMEs can record technical documentation in a simplified way. The European Commission will provide a form for this at a later date.

4. Logging

Automatic logs should be kept during the life of the AI system to allow timely detection of risks and monitoring of system operation.

Logs must be kept for at least six months. At least the following must be logged:

- The duration of each use of the AI system;
- The input data and its control by the AI system (and reference database);
- The identification of those involved in the verification of results.

5. Transparency and information

The provider of the AI system knows how the system works and how to handle it. It must therefore ensure that the AI system is transparent enough that those responsible for using it understand how to properly use its output.

This requires the preparation of **instructions for use** that include the following:

- Contact details;
- The purpose, characteristics, capabilities and performance limitations of the AI system;
- Human supervision measures.

6. Human supervision

High-risk AI systems must be designed by the provider to allow human supervision during use, thereby reducing risks to individuals.

Human supervision is context-dependent - the greater the risks the stronger the supervision must be. Supervisory measures can be technical in nature (for example, a clear human-machine interface), or measures to be implemented by those responsible for use (for example, mandatory training for their staff).

Ultimately, the measures should ensure that those who will use the AI system can do the following:

- Understand system capabilities and monitor operation;
- Be aware of "automation bias.
- Correctly interpret the output and ignore or replace if necessary;
- Shutting down the system.

7. Accuracy, robustness and cyber security

High-risk AI systems must provide an appropriate level of accuracy, robustness and cybersecurity. Benchmarks and metrics are being developed for this purpose, including by the European Commission.

At least the following measures are mentioned for this purpose:

- Technical and organizational measures to prevent errors arising from the use of the AI system by individuals;
- Robustness solutions, such as backups or security measures in case of failure;
- Remove or reduce negative influence on the system by limiting feedback loops;
- Cyber security that prevents unwanted third-party access by tracking, responding to and resolving attacks. These include attacks aimed at data contamination, model contamination, modifying inputs or obtaining confidential data.

8. Quality management system

The quality management system must ensure compliance with the requirements of the AI regulation. How extensive this system of quality management should be depends on the size of the organization. Among other things, by establishing :

- A strategy for compliance;
- Techniques, procedures and measures for the design, development and quality control of the AI system;
- Using or not using standardization;
- Systems and procedures for data management, risk management, monitoring, incident reporting and documentation.

9. Monitoring

Once an AI system is marketed or in use, providers must monitor the system based on usage data to use this to verify that the system continues to meet the requirements of the AI regulation. To do this, providers must establish a plan for monitoring.

If the provider of the high-risk AI system learns that the system is no longer operating in accordance with the AI regulation, corrective action must be taken immediately to fix it. This may even include recalling the system if necessary. Also, the provider must cooperate with usage managers and inform regulators about this.

Serious incidents involving the AI system must be reported to the supervisor(s).

Other requirements

- The registration of the high-risk AI system in the EU database.
- The provider's contact information must be included with the AI system;
- Technical documentation, quality management documentation and conformity assessment documentation must be retained for 10 years.

Rules for persons responsible for use of high-risk AI systems

Not only providers, but also those responsible for using high-risk AI systems must start meeting requirements. After all, they are the ones who have control over how the AI system is used in practice and thus have a major impact on the risks that may arise.

Usage managers must:

- Take technical and organizational measures to ensure that the high-risk AI system is used according to the instructions for use;
- Ensure that those overseeing the system have appropriate knowledge, skill and authority;
- Ensuring that, to the extent possible, input data is sufficiently relevant and representative;
- Monitor the operation of the AI system based on the instructions for use;
- If a user manager assumes that the system no longer meets the requirements of the AI regulation, he or she should notify the provider and suspend use of the system;
- Notify the provider and supervisor(s) of potential risks and serious incidents that have occurred;
- Retain the logs over which they have control for at least six months;
- Inform employee representation if the AI system is deployed in the workplace;
- When decisions are made about people using the high-risk AI system, inform these people about it.
- If AI is used for emotion recognition or biometric categorization, the people to whom it is applied should be informed about it.

Specific obligations for government organizations as use cases Government organizations must meet a number of additional obligations in addition to those listed above:

- Registering the use of a high-risk system in the EU database;
- Assess what the impact on fundamental rights may be when using the high-risk AI system, taking into account the specific context in which the use takes place (a **'fundamental rights impact assessment'**). For example, they look at the duration of use, the processes within which the system is used, and the impact the use may have on the fundamental rights of individuals and groups. After identifying the risks, use managers must take measures for human supervision and dealing with the potential risks. A notification must also be made to the market regulator.
Note: This requirement also applies to individuals providing public services, the use of AI systems for assessing people's creditworthiness, and AI systems for risk assessments for life and health insurance.

4.3. AI models and general purpose systems

Obligations for general purpose AI model providers

General purpose AI models can be integrated into a variety of different AI systems. This does require that the providers of these AI systems know what the AI models can and cannot. There are also specific requirements for training these models, as they often use large data sets. The providers of these models must:

- Prepare technical documentation including the training and testing process and its results and evaluation;
- Prepare and update information and documentation for AI system providers who wish to integrate the model into their AI system. The information should provide an understanding of the capabilities and limitations of the AI model and enable the AI system provider to self-fulfill the obligations of the AI Regulation.
- Establish policies to ensure that they train the model without infringing on the copyrights of people and organizations;
- Prepare and make public a sufficiently detailed summary about the content used to train the AI model.

Providers of open-source models do not have to meet these obligations. The AI regulation determines when a model is "truly" open-source.

Obligations for general purpose AI model providers with systemic risks

In some cases, general purpose AI models can create systemic risks. This is the case when the model has high-impact capabilities. In any case, this is assumed to be the case when at least 10^{25} floating-point operations (FLOPs) have been used to train the model. In addition, based on certain criteria, the European Commission may determine that the model has a similarly high impact in another way. These models must:

- Meet the obligations for general purpose AI models;
- Conduct model evaluations to identify systemic risks;
- Reducing systemic risk;
- Maintain and report information on serious incidents to the AI Office;
- Ensure appropriate cybersecurity.

Note that these obligations apply only to the largest AI models.

Providers of these models with systemic risks cannot rely on an exception for open source.

What rights do you have if you integrate a general purpose AI model into your (high-risk) AI system? As also indicated above, you should at least receive information and documentation so that you can make your own assessment of how to use the model in your AI system for your chosen purpose. If you incorporate the model into a high-risk AI system, you must then still comply with the obligations of the AI regulation as a provider.

How should you handle general purpose AI systems? As indicated under [1.3. AI models and general purpose AI systems on page 8](#), there are also AI systems that are

can serve various purposes. Consider, for example, the well-known AI chatbots. **Note:** If you deploy these systems for high-risk purposes, according to the AI regulation, you yourself become the provider of a high-risk AI system. You will then have to meet the associated obligations. In this situation, it is very difficult to meet the obligations for a high-risk AI system, putting you at risk of receiving a fine.

4.4. Generative AI and Chatbots

To ensure that people know whether they are talking to an AI system or seeing content generated by AI, there are transparency requirements for generative AI and chatbots.

Rules for chatbot providers

Providers of systems made for direct interaction with people should ensure that these people are informed that they are dealing with an AI system.

Rules for providers of generative AI

Providers of systems that generate audio, image, video or text must ensure that the output is marked in a machine-readable format so that the output can be recognized as artificially generated or manipulated.

Rules for use cases of generative AI

Those responsible for using systems that generate audio, image, video should ensure that it is clear that the content has been artificially generated or manipulated. This can be done, for example, by watermarking. For creative, satirical, fictional or analog work, this may be done in a way that does not spoil the work.

Artificially generated text is subject to a special regime. Only for text used to "inform the public on matters of public interest" must it be made clear that it was artificially generated or edited. When there is editorial control and responsibility, this need not be done.

Rules for users of emotion recognition systems or biometric categorization systems

Those in charge of using these AI systems should inform those exposed to the system about how the system works.

4.5. Other AI

AI systems that fall outside the above categories are not required to meet requirements under the AI Regulation.

But please note that if you, as a **person responsible for use**, deploy the "other category" AI system for a high-risk application as mentioned in the AI Regulation (see [1.2 High-risk AI systems on page 5](#)), then it automatically becomes a high-risk AI system and you, as the **provider** of the system, must comply with the requirements of the AI Regulation.

This brochure is a publication of: Central

Government

P.O. Box 00000 | 2500 AA The Hague

T 0800 646 39 51 (Mon to Fri 9 a.m. - 9 p.m.)

October 2024 | Publication no. 104969