



ELEVATE

EU AI Act-Ready AI Governance Policy Suite

Policies + procedures to classify AI, manage risk, and stay audit-ready
(ISO 42001-aligned).

CONTENT

Overview	3
Compliance Requirements	12
Prohibited AI Systems	14
AI Systems Classification	17
General Purpose AI (GPAI) Model	21
Risk Management System.....	28
Data and Data Governance.....	33
Technical Documentation	36
Record-Keeping.....	38
Transparency.....	40
Human Oversight.....	45
Accuracy, Robustness, and Cybersecurity	47
Provider Obligations	49
Further Processing of Personal Data in AI Regulatory Sandbox.....	51
Testing of High-Risk AI Systems in Real World Conditions (Outside of AI Regulatory Sandboxes).....	54
Testing in Real-World Conditions.....	58
Quality Management System.....	60
Codes of Conduct.....	62
Conformity Assessment.....	64
Declaration of Conformity	68
CE Marking of Conformity	70
Product Manufacturer Obligations	74
Importer Obligations	76
Distributors Obligations	79
Deployers Obligations.....	82
Fundamental Rights Impact Assessment.....	87
Other Third-Party Obligations.....	89
Post-Market Monitoring.....	92
Access to Data and Documentation	95



AI Governance Policy

Overview

Introduces the scope, purpose, and regulatory context of this policy.

AI Governance Policy

Overview

Artificial Intelligence (AI) system is a machine-based system (or software) developed to generate outputs for an explicit (or implicit) human-defined objectives. AI systems generate outputs such as predictions, content, recommendations, or decisions inferred by input received, which can influence physical (or virtual) environments. AI systems may vary in autonomy or adaptiveness after being deployed.

This AI Governance Policy is informed by: business strategy; organizational values and culture along with the amount of risk the organization is willing to pursue/retain; level of risk posed by the AI systems; legal requirements (including contracts); the organization's risk environment; and impact to relevant interested parties.

This AI Governance Policy includes principles guiding all activities of the organization related to AI and processes for handling deviations/exceptions to the policy. The organization considers topic-specific aspects where necessary to provide additional guidance or to provide cross-references to other policies dealing with these aspects. Topics considered include: AI resources/assets; AI system impact assessments; and AI system development. The organization will maintain relevant policies to guide the development, purchase, operation, and use of AI systems.

The organization places (or plans to place) on the market, puts into service (or plans to put into service), and/or places on the market general-purpose AI models in the European Union (EU) [irrespective of whether the organization is established or located within the EU or in a third country].

The organization is an AI systems' deployer having a place of establishment or located within the Union.

The organization is an AI systems' provider or deployer established or located outside of the EU, but the output produced by the AI system is used in the EU.

The organization is an importer or distributor of AI systems.

The organization is a product manufacturer placing on the market or putting into service an AI system together with their product and under their own name/trademark.

The organization is an authorized representative of providers not established in the EU.

The AI system affects persons located in the EU.

AI Literacy

The organization will take measures to ensure their staff (and other persons) dealing with the operation and use of AI systems on their behalf have a sufficient level of AI literacy considering technical knowledge, experience, education, and training as well as the context the AI systems are to be used in along with considering persons (or groups of persons) on which the AI systems are to be used.

Use of AI Systems

The organization must use AI systems responsibility and per organizational policies and procedures.

Process for Responsible Use of AI Systems

The organization will define and document the processes for the responsible use of AI systems. The organization will determine whether to use a particular AI system depending on its context and other considerations. The organization will be clear on what the considerations are and ensure policies/processes are in place to address them whether the AI system is developed by the organization or sourced from a third party. The organization will consider required approvals, cost (including ongoing monitoring and maintenance), approved sourcing requirements, and legal requirements applicable to the organization. The organization may incorporate other accepted policies for the use of other systems/assets as applicable.

Objectives for Responsible Use of AI System

The organization will identify and document objectives to guide the responsible use of AI systems. Depending on the AI system's context, the organization will identify objectives related to responsible use to include:

- fairness;
- accountability;
- transparency;
- explainability;
- reliability;
- safety;
- robustness and redundancy;
- privacy and security; and
- accessibility.

Once defined, the organization will implement mechanisms to achieve its objectives and determine if a third-party solution meets objectives or if an internally developed solution is applicable for the intended use.

The organization will determine at which stages of the AI system lifecycle meaningful human oversight objectives will be incorporated to include:

- involving human reviewers to check the outputs of the AI system (including having authority to override decisions made by the AI system);
- ensuring human oversight is included (if required) for acceptable use of the AI system according to instructions (or other documentation associated with the intended deployment of the AI system);
- monitoring the performance of the AI system (including the accuracy of the AI system outputs);
- reporting concerns related to the outputs of the AI system and their impact to relevant interested parties;
- reporting concerns with changes in the performance (or ability) of the AI system to make correct outputs on the production data; and
- considering whether automated decision-making is appropriate for a responsible approach to the use of an AI system and the intended use of the AI system

The need for human oversight will be informed by the AI system impact assessments. The personnel involved in human oversight activities will be informed, trained, and understand the instructions/documentation of the AI system. The personnel will understand the duties they carry out to satisfy human oversight objectives. Human oversight will augment automated monitoring when reporting performance issues.

Intended Use of AI System

The organization will ensure the AI system is used according to the intended uses of the AI system as well as its accompanying documentation. The AI system will be deployed according to the instructions and other documentation associated with the AI system. The deployment may require specific resources to support the deployment, (including the need to ensure that human oversight is applied as required). To ensure the AI system performance is accurate and the AI system is acceptable for use, the data the AI system uses needs to align with the documentation associated with the AI system.

AI system's operations will be monitored. If AI system deployment causes concerns regarding the impact on relevant parties (or any legal requirements), the organization will communicate these concerns to relevant personnel or to any third-party suppliers of the AI system.

The organization will keep event logs (or other documentation) related to the deployment/operation of the AI system, which can be used to demonstrate the AI system is being used as intended (or to help with communicating concerns related to the intended use of the AI system). The organization will retain event logs (and other documentation) accordingly to the AI system's intended use, the organization's data retention schedule, and applicable legal requirements for data retention.

Alignment with Other Organizational Policies

The organization will determine other policies and processes impacted by (or may apply to) the organization's objectives with respect to AI systems. Several domains may intersect with AI such as quality, security, privacy, and safety. The organization will conduct a thorough analysis to determine where current policies intersect or require updates.

Review of the AI Governance Policy

The AI Governance Policy will be reviewed at planned intervals (or at least annually) or as needed to ensure continued suitability, adequacy, and effectiveness. The Chief AI Officer is approved by management responsible for the development, review, and evaluation of the AI Governance Policy along with components. The review will include assessing opportunities for improvement over policies and approach to managing AI systems in response to changes in the environment, business circumstances, legal conditions, or technical environment. The AI Governance Policy review will take into consideration the results of management reviews.

OECD Principles

The organization incorporates the Organization for Economic Co-Operation and Development (OECD) Principles related to AI to include the following:

1. **Inclusive growth, sustainable development, and well-being:** The organization is committed to working with trustworthy AI to contribute to the growth and prosperity for all (including individuals, society, and the planet). The organization is committed to advancing global development objectives through the use of AI.
2. **Human-centered values and fairness:** The organization is committed to designing AI systems respecting the rule of law, human rights, democratic values, and diversity. The Organization is committed to include appropriate safeguards to ensure a fair and just society.
3. **Transparency and explainability:** The organization is committed to transparency and responsible disclosure involving AI systems to ensure people understand when they are engaging with them as well as the ability to challenge the AI system's outcome.
4. **Robustness, security, and safety:** The organization is committed to ensure AI systems function in a robust, secure, and safe manner throughout its lifecycle. The organization will continually assess and manage potential risks of the AI systems.
5. **Accountability:** The organization (and individuals) developing, deploying, or operating AI systems are held accountable for their proper functioning according to these OECD AI principles.

The organization may also incorporate the following principles, concepts, and values related to AI

as applicable:

- Asilimoar AI Principles
- IEEE's Ethically Aligned Design
- Montreal Declaration for a Responsible Development of Artificial Intelligence

Management Commitment to Responsible AI (RAI) Governance

The purpose of Responsible AI Governance is to promote human centric and trustworthy AI while ensuring a high level of protection over health, safety, and fundamental rights to include democracy, rule of law, and environmental protection against harmful effects of AI systems. Executive Management is committed to provide direction and support for AI systems according to business requirements and abide by:

- harmonized rules for selling, deploying, using, and development of AI systems;
- prohibitions of certain AI practices;
- specific requirements and obligations for AI systems (to include high-risk AI systems) as well as operations for these systems;
- harmonized transparency rules;
- harmonized rules for selling general-purpose AI models;
- rules on market monitoring along with surveillance, governance, and enforcement; and
- measures to support innovation focusing on SMEs.

The organization will establish accountability to uphold its responsible approach for the implementation, operation, and management of AI systems.

Roles and Responsibilities

The organization will consider related AI policies, AI objectives, and identified risks when assigning roles/responsibilities to ensure all relevant areas are covered. The organization will prioritize how the roles/responsibilities are assigned.

Institutional Review Board

If the AI system relates to people, the organization will maintain an Institutional Review Board (IRB) responsible for any ethical review applications, reviews, and/or approvals.

AI Governance Committee

The organization will establish an AI Governance Committee of senior leaders to oversee responsible AI and AI governance programs development as well as implementation of AI within

the organization. The AI Governance Committee is responsible to integrate AI governance into the existing corporate structure to prevent risks from shadow AI and ensure clear decision-making authority.

Chief AI Officer

Executive Management of the organization will appoint and designate a Chief AI Officer (CAIO) responsible for oversight of AI systems to include overall implementation, maintenance, and enforcement of related AI policies, procedures, and practices. The CAIO will report to the Chief Executive Officer (CEO) and collaborate cross-functionally with other Executive Management staff members, the AI Governance Committee, and any Institutional Review Boards. The CAIO is ultimately responsible for the decisions of the AI and is aware of the intended uses as well as limitations of any related AI analytics. The CAIO is accountable for the ethical considerations during all states of the AI lifecycle.

The CAIO will be responsible to develop principles, policies, and guardrails governing the use of AI across the organization, which is aligned to the missions and values of the organization. The CAIO is responsible to create general operating structures, objectives, and definitions (i.e., framework) such as the NIST AI Risk Management Framework and/or the OECD Framework for the Classification of AI Systems to identify high-risk AI applications needing to undertake enhanced security and privacy reviews.

The CAIO is also responsible for security, privacy, quality, development, performance, supplier relationships, and safety when it comes to AI as well as the following:

- Ensuring adequate performance of the organization's AI portfolio (such predictive modeling or machine learning);
- Collaborate with senior leaders to determine risk tolerances for developing/using AI systems;
- Support AI risk management efforts and play an active role in such efforts;
- Integrate a risk and harm prevention mindset throughout the AI lifecycle as part of the organization's culture;
- Support competent risk management executives along with developing assessment standards covering the necessary skills, training, resources, and domain knowledge for AI personnel to fulfill their assigned responsibilities;
- Define/differentiate various human roles/responsibilities when using, interacting with, or monitoring AI systems;
- Capture/track risk information related to human-AI configurations and associated outcomes;
- Develop proficiency standards for AI actors carrying out system operation tasks and system oversight tasks;
- Specify risk management training protocols for AI actors carrying out system operation tasks and system oversight tasks;

- Establish procedures to address AI actor roles/responsibilities for human oversight of deployed systems;
- Define human-AI configurations (configurations where AI systems are explicitly designated and treated as team members in primarily human teams) in relation to organizational risk tolerance along with associated documentation;
- Enhance the explanation, interpretation, and overall transparency of AI systems;
- Establish impact assessment processes for AI systems;
- Manage risks regarding known difficulties in human-AI configurations, human-AI teaming, and AI system user experience and user interactions (UI/UX);
- Establish processes regarding public disclosure of incidents and information sharing; and
- Delegate the power, resources, and authorization to perform risk management to each appropriate level throughout the management chain.

The CAIO will be responsible for receiving concerns regarding AI no longer meeting an ethical framework. The CAIO will appropriately investigate and remediate the issue, if possible. The CAIO maintains the authority to modify limit, or stop the use of the AI.

The CAIO will adhere to and consult AI regulations, requirements, contractual obligations, standards, guidelines, and best industry practices involving AI governance. The CAIO will stay up-to-date on high-profile litigation cases involving AI to prepare for evolving legal issues such as issues around generative AI and intellectual property rights.

Other roles, responsibilities, and delegation of authority:

- AI Design
- AI Development
- AI Acquisition/Procurement
- Human-AI Interaction - need to define and differentiate the various human roles/responsibilities when using, interacting with, or monitoring AI systems; capture and track risk information related to human-AI configurations and associated outcomes;
- Product/Project Management
- Impact Assessment
- AI Testing and Evaluation
- AI Monitoring - who is responsible for maintaining, re-verifying, monitoring, and updating AI once deployed?
- AI Audit/Assessment

Diversity Requirements

The organization will utilize a diverse team throughout the AI system's lifecycle to inform on decisions made related to mapping, measuring, and managing AI risks. The team should be made up of a diverse group of individuals reflecting a diverse perspective of demographics, disciplines, experience, expertise, opinions, and backgrounds to enhance the organization's capacity and

capabilities in anticipating risks. The diverse perspectives may come from technical and non-technical communities through the AI systems' life cycle in order to anticipate/mitigate unintended consequences to include potential bias/discrimination. The organization may need to consult with external parties when internal teams lack the necessary diversity. External expertise should supplement organizational diversity, equity, inclusion, and accessibility where internal expertise is lacking.



AI Governance Policy

Compliance Requirements

Outlines mandatory legal and regulatory obligations for AI systems.

Compliance Requirements

The organization, maintaining high-risk AI systems, will comply with the requirements of the EU AI Act. The organization will take into account the intended purpose of the high-risk AI system, the generally acknowledged state of the art on AI (and AI related technologies), and the organization's risk management system.

The organization may also need to comply with the following regulations, as applicable, such as:

- The Artificial Intelligence and Data Act (Canada);
- NYC Local Law 144 of 2021 (regulating automated employment decision tools);
- American Data Privacy and protection Act, Section 207 (US); and/or
- Internet Information Service Algorithm Recommendation Management Regulations (China).
-

The organization may voluntarily abide by the following practices, structures, guidelines, and actions involving AI:

- White House's voluntary commitments from leading AI companies; and/or
- Canada's generative AI code of conduct.

The organization may also demonstrate or provide assurance with compliance to regulations by abiding by the following standards/certifications:

- ISO/IEC JTC 1/SC 42;
- IEEE P7000 series of standards projects;
- CEN/CENELECT standards development; and/or
- RAI Institute's Certification Program for AI Systems.



AI Governance Policy

Prohibited AI Systems

Defines AI practices that are banned due to unacceptable risk.

Prohibited AI Systems

The organization is prohibited from placing on the market, putting into service, or using the following AI systems:

- **Subliminal manipulation resulting in physical/psychological harm.** Subliminal techniques include techniques beyond a person's consciousness (or purposefully manipulative/deceptive techniques) with the objective to (or the effect of) materially distorting a person's (or a group of persons') behavior by appreciably impairing the person's ability to make an informed decision, thereby causing the person to take a decision the person would not have otherwise taken in a manner causing (or is likely to cause) the person, another person, or group of persons significant harm. For example, an inaudible sound played in truck drivers' cabins to push them to drive longer than healthy and safe. AI used to find the frequency maximizing this effect on drivers.
- **Exploitation of children/mentally disabled persons resulting in physical/psychological harm.** An AI system exploiting any of the vulnerabilities of a person (or a specific group of persons) due to their age, disability or a specific social/economic situation, with the objective to (or the effect of) materially distorting the behavior of the person (or a person pertaining to the group) in a manner causing (or is reasonably likely to cause) the person (or another person) significant harm. For example, a doll with an integrated voice assistant encouraging a minor to engage in progressively dangerous behavior or challenges in the guise of a fun/cool game.
- **General purpose social scoring.** AI systems for the evaluation (or classification) of natural persons (or groups thereof) over a certain period of time based on their social behavior (or known, inferred or predicted personal or personality characteristics), with the social score leading to either (or both) of the following:
 - o detrimental/unfavorable treatment of certain natural persons (or whole groups thereof) in social contexts unrelated to the contexts in which the data was originally generated/collected; and/or
 - o detrimental/unfavorable treatment of certain natural persons)or groups thereof) unjustified/disproportionate to their social behavior (or its gravity).
 - o For example, an AI system identifies at-risk children in need of social care based on insignificant or irrelevant social 'misbehavior' of parents (e.g., missing a doctor's appointment or divorce).
- **Criminal prediction.** AI system for making risk assessments of natural persons in order to assess (or predict) the risk of a natural person to commit a criminal offense, based solely on the profiling of a natural person (or on assessing their personality traits and characteristics). Note: This prohibition shall not apply to AI systems used to support the human assessment of the involvement of a person in a criminal activity, which is already based on objective and verifiable facts directly linked to a criminal activity.
- **Facial recognition through untargeted scraping.** AI systems creating (or expanding)

facial recognition databases through the untargeted scraping of facial images from the internet (or CCTV footage).

- **Determining emotions in workplace/education.** AI systems to infer emotions of a natural person in the areas of workplace and education institutions except in cases where the use of the AI system is intended to be put in place or into the market for medical or safety reasons.
- **Biometric Categorization.** AI systems categorizing individual natural persons based on their biometric data to deduce or infer their race, political opinions, trade union membership, religious or philosophical beliefs, sex life or sexual orientation. Note: this prohibition does not cover any labelling or filtering of lawfully acquired biometric datasets, such as images, based on biometric data or categorizing of biometric data in the area of law enforcement;
- **Remote biometric identification (RBI) for law enforcement purpose in publicly accessible spaces (with some exceptions).** ‘Real-time’ remote biometric identification systems in publicly accessible spaces for the purpose of law enforcement unless (and in as far as) such use is strictly necessary for one of the following objectives:
 - o the targeted search for specific victims of abduction, trafficking in human beings, and sexual exploitation of human beings as well as search for missing persons;
 - o the prevention of a specific, substantial, and imminent threat to the life (or physical safety) of natural persons or a genuine and present (or genuine and foreseeable) threat of a terrorist attack;
 - o the localization (or identification) of a person suspected of having committed a criminal offence, for the purposes of conducting a criminal investigation, prosecution, or executing a criminal penalty for offences [referred to in Annex IIa and punishable in the Member State concerned by a custodial sentence or a detention order for a maximum period of at least four years]. Note: The prohibition is without prejudice to the provisions in Article 9 of the GDPR for the processing of biometric data for purposes other than law enforcement.
 - o For example, all faces captured live by video cameras checked, in real time, against a database to identify a terrorist. Some exceptions include: Searching for victims of crime; Threat to life or physical integrity or of terrorism; Serious crime (EU Arrest Warrant); Ex-ante authorization by judicial authority (or independent administrative body); or putting on market of RBI systems (real-time and ex-post) with ex-ante third party conformity assessment, enhanced logging requirements, and abiding by the “Four eyes” principle (i.e. a certain activity such as a decision or transaction is approved by at least two (2) people).

AI Governance Policy

AI Systems Classification

Categorizes AI systems based on risk and regulatory requirements.

AI Systems Classification

The organization will perform a review of all AI systems and classify the AI systems based on the following classification:

1. **Prohibited AI systems** (see above)
2. **High-Risk AI systems:** an AI system is intended to be used as a safety component of a product (or the AI system is itself a product) covered under the EU harmonized legislation summarized as follows AND the product whose safety component is the AI system (or the AI system itself as a product) is required to undergo a third-party conformity assessment (with the intent to place on the market or put into service the product pursuant to the EU harmonized legislation summarized as follows):
 - a. As listed in Annex II Section A:
 - i. Machinery
 - ii. Toy Safety
 - iii. Recreational Craft and Personal Watercraft
 - iv. Lifts and Safety Components for Lifts
 - v. Equipment and Protective Systems Intended for Use in Potentially Explosive Atmospheres
 - vi. Radio Equipment
 - vii. Pressure Equipment
 - viii. Cableway Installations
 - ix. Personal Protective Equipment
 - x. Appliances Burning Gaseous Fuels
 - xi. Medical Devices
 - xii. In Vitro Diagnostic Medical Devices
 - b. As listed in Annex II Section B:
 - i. Civil Aviation Security
 - ii. Two- or Three-Wheel Vehicles and Quadricycles
 - iii. Agricultural and Forestry Vehicles
 - iv. Marine Equipment
 - v. Rail System
 - vi. Motor Vehicles and Trailers
 - vii. Components and Separate Technical Units of Motor Vehicles and Trailers
 - viii. Unmanned Aircrafts and Engines, Propellers, Parts, and Equipment to Control Them Remotely
 - c. Annex III:

- i. Biometric identification and categorization of natural persons
 - ii. Management and operation of critical infrastructure
 - iii. Education and vocational training
 - iv. Recruitment: Employment and workers management, access to self-employment
 - v. Access to and enjoyment of essential private services and public services/benefits
 - vi. Law enforcement
 - vii. Migration, asylum, and border control management
 - viii. Administration of justice and democratic processes
3. **Non-High-Risk (Limited Risk) AI systems:** an AI system not considered high-risk AI systems for example, impersonation (bots), biometric categorization systems, and emotional (text audio or visual) (i.e., 'deep fake').
 4. **Minimal (or No Risk) AI system:** any other AI system not covered by other classifications for example, AI-enabled videogames, spam filters, and predictive resource optimization.

The organization may consider an AI system following under Annex III as NOT high-risk if the AI system does not pose a significant risk of harm to the health, safety, or fundamental rights of natural persons (including by not materially influencing the outcome of decision making). One or more of the following criteria must be fulfilled in order to classify an AI system following under Annex III as NOT being high risk:

- the AI system is intended to perform a narrow procedural task;
- the AI system is intended to improve the result of a previously completed human activity;
- the AI system is intended to detect decision-making patterns (or deviations) from prior decision-making patterns and is not meant to replace/influence the previously completed human assessment (without proper human review); or
- the AI system is intended to perform a preparatory task to an assessment relevant for the purpose of the use cases listed in Annex III. Note: AI systems performing profiling of natural persons are always considered high-risk.

The organization (as a provider of an AI system) must document the assessment performed to determine an AI system falling under Annex III is NOT high-risk before the AI system is placed on the market or put into service. The organization is still subject to registration obligations and must provide documentation of the assessment upon request by a national competent authority. High-risk AI systems are subject to conformity assessments either internal or required to be performed by a third-party. Note: These policies are focused on high-risk AI systems. The intended purpose and the risk management system of high-risk AI system will be taken into account when ensuring compliance with high-risk AI system requirements.

Non-high risk (limited risk) AI systems are subject to transparency requirements such as notifying humans of their interaction with the AI system, notifying humans of the application of emotional recognition or biometric categorization systems, and applying labels to 'deep fakes' (unless

necessary for the exercise of fundamental rights, freedoms, or public interest).

Minimal (or no risk) AI systems have no mandatory obligations; however, it strongly suggested a code of conduct is followed with transparency requirements being met.



AI Governance Policy

General Purpose AI (GPAI) Model

Defines obligations and controls for general-purpose AI models.

General Purpose AI (GPAI) Model

A general purpose AI model is classified as general-purpose AI model with systemic risk if it meets any of the following criteria:

- it has high impact capabilities evaluated on the basis of appropriate technical tools and methodologies (including indicators and benchmarks); or
- based on a decision of the Commission, ex officio, or following a qualified alert by the scientific panel a general purpose AI model has capabilities (or impact equivalent) to those above.

A general purpose AI model is presumed to have high impact capabilities when the cumulative amount of compute used for its training is greater than 10^{25} floating point operations (FLOPs).

GPAI Procedures

The organization, as the relevant provider, will notify the Commission without delay and within two (2) weeks of meeting requirements (or becoming known to meet the requirements) for a GPAI model. The notification will include the information necessary to demonstrate the relevant requirements have been met. If the Commission becomes aware of a GPAI model presenting systemic risks of which it has not been notified, it may decide to designate it as a model with systemic risk.

The organization may attempt to have the AI model exempted from being designated a GPAI by presenting (with the notification) sufficient substantiated arguments to exceptionally demonstrate that although it meets the requirements, the GPAI does not present systemic risks (due to its specific characteristics) and therefore, should be classified as a GPAI model with systemic risk.

The Commission will decide whether the arguments are sufficiently substantiated (or not) and reject the arguments classifying the GPAI as a GPAI with systemic risk. The Commission may designate a GPAI as presenting systemic risks, ex officio, or following a qualified alert of the scientific panel on the basis of criteria as follows:

- number of parameters of the model;
- quality or size of the data set, for example measured through tokens;
- the amount of compute used for training the model, measured in FLOPs or indicated by a combination of other variables such as estimated cost of training, estimated time required for the training, or estimated energy consumption for the training;
- input and output modalities of the model, such as text to text (large language models),

text to image, multi-modality, and the state-of-the-art thresholds for determining high-impact capabilities for each modality, and the specific type of inputs and outputs (e.g. biological sequences);

- benchmarks and evaluations of capabilities of the model, including considering the number of tasks without additional training, adaptability to learn new, distinct tasks, its degree of autonomy and scalability, the tools it has access to;
- it has a high impact on the internal market due to its reach, which shall be presumed when it has been made available to at least 10,000 registered business users established in the Union; and number of registered end-users

If the organization's AI model has been designated as a GPAI with systemic risks, the organization may request the model to be reassessed by the Commission. The request should contain objective, concrete, and new reasons arising since the designation decision. The organization may request reassessment as early as six (6) months after the designation decision.

Obligations for Providers of GPAI Models

The organization, as a provider of GPAI models, will develop and maintain up-to-date technical documentation of the model including its training/testing process and the results of evaluation. The technical documentation will be provided, upon request, to the AI Office and the national competent authorities. The technical documentation will contain, at a minimum, the following:

- Section 1: Information to be provided by all providers of general-purpose AI models
 - o A general description of the general purpose AI model including:
 - o the tasks that the model is intended to perform and the type and nature of AI systems in which it can be integrated;
 - o acceptable use policies applicable;
 - o the date of release and methods of distribution;
 - o the architecture and number of parameters;
 - o modality (e.g. text, image) and format of inputs and outputs;
 - o the license.
 - o A detailed description of the elements of the model and relevant information of the process for the development, including the following elements:
 - o the technical means (e.g. instructions of use, infrastructure, tools) required for the general-purpose AI model to be integrated in AI systems;
 - o the design specifications of the model and training process, including training methodologies and techniques, the key design choices including the rationale and assumptions made; what the model is designed to optimize for and the relevance of the different parameters, as applicable;

- o information on the data used for training, testing and validation, where applicable, including type and provenance of data and curation methodologies (e.g. cleaning, filtering etc), the number of data points, their scope and main characteristics; how the data was obtained and selected as well as all other measures to detect the unsuitability of data sources and methods to detect identifiable biases, where applicable;
 - o the computational resources used to train the model (e.g. number of floating point operations – FLOPs), training time, and other relevant details related to the training;
 - o known or estimated energy consumption of the model; in case not known, this could be based on information about computational resources used
- Section 2: Additional information to be provided by providers of general purpose AI model with systemic risk
 - o Detailed description of the evaluation strategies, including evaluation results, on the basis of available public evaluation protocols and tools or otherwise of other evaluation methodologies. Evaluation strategies shall include evaluation criteria, metrics and the methodology on the identification of limitations.
 - o Where applicable, detailed description of the measures put in place for the purpose of conducting internal and/or external adversarial testing (e.g. red teaming), model adaptations, including alignment and fine-tuning.
 - o Where applicable, detailed description of the system architecture explaining how software components build or feed into each other and integrate into the overall processing.

The organization, as a provider of GPAI models, will develop and maintain up-to-date as well as make available information/documentation to providers of AI systems who intend to integrate the GPAI model into their AI system. Without prejudice to the need to respect/protect intellectual property rights and confidential business information (or trade secrets) according to laws, the information/documentation will enable providers of AI systems to have a good understanding of the capabilities and limitations of the GPAI model and to comply with their obligations pursuant to the EU AI Act and contain, at a minimum, the following:

- A general description of the general purpose AI model including:
 - o the tasks that the model is intended to perform and the type and nature of AI systems in which it can be integrated;
 - o acceptable use policies applicable;
 - o the date of release and methods of distribution;
 - o how the model interacts or can be used to interact with hardware or software that is not part of the model itself, where applicable;
 - o the versions of relevant software related to the use of the general purpose AI model, where applicable;
 - o architecture and number of parameters,

- o modality (e.g., text, image) and format of inputs and outputs; (h) the license for the model.
- o A description of the elements of the model and of the process for its development, including:
 - o
 - o the technical means (e.g. instructions of use, infrastructure, tools) required for the general-purpose AI model to be integrated in AI systems;
 - o modality (e.g., text, image, etc.) and format of the inputs and outputs and their maximum size (e.g., context window length, etc.)
 - o information on the data used for training, testing and validation, where applicable, including, type and provenance of data and curation methodologies.

The organization, as a provider of GPAI models, will implement a policy to respect Union copyright law including through state of the art technologies and reservations of rights. The organization will develop and make publicly available a sufficiently detailed summary about the content used for training of the GPAI model according to the AI Office's template.

The organization, as a provider of GPAI models, will cooperate as necessary with the Commission and the national competent authorities in the exercise of their competences and powers pursuant to this Regulation.

The organization, as a provider of GPAI models, may rely on codes of practice to demonstrate compliance with obligations, until a harmonized standard is published. Compliance with a European harmonized standard grants providers the presumption of conformity. The organization, as a provider of GPAI models with systemic risks, who does not adhere to an approved code of practice will demonstrate alternative adequate means of compliance for approval of the Commission.

Note: The Commission is empowered to adopt delegated acts to detail measurement and calculation methodologies with a view to allow comparable/verifiable documentation related to the computational resources used to train the model, training time, and other relevant details related to training as well as know (or estimated) energy consumption of the model. The Commission is also empowered to adopt delegated acts to amend the elements of the technical documentations in light of evolving technological developments

Any information and documentation, including trade secrets, will be treated in compliance with the confidentiality obligations.

Authorized Representative for GPAI Model

The organization, as a provider of GPAI model outside of the union, will appoint an authorized

representative, which is established in the union, by written mandate and will enable the authorized representative to perform tasks under the EU AI Act prior to placing a GPAI model on the Union market.

The authorized representative will perform the tasks specified in the mandate received from the organization, as a provider of GPAI model. The authorized representative will provide a copy of the mandate to the AI Office upon request, in one of the official languages of the institutions of the Union. The mandate will empower the authorized representative to carry out the following tasks:

- verify the technical documentation has been developed and all obligations, where applicable, have been fulfilled by the organization;
- keep a copy of the technical documentation at the disposal of the AI Office and national competent authorities, for a period ending ten (10) years after the model has been placed on the market and the contact details of the provider by which the authorized representative has been appointed;
- provide the AI Office, upon a reasoned request, with all the information and documentation necessary to demonstrate the compliance with obligations; and
- cooperate with the AI Office and national competent authorities, upon a reasoned request, on any action the latter takes in relation to the general purpose AI model with systemic risks, including when the model is integrated into AI systems placed on the market (or put into service) in the Union.

The mandate will empower the authorized representative to be addressed, in addition to (or instead of the organization), by the AI Office (or the national competent authorities), on all issues related to ensuring compliance with regulations.

The authorized representative will terminate the mandate if it considers (or has reason to consider) the organization acts contrary to its obligations under regulations. In such a case, the authorized representative will also immediately inform the AI Office about the termination of the mandate and the reasons thereof.

Note: *This obligation does not apply to providers of GPAI models made accessible to the public under a free and open source license allowing for the access, usage, modification, and distribution of the model, and whose parameters (including the weights, the information on the model architecture, and the information on model usage) are made publicly available, unless the GPAI models present systemic risks.*

Obligations for Providers of GPAI Models with Systemic Risk

The organization, as a provider of GPAI models with systemic risks, will:

- perform model evaluation in accordance with standardized protocols and tools reflecting

the state of the art, including conducting and documenting adversarial testing of the model with a view to identify and mitigate systemic risk;

- assess and mitigate possible systemic risks at Union level, including their sources, that may stem from the development, placing on the market, or use of general purpose AI models with systemic risk;
- 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; and
- ensure an adequate level of cybersecurity protection for the general purpose AI model with systemic risk and the physical infrastructure of the model.

The organization, as a provider of GPAI models with systemic risks, may rely on codes of practice to demonstrate compliance with obligations, until a harmonized standard is published. Compliance with a European harmonized standard grants providers the presumption of conformity. The organization, as a provider of GPAI models with systemic risks, who does not adhere to an approved code of practice will demonstrate alternative adequate means of compliance for approval of the Commission.

Any information and documentation, including trade secrets, will be treated in compliance with the confidentiality obligations.



AI Governance Policy

Risk Management System

Establishes processes to identify, assess, and mitigate AI risks.

Risk Management System

The organization will establish, implement, document, and maintain a risk management system related to AI systems (including high-risk AI systems). The risk management system will consist of a continuous iterative process running throughout the entire AI system lifecycle and covering more specifically high-risk AI system lifecycles. The risk management system requires regular systematic updates and will comprise of the following:

1. identification and analysis of the known and the reasonably foreseeable risks the AI system (such as high-risk AI system) can pose to the health, safety, or fundamental rights when the AI system is used according to its intended purpose;
2. estimation and evaluation of the risks emerging when the AI system is used according to its intended purpose and under conditions of reasonably foreseeable misuse;
3. evaluation of other possible risks based on the analysis of data gathered from post-market monitoring; and
4. adoption of suitable risk management measures.

The organization will be concerned with those risks which may be reasonably mitigated (or eliminated) through the development (or design) of the AI system (such as high-risk AI system) as well as the provision of adequate technical information.

Effects and possible interactions resulting from the combined application of requirements will be considered in risk management measures with a view to minimize risks more effectively while achieving an appropriate balance in implementing the measure to fulfill those requirements taking into account generally acknowledged state-of-the-art as well as relevant standards or specifications of the AI system. Specific considerations will be made within the risk management system for AI systems being accessed or having an impact on children (such as persons under the age of 18) and other vulnerable groups of people as applicable.

Only acceptable residual risks associated with each hazard or the overall residual risk of an AI system will be approved based on the intended purpose (or under reasonably foreseeable misuse conditions). Acceptable residual risks will be communicated to the user.

The following risk management measures will be ensured:

- elimination or reduction of risks as far as possible through adequate design and development;
- where appropriate, implementation of adequate mitigation and control measures in relation to risks, which cannot be eliminated;
- provision of adequate information to users regarding risks and training users, where appropriate.

To reduce or eliminate risks, the organization will consider the technical knowledge, experience, education, or training expected by the user as well as the environment in which the high-risk AI system is intended to be used.

To identify the most appropriate risk management measures, high-risk AI systems will be tested. Testing will ensure high-risk AI systems perform consistently for its intended purpose and it is compliant with this policy along with other regulations, obligations, and requirements. Testing procedures must be suitable to achieve the intended purpose and do not need to go beyond what is necessary to achieve this purpose.

Testing will be performed on high-risk AI systems, as appropriate, at any point in the development process and prior to placing the high-risk AI system on the market (or put into service). Testing will be made against appropriate preliminarily defined metrics and probabilistic thresholds.

AI System Resources and Documentation

The organization must account for resources (including AI system components and assets) of the AI system in order to fully understand, address risks, and address impacts. The organization will identify and document relevant resources required for the activities at given AI system lifecycle stages as well as other relevant related activities.

The organization may utilize data flow diagrams and/or system architecture diagrams to inform the AI system impact assessments. Additional resources may include the following: AI system components; data resources (such as data used at any stage in the AI system lifecycle); tooling resources (such as AI algorithms, models, or tools); system and computing resources (such as hardware to develop/run AI models, storage for data and tooling resources); human resources (such as experts for development, sales, training, operation, and maintenance of AI systems) in relation to roles throughout the AI system lifecycle. The organization may obtain resources internally, by customers, or by third parties.

Data Resources

Data resources utilized for the AI system will be documented. Documentation on data will include, but not limited to:

- data provenance;
- last updated/modified data data (such as data tag in metadata);
- for machine learning, data categories (such as training, validation, test, and production data);
- data categories (as may be defined in ISO/IEC 19944-1);
- labelling data processes;

- intended data use;
- data quality (as described in the ISO/IEC 5259 series2);
- applicable data retention and disposal policies;
- known/potential bias issues in the data; and
- data preparation.

Tooling Resources

Tooling resources utilized for the AI system will be documented. Documentation on tooling resources for AI systems (such as for machine learning) will include, but not limited to:

- algorithm types and machine learning models;
- data conditioning tools or processes;
- optimization methods;
- evaluation methods;
- provisioning tools for resources;
- tools to aid model development; and
- software and hardware for AI system design, development and deployment.

System and Computing Resources

System and computing resources utilized for the AI system will be documented. Information about system and computing resources for an AI system may include, but not limited to:

- AI system resource requirements (such as ensuring the system can run on constrained resource devices);
- system and computing resources location (such as on-premises, cloud computing, or edge computing);
- processing resources (including network and storage); and
- the impact of the hardware used to run the AI system workloads (such as the impact to the environment - either through use or the manufacturing of the hardware - or cost of using the hardware).

The organization will consider different resources required to permit continual improvement of AI systems. Development, deployment, and operation of the system can have different system needs and requirements.

Human Resources

Human resources (and their competences) utilized for the development, deployment, operation, change management, maintenance, transfer, and decommissioning (as well as verification and integration) of the AI system will be documented. The organization will consider the need for diverse expertise and include types of roles necessary for the system. If inclusion is a necessary

component of system design, the organization will include specific demographic groups related to data sets used to train machine learning models. Necessary human resources may include, but not limited to:

- data scientists;
- roles related to human oversight of AI systems;
- experts on trustworthiness topics such as safety, security and privacy; and
- AI researchers and specialists, and domain experts relevant to the AI systems.

Different resources may be necessary at different stages of the AI system lifecycle.

Corrective Actions and Notification

The organization will immediately take necessary corrective actions to mitigate risks associated with non-conformities. If the organization considers (or have reason to consider) a high-risk AI system placed on the market (or put into service) doesn't conform with these policies or relevant regulations, the organization will correct the non-conformity, withdraw the system, disable the system, or recall the system, as appropriate.

Where the high-risk AI system presents a risk and the organization becomes aware of the risk, the organization will immediately investigate the causes, in collaboration with the reporting deployer, where applicable.

The organization will inform the distributors of the high-risk AI system in question and, where applicable, the deployers, the authorized representative and importers accordingly. The organization will also inform the market surveillance authorities in which the high-risk AI system is available and the notified body issuing a certification for the high-risk AI system of the nature of the non-compliance along with any relevant corrective actions taken, where applicable.



AI Governance Policy

Data and Data Governance

Defines requirements for data quality, governance, and controls.

Data and Data Governance

Appropriate data governance and management practices will be applied to the development of AI systems to ensure compliance with these policies and other related regulations.

If the high-risk AI system involves training of models with data, the organization will ensure the data is developed on the basis of training, validation, and testing data sets meeting quality criteria. Training, validation, and testing data sets will be subject to appropriate data governance and management practices to include:

- the relevant design choices;
- data collection processes and origin of data (and in the case of personal data, the original purpose of data collection);
- relevant data preparation processing operations, such as annotation, labelling, cleaning, updating, enrichment, and aggregation;
- the formulation of relevant assumptions, notably with respect to the information the data are supposed to measure and represent;
- a prior assessment of the availability, quantity and suitability of the data sets are needed;
- examination in view of possible biases likely to impact the health and safety of persons, negatively impact fundamental rights, or lead to discrimination prohibited under regulations (especially where data outputs influence inputs for future operations);
- appropriate measures to detect, prevent, and mitigate possible biases identified;
- the identification of any possible data gaps or shortcomings, and how those gaps and shortcomings can be addressed.

Training, validation, and testing data sets will be relevant, representative, free of errors and complete. Data sets will have the appropriate statistical properties, including, where applicable, as regards the persons or groups of persons on which the high-risk AI system is intended to be used. These characteristics of the data sets may be met at the level of individual data sets or a combination thereof.

Training, validation, and testing data sets will take into account, to the extent required by the intended purpose, the characteristics or elements particular to the specific geographical, contextual, behavioral or functional setting within which the high-risk AI system is intended to be used.

To the extent it is strictly necessary for the purposes of ensuring bias monitoring, detection and correction in relation to the high-risk AI systems, the organization may process special categories of personal data referred to in the General Data Protection Regulation (GDPR), subject to appropriate safeguards for the fundamental rights and freedoms of natural persons. The following conditions must apply in order to process special categories of personal data:

- the bias detection and correction cannot be effectively fulfilled by processing other data (including the use of synthetic or anonymized data);
- the special categories of personal data processed are subject to technical limitations on the re-use of the personal data and state of the art security and privacy-preserving measures (including pseudonymization);
- the special categories of personal data processed are subject to measures to ensure the personal data processed are secured, protected, subject to suitable safeguards (including strict controls and documentation of the access to avoid misuse), and ensure only authorized persons have access to those personal data with appropriate confidentiality obligations;
- the special categories of personal data processed are not to be transmitted, transferred, or otherwise accessed by other parties;
- the special categories of personal data processed are deleted once the bias has been corrected or the personal data has reached the end of its retention period, whatever comes first; and
- the records of processing activities pursuant to the GDPR includes justification why the processing of special categories of personal data was strictly necessary to detect and correct biases and this objective could not be achieved by processing other data.

For the development of high-risk AI systems NOT using techniques involving the training of models, the above criteria shall apply only to the testing data sets (not to the training or validation data sets).



AI Governance Policy

Technical Documentation

Specifies documentation needed to demonstrate AI compliance.

Technical Documentation

The organization will draft technical documentation prior to the system being placed on the market or put into service. This technical documentation must be kept up to date. The technical documentation will be drafted to demonstrate the high-risk AI system complies with the requirements of this policy and relevant regulations as well as provide national competent authorities (and notified bodies) with all the necessary information to assess compliance of the AI system with these requirements. The technical documentation will contain at least the following elements [see Technical Documentation Policy/Procedures for additional information]:

- A general description of the AI system;
- A detailed description of the elements of the AI system and the process for its development;
- Detailed information about monitoring, functioning, and control of the system;
- A description of the appropriateness of the performance metrics for the specific AI system;
- A detailed description of the risk management system;
- A description of any change made by the provider to the system through its lifecycle;
- A list of the harmonized standards applied in full (or in part) (or where no harmonized standards have been applied, a detailed description of the solutions adopted to meet requirements including a list of other relevant standards and technical specifications applied);
- A copy of the EU declaration of conformity; and
- A detailed description of the system in place to evaluate the AI system performance in the post market phase.

In general, one single technical documentation will be drafted containing all of this information along with any information required under specific legal acts.

Documentation Keeping

The organization, as a provider, must retain for a period of ten (10) years after the AI system has been placed on the market (or put into service) the following:

- the technical documentation;
- the documentation concerning the quality management system;
- the documentation concerning the changes approved by notified bodies where applicable;
- the decisions and other documents issued by the notified bodies, where applicable; and
- the EU declaration of conformity.



AI Governance Policy

Record-Keeping

Defines requirements for logging and retention of AI records.

Record-Keeping

The organization will design and develop high-risk AI systems with the capability to automatically record events (logs) over the duration of the lifetime of the system. In order to ensure a level of traceability of the AI system's functioning appropriate to the intended purpose of the system, logging capabilities will enable the recording of events relevant for:

- identification of situations resulting in the AI system presenting a risk or in a substantial modification;
- facilitation of the post-market monitoring; and
- monitoring of the operation of high-risk AI systems.

Logging capabilities will provide at least the following:

-
- recording of the period of each use of the system (start date and time and end date and time of each use);
- the reference database against which input data has been checked by the system;
- the input data for which the search has led to a match; and
- the identification of the natural persons involved in the verification of the results.

Automatically Generated Logs

The organization, as a provider of high-risk AI systems, will keep the logs, automatically generated by the high-risk AI systems, to the extent such logs are under their control. The logs will be kept for a period appropriate to the intended purpose of the high-risk AI system of at least six (6) months unless provided otherwise by regulations.



AI Governance Policy

Transparency

Establishes transparency obligations toward users and authorities.

Transparency

The organization will design and develop high-risk AI systems to ensure its operation is sufficiently transparent to enable users to interpret the system's output and use it appropriately. The organization will ensure an appropriate type and degree of transparency with a view to achieving compliance with relevant obligations of the user and the organization.

High-risk AI systems will be accompanied by instructions for use in an appropriate digital format (or otherwise) to ensure the instructions are concise, complete, correct, and provide clear information. The information will be relevant, accessible, and comprehensible to users. Instructions will specify:

- the identity and the contact details of the organization and, where applicable, of the organization's authorized representative;
- the characteristics, capabilities and limitations of performance of the high-risk AI system, including:
 - o its intended purpose;
 - o the level of accuracy, robustness, and cybersecurity against which the high-risk AI system has been tested/validated along with expected (or any known/foreseeable circumstances) impacting the expected level of accuracy, robustness, and cybersecurity;
 - o any known (or foreseeable) circumstance related to the use of the high-risk AI system according to its intended purpose or under conditions of reasonably foreseeable misuse, which may lead to risks to the health/safety or fundamental rights;
 - o where applicable, the technical capabilities and characteristics of the AI system to provide information relevant to explain its output;
 - o its performance regarding the persons (or groups of persons) on which the system is intended to be used;
 - o when appropriate, specifications for the input data (or any other relevant information) in terms of the training, validation, and testing data sets used taking into account the intended purpose of the AI system;
 - o where applicable, information to enable deployers to interpret the system's output and use it appropriately;
 - o the changes to the high-risk AI system and its performance, which have been pre-determined by the organization at the moment of the initial conformity assessment, if any;
 - o the human oversight measures, including the technical measures put in place to facilitate the interpretation of the outputs of AI systems by the deployers;
 - o the computational and hardware resources needed, the expected lifetime of the high-risk AI system, and any necessary maintenance and care measures (including their frequency) to ensure the proper functioning of that AI system (including as

- regards software updates); and
- o where relevant, a description of the mechanisms included within the AI system allowing users to properly collect, store, and interpret the logs.

The organization, as a provider, will ensure natural persons are informed they are interacting with a non-high-risk AI system when the AI system is intended to interact with natural persons unless this is obvious from the circumstances and the context of use. Note: This obligation does not apply to AI systems authorized by law to detect, prevent, investigate, and prosecute criminal offenses, subject to appropriate safeguards for the rights and freedoms of third parties (unless those systems are available for the public to report a criminal offense).

The organization, as a provider of AI systems (including GPAI systems) generating synthetic audio, image, video or text content, will ensure the outputs of the AI system are marked in a machine-readable format (and detectable as artificially generated or manipulated).

Considering specificities/limitations of different types of content, costs of implementation, and the generally acknowledged state-of-the-art (as may be reflected in relevant technical standards), the organization (as a provider) will ensure their technical solutions are effective, interoperable, robust, and reliable as far as this is technically feasible. Note: This obligation does not apply to the extent the AI systems perform an assistive function for standard editing, do not substantially alter the input data provided by the deployer (or the semantics thereof), or where authorized by law to detect, prevent, investigate and prosecute criminal offences.

The organization, as a deployer, of an emotion recognition system (or a biometric categorization system) will inform of the operation of the system the natural persons exposed thereto and process the personal data according to regulations, as applicable. Note: This obligation does not apply to AI systems used for biometric categorization and emotion recognition, which are permitted by law to detect, prevent and investigate criminal offences, subject to appropriate safeguards for the rights and freedoms of third parties, and in compliance with Union law. The organization, as a deployer, of a system generating (or manipulating image, audio or video content constituting a deep fake, will disclose the content has been artificially generated (or manipulated). Note: This obligation does not apply where the use is authorized by law to detect, prevent, investigate and prosecute criminal offence.

Where the content forms part of an evidently artistic, creative, satirical, fictional analogous work or program, the transparency obligations are limited to disclosure of the existence of such generated (or manipulated) content in an appropriate manner not hampering the display (or enjoyment) of the work.

The organization, as a deployer, of an AI system generating (or manipulating) text, which is published with the purpose of informing the public on matters of public interest, will disclose the text has been artificially generated (or manipulated). Note: This obligation does not apply where the use is authorized by law to detect, prevent, investigate and prosecute criminal offences or

where the AI-generated content has undergone a process of human review (or editorial control) and where a natural (or legal person) holds editorial responsibility for the publication of the content.

The information will be provided to the concerned natural persons in a clear and distinguishable manner at the latest at the time of the first interaction or exposure. The information will respect the applicable accessibility requirements.

Authorized Representative

If the organization, as a provider, is established outside the Union, the organization will appoint an authorized representative established in the Union by written mandate prior to making the AI system available on the union market (or when an importer cannot be identified).

The authorized representative will perform the tasks specified in the mandate received from the organization. The organization will provide a copy of the mandate to the market surveillance authorities upon request in one of the official languages of the institution of the Union determined by the national competent authority.

The mandate will empower the authorized representative to carry out the following tasks:

- keep at the disposal of the national competent authorities and national authorities for a period ending ten (10) years after the high-risk AI system has been placed on the market (or put into service), the contact details of the provider by which the authorized representative has been appointed, a copy of the EU declaration of conformity, the technical documentation and, if applicable, the certificate issued by the notified body;
- provide a national competent authority, upon a reasoned request, with all the information and documentation, including the above, necessary to demonstrate the conformity of a high-risk AI system with the requirements, including access to the logs, automatically generated by the high-risk AI system to the extent such logs are under the control of the organization;
- cooperate with competent authorities, upon a reasoned request, on any action the latter takes in relation to the high-risk AI system, in particular to reduce and mitigate the risks posed by the high-risk AI system; and
- where applicable, comply with the registration obligations, or, if the registration is carried out by the organization itself, ensure the information is correct.

The mandate will empower the authorized representative to be addressed, in addition to or instead of the provider, by the competent authorities, on all issues related to ensuring compliance with regulations.

The authorized representative may terminate the mandate if the authorized representative

considers (or has reason to consider) the organization is acting contrary to the organization's obligations under regulations. In such a case, the authorized representative will also immediately inform the market surveillance authority in which the authorized representative is established, as well as, where applicable, the relevant notified body, about the termination of the mandate and the reasons thereof.



AI Governance Policy

Human Oversight

Defines mechanisms for effective human supervision of AI systems.

Human Oversight

The organization will design and develop high-risk AI systems in such a way to include appropriate human-machine interface tools to be effectively overseen by natural persons during the period the AI system is in use. Human oversight will aim at preventing (or minimizing) risks to health, safety, or fundamental rights emerging when a high-risk AI system is being used according to its intended purposes and under conditions of reasonably foreseeable misuse particularly when such risks persists.

Human oversight will be ensured through one (1) or all of the following before placing the high-risk AI system on the market (or put into service):

- when technically feasible, identified and built into the high-risk AI system by the organization; and/or
- identified by the organization and appropriately implemented by the user.

Individuals assigned to human oversight responsibilities must:

- to properly understand the relevant capacities and limitations of the high-risk AI system and be able to duly monitor its operation where signs of anomalies, dysfunctions, and unexpected performance can be detected and addressed as soon as possible;
- remain aware of the possible tendency of automatically relying (or over-relying) on the output produced by a high-risk AI system ('automation bias'), in particular for high-risk AI systems used to provide information (or recommendation)s for decisions to be taken by natural persons;
- be able to correctly interpret the high-risk AI system's output, taking into account the characteristics of the system and the interpretation tools/methods available;
- be able to decide, in any particular situation, not to use the high-risk AI system (or otherwise) disregard, override, or reverse the output of the high-risk AI system; and
- be able to intervene on the operation of the high-risk AI system or interrupt the system through a "stop" button (or a similar procedure).

For high-risk AI systems intended to be used for the 'real-time' and 'post' remote biometric identification of natural persons, no action or decision should be taken by the user on the basis of the identification resulting from the system unless the output has been verified/confirmed by at least two (2) natural persons with the necessary competence, training, and authority. The requirement for a separate verification by at least two (2) natural persons does NOT apply to high risk AI systems used for the purpose of law enforcement, migration, border control or asylum, in cases where regulations consider the application of this requirement to be disproportionate.



AI Governance Policy

Accuracy, Robustness, and Cybersecurity

Sets requirements for reliable, secure, and resilient AI systems.

Accuracy, Robustness, and Cybersecurity

The organization will design and develop high-risk AI systems to achieve an appropriate level of accuracy, robustness, and cybersecurity based on its intended purpose. The organization will design and develop high-risk AI systems to perform consistently with respect to accuracy, robustness, and cybersecurity throughout its lifecycle. The organization will abide by the benchmarks and measurement methodologies to measure the appropriate levels of accuracy and robustness as defined by metrology and benchmarking authorities.

The organization will declare levels of accuracy and relevant accuracy metrics of high-risk AI systems in accompanying instructions of use.

The organization will ensure high-risk AI systems are resilient against errors, faults, or inconsistencies occurring within the system (or operational environment) due to its interaction with natural persons (or other systems). The organization will take technical and organizational measures as appropriate.

The organization will achieve robustness of high-risk AI systems through technical redundancy solutions such as backup or fail-safe plans.

If the high-risk AI system continues to learn after being placed on the market (or put into service), the organization will ensure the high-risk AI system is developed to address possible biased outputs due to outputs being used as an input for future operations (i.e., 'feedback loops') along with implementing appropriate mitigation measures.

The organization will develop the high-risk AI systems to be resilient in regards to attempts by unauthorized third parties to alter use, outputs, or performance) by exploiting the system vulnerabilities. Technical solutions ensuring cybersecurity will be appropriate to relevant circumstances and risks. Technical solutions addressing AI specific vulnerabilities will include, where appropriate, measures to prevent, detect, respond to, resolve, and control attacks trying to manipulate the training dataset (i.e., 'data poisoning'), pre-trained components used in training (i.e., 'model poisoning'), inputs designed to cause the model to make a mistake (i.e., 'adversarial examples' or 'model evasion'), confidentiality attacks, or model flaws.



AI Governance Policy

Provider Obligations

Defines responsibilities of AI system providers under regulation.

Provider Obligations

The organization, as a provider of high-risk AI systems, will:

- ensure the high-risk AI systems are compliant with applicable requirements;
- indicate their name, registered trade name (or registered trade mark), the address at which they can be contacted on the high-risk AI system or, where not possible, on its packaging (or its accompanying documentation) as applicable;
- have a quality management system implemented;
- maintain and retain required documentation;
- when under the organization's control, keep the logs automatically generated by the high-risk AI systems;
- ensure the high-risk AI system undergoes the relevant conformity assessment procedure prior to its placing on the market (or putting into service);
- draw up an EU declaration of conformity as required;
- affix the CE marking to the high-risk AI system to indicate conformity with regulations;
- comply with the registration obligations;
- take the necessary corrective actions and provide information as required;
- upon a reasoned request of regulators, demonstrate the conformity of the high-risk AI system with requirements; and
- ensure the high-risk AI system complies with accessibility requirements.

Corrective Actions

The organization, as a provider, will ensure all necessary actions are taken to bring the AI system into compliance with the requirements and regulatory obligations. Where the organization does not bring the AI system into compliance with requirements/regulations within the required period, the organization will be subject to fines.

The organization will ensure all appropriate corrective actions are taken in respect of all the AI systems concerned the organization has made available on the market throughout the Union. Where, in the course of an evaluation by the market surveillance authority, establishes the AI system was misclassified by the organization as not high-risk to circumvent the application of requirements the organization will be subject to fines.



AI Governance Policy

Further Processing of Personal Data in AI Regulatory Sandbox

Regulates secondary data use within AI regulatory sandboxes.

Further Processing of Personal Data in AI Regulatory Sandbox

In the AI regulatory sandbox personal data lawfully collected for other purposes may be processed solely for the purposes of developing, training and testing certain AI systems in the sandbox when all of the following conditions are met:

- AI systems will be developed for safeguarding substantial public interest by a public authority (or another natural/legal person governed by public/private law) and in one (or more) of the following areas:
 - o public safety and public health (including disease detection, diagnosis prevention, control and treatment, and improvement of healthcare systems);
 - o a high level of protection and improvement of the quality of the environment, protection of biodiversity, pollution as well as green transition, climate change mitigation and adaptation;
 - o energy sustainability;
 - o safety and resilience of transport systems and mobility, critical infrastructure, and networks;
 - o efficiency and quality of public administration and public services;
- the data processed are necessary for complying with one (or more) of the requirements referred to in Title III, Chapter 2 where those requirements cannot be effectively fulfilled by processing anonymized, synthetic, or other non-personal data;
- there are effective monitoring mechanisms to identify if any high risks to the rights and freedoms of the data subjects, may arise during the sandbox experimentation as well as response mechanism to promptly mitigate those risks and, where necessary, stop the processing;
- any personal data to be processed in the context of the sandbox are in a functionally separate, isolated, and protected data processing environment under the control of the prospective provider and only authorized persons have access to the data;
- providers can only further share the originally collected data in compliance with EU data protection law. Note: Any personal data created in the sandbox cannot be shared outside the sandbox;
- any processing of personal data in the context of the sandbox do not lead to measures or decisions affecting the data subjects nor affect the application of their rights laid down in Union law on the protection of personal data;
- any personal data processed in the context of the sandbox are protected by means of appropriate technical and organizational measures and deleted once the participation in the sandbox has terminated or the personal data has reached the end of its retention period;

- the logs of the processing of personal data in the context of the sandbox are kept for the duration of the participation in the sandbox, unless provided otherwise by Union or national law;
- complete and detailed description of the process and rationale behind the training, testing and validation of the AI system is kept together with the testing results as part of the technical documentation;
- a short summary of the AI project developed in the sandbox, its objectives and expected results published on the website of the competent authorities. Note: This obligation does not cover sensitive operational data in relation to the activities of law enforcement, border control, immigration, or asylum authorities.



AI Governance Policy

Testing of High-Risk AI Systems in Real World Conditions (Outside of AI Regulatory Sandboxes)

Defines conditions for testing high-risk AI systems outside sandboxes.

Testing of High-Risk AI Systems in Real World Conditions (Outside of AI Regulatory Sandboxes)

The organization, as a provider or prospective provider of high-risk AI systems, may test AI systems in real world conditions outside AI regulatory sandboxes according to the EU AI Act and real-world testing plan without violating the EU AI Act.

The detailed elements of the real world testing plan will be specified by the Commission according to examination procedures and will be without prejudice to Union (or national law) for the testing in real world conditions of high-risk AI systems related to products covered by legislation listed in the Union harmonization legislation.

The organization, as a provider or prospective provider, may conduct testing of high-risk AI systems in real world conditions at any time before the placing on the market (or putting into service) of the AI system on their own or in partnership with one (or more) prospective deployers. The testing of high-risk AI systems in real world conditions under will be without prejudice to ethical review required by national or Union law.

The organization, as a provider or prospective provider, may conduct the testing in real world conditions only where all of the following conditions are met:

- has drawn up a real world testing plan and submitted it to the market surveillance authority in the Member State(s) where the testing in real world conditions is to be conducted;
- the market surveillance authority in the Member State(s) where the testing in real world conditions is to be conducted has approved the testing in real world conditions and the real world testing plan. Where the market surveillance authority in that Member State has not provided with an answer in thirty (30) days, the testing in real world conditions and the real world testing plan shall be understood as approved. In cases where national law does not foresee a tacit approval, the testing in real world conditions shall be subject to an authorization;
- with the exception of biometrics, law enforcement, and migration, asylum, and border control management along with critical infrastructure high-risk AI systems has registered the testing in real world conditions in the non-public part of the EU database with a Union-wide unique single identification number and the information specified as follows:
 - o Union-wide unique single identification number of the testing in real world conditions;
 - o Name and contact details of the provider or prospective provider and users involved

- o in the testing in real world conditions;
 - o A brief description of the AI system, its intended purpose and other information necessary for the identification of the system;
 - o A summary of the main characteristics of the plan for testing in real world conditions;
 - o Information on the suspension or termination of the testing in real world conditions
- conducting the testing in real world conditions is established in the Union or it has appointed a legal representative who is established in the Union;
- data collected and processed for the purpose of the testing in real world conditions shall only be transferred to third countries outside the Union provided appropriate and applicable safeguards under Union law are implemented;
- the testing in real world conditions does not last longer than necessary to achieve its objectives and in any case not longer than six (6) months, which may be extended for an additional amount of six (6) months, subject to prior notification by the provider to the market surveillance authority, accompanied by an explanation on the need for such time extension;
- persons belonging to vulnerable groups due to their age, physical or mental disability are appropriately protected;
- where the organization organizes the testing in real world conditions in cooperation with one or more prospective deployers, the latter have been informed of all aspects of the testing relevant to their decision to participate, and given the relevant instructions on how to use the AI system; the organization and the deployer(s) shall conclude an agreement specifying their roles and responsibilities with a view to ensuring compliance with the provisions for testing in real world conditions under this Regulation and other applicable Union and Member States legislation;
- the subjects of the testing in real world conditions have given informed consent, or in the case of law enforcement, where the seeking of informed consent would prevent the AI system from being tested, the testing itself and the outcome of the testing in the real world conditions shall not have any negative effect on the subject and his or her personal data shall be deleted after the test is performed;
- the testing in real world conditions is effectively overseen by the organization and deployer(s) with persons who are suitably qualified in the relevant field and have the necessary capacity, training and authority to perform their tasks;
- the predictions, recommendations or decisions of the AI system can be effectively reversed and disregarded.

Any subject of the testing in real world conditions (or his/her legally designated representative, as appropriate) may, without any resulting detriment and without having to provide any justification, withdraw from the testing at any time by revoking his/her informed consent and request the immediate/permanent deletion of their personal data. The withdrawal of the informed consent will not affect the activities already carried out.

Any serious incident identified in the course of the testing in real world conditions will be reported

to the national market surveillance authority. The organization, as a provider or prospective provider, will adopt immediate mitigation measures or, failing that, suspend the testing in real world conditions until such mitigation takes place (or otherwise terminate it). The organization, as a provider or prospective provider, will establish a procedure for the prompt recall of the AI system upon such termination of the testing in real world conditions.

The organization, as a provider or prospective provider, will notify the national market surveillance authority in the Member State(s) where the testing in real world conditions is to be conducted of the suspension or termination of the testing in real world conditions and the final outcomes. The organization, as a provider or prospective provider, will be liable under applicable Union and Member States liability legislation for any damage caused in the course of their participation in the testing in real world conditions.



AI Governance Policy

Testing in Real-World Conditions

Establishes safeguards for real-world AI system testing.

Testing in Real-World Conditions

For the purpose of testing in real world conditions, informed consent shall be freely given by the subject of testing prior to his/her participation in such testing and after having been duly informed with concise, clear, relevant, and understandable information regarding:

- the nature and objectives of the testing in real world conditions and the possible inconvenience linked to his/her participation;
- the conditions under which the testing in real world conditions is to be conducted, including the expected duration of the subject's participation;
- the subject's rights and guarantees regarding participation, in particular his or her right to refuse to participate in and the right to withdraw from testing in real world conditions at any time without any resulting detriment and without having to provide any justification;
- the modalities for requesting the reversal or the disregard of the predictions, recommendations or decisions of the AI system; and
- the Union-wide unique single identification number of the testing in real world conditions and the contact details of the provider (or its legal representative) from whom further information can be obtained.

The informed consent will be dated and documented and a copy will be given to the subject or his/her legal representative.



AI Governance Policy

Quality Management System

Defines quality controls across the AI system lifecycle.

Quality Management System

The organization, as a provider of high-risk AI systems, will implement a quality management system (QMS) in accordance to ISO 9001 standards to ensure compliance with this policy and related regulations. The QMS will be documented in a systematic and orderly manner through written policies, procedures, and instructions as well as include at least the following aspects:

- a strategy for regulatory compliance, including compliance with conformity assessment procedures and procedures for the management of modifications to the high-risk AI system;
- techniques, procedures, and systematic actions to be used for the design, design control, and design verification of the high-risk AI system;
- techniques, procedures, and systematic actions to be used for the development, quality control, and quality assurance of the high-risk AI system;
- examination, test, and validation procedures to be carried out before, during, and after the development of the high-risk AI system as well as the frequency with which they have to be carried out;
- technical specifications, including standards, to be applied as well as where the relevant harmonized standards are not applied in full (or do not cover all of the relevant requirements) the means to be used to ensure the high-risk AI system complies with requirements;
- systems and procedures for data management, including data acquisition, data collection, data analysis, data labelling, data storage, data filtration, data mining, data aggregation, data retention, and any other operation regarding the data performed before and for the purposes of the placing on the market (or putting into service) of high-risk AI systems;
- the risk management system;
- the setting-up, implementation, and maintenance of a post-market monitoring system;
- procedures related to the reporting of serious incidents and of malfunctioning;
- the handling of communication with national competent authorities, competent authorities, including those providing (or supporting) the access to data, notified bodies, other operators, customers, or other interested parties;
- systems and procedures for record keeping of all relevant documentation and information;
- resource management, including security of supply related measures; and
- an accountability framework setting out the responsibilities of the management and other staff with regard to all aspects listed above.

The implementation of these aspects will be proportionate to the size of the organization. The organization will, in any event, respect the degree of rigor and the level of protection required to ensure compliance of their AI systems with regulations.



AI Governance Policy

Codes of Conduct

Encourages voluntary commitments to responsible AI practices.

Codes of Conduct

Codes of conduct may be drawn up by individual organizations (acting as providers or deployers of AI systems). Codes of conduct may be drawn up by organizations representing providers or deployers of AI systems (or by both), including with the involvement of deployers and any interested stakeholders along with their representative organizations (including civil society organizations and academia).

Codes of conduct may cover one or more AI systems taking into account the similarity of the intended purpose of the relevant systems.



AI Governance Policy

Conformity Assessment

Describes procedures to assess compliance of AI systems.

Conformity Assessment

The organization will ensure high-risk AI systems will undergo relevant conformity assessment procedures prior to placing on the market (or putting into service). Where the compliance of the AI systems meet requirements and has been demonstrated following a conformity assessment, the organization will draw up an EU declaration of conformity and affix the CE marking of conformity accordingly.

If the organization demonstrated compliance of a high-risk AI system applying harmonized standards or common specifications, where applicable, the organization will follow one of these procedures:

- the following conformity assessment procedure based on internal control:
 - o the organization, as a provider, verifies the established quality management system is in compliance;
 - o the organization, as a provider, examines the information contained in the technical documentation in order to assess the compliance of the AI system with the relevant essential requirements; and
 - o the organization, as a provider, also verifies the design and development process of the AI system and its post-market monitoring is consistent with the technical documentation;
 - o the conformity assessment procedure based on assessment of the quality management system and assessment of the technical documentation, with the involvement of a notified body,

If the organization demonstrated compliance of a high-risk AI system NOT applying (or applying only in part) harmonized standards (or where such harmonized standards do not exist or where standards exist, but the organization, as a provider, has not applied them) and common specifications will follow the conformity assessment procedures based on Annex VII. The organization may choose to use a notified body; however, if the system is intended to be put into service by law enforcement, immigration, or asylum authorities as well as EU institutions, bodies (or agencies), the market surveillance authority, as applicable, shall act as a notified body.

For the following high-risk AI systems, the organization will follow this conformity procedure based on internal controls, which does not provide for the involvement of a notified body:

- a. Management and operation of critical infrastructure
- b. Education and vocational training
- c. Recruitment: Employment and workers management, access to self-employment
- d. Access to and enjoyment of essential private services and public services/benefits

- e. Law enforcement
- f. Migration, asylum, and border control management
- g. Administration of justice and democratic processes

The following conformity assessment procedure based on internal control:

- the organization verifies the established quality management system is in compliance;
- the organization examines the information contained in the technical documentation in order to assess the compliance of the AI system with the relevant essential requirements; and
- the organization also verifies the design and development process of the AI system and its post-market monitoring is consistent with the technical documentation.

For high-risk AI systems, to which legal acts listed in Annex II, section A, apply, the organization will follow the relevant conformity assessment as required under those legal acts. The requirements will apply to those high-risk AI systems and will be part of that assessment. Points 4.3., 4.4., 4.5. and the fifth paragraph of point 4.6 of Annex VII will also apply.

For the purpose of that assessment, notified bodies which have been notified under those legal acts will be entitled to control the conformity of the high-risk AI systems with the requirements, provided the compliance of those notified bodies with requirements laid down has been assessed in the context of the notification procedure under those legal acts.

Where the legal acts listed in Annex II, section A, enable the manufacturer of the product to opt out from a third-party conformity assessment, provided the manufacturer has applied all harmonized standards covering all the relevant requirements, the manufacturer may make use of that option only if harmonized standards (or, where applicable, common specifications) covering the requirements have also been applied.

High-risk AI systems will undergo a new conformity assessment procedure whenever the systems are substantially modified, regardless of whether the modified system is intended to be further distributed or continues to be used by the current user.

For high-risk AI systems continuing to learn after being placed on the market (or put into service) where changes are done on the high-risk AI system as well as its performance have been pre-determined by the organization at the moment of the initial conformity assessment and part of the information contained in the technical documentation will NOT constitute a substantial modification.

Presumption of Conformity

Taking into account the intended purpose, high-risk AI systems having been trained and tested on data concerning the specific geographical, behavioral, contextual, and functional setting within

which they are intended to be used are presumed to be in compliance with the requirements. High-risk AI systems having been certified (or for which a statement of conformity has been issued under a cybersecurity scheme pursuant to Regulation (EU) 2019/881 of the European Parliament and of the Council and the references of which have been published in the Official Journal of the European Union) are presumed to be in compliance with the cybersecurity requirements so far as the cybersecurity certificate (or statement of conformity or parts thereof) cover those requirements.



AI Governance Policy

Declaration of Conformity

Formal statement confirming AI system regulatory compliance.

Declaration of Conformity

The organization, as a provider, will draw up a written machine readable (physical or electronically signed) EU declaration of conformity for each high-risk AI system. The EU declaration of conformity will identify the high-risk AI system for which it has been drawn up. A copy of the EU declaration of conformity will be submitted to the relevant national competent authorities upon request. The organization will retain this declaration of conformity for ten (10) years after the high-risk AI system has been placed on the market (or put into service) and make it available to the national competent authorities.

The EU declaration of conformity will state the high-risk AI system meets requirements. The EU declaration of conformity will contain the following and be translated into a language easily understood by the national competent authorities of the Member State(s) in which the high-risk AI system is placed on the market (or made available):

- AI system name and type and any additional unambiguous reference allowing identification and traceability of the AI system;
- Name and address of the provider or, where applicable, their authorized representative;
- A statement the EU declaration of conformity is issued under the sole responsibility of the provider;
- A statement the AI system in question is in conformity with the EU AI Act and, if applicable, with any other relevant Union legislation that provides for the issuing of an EU declaration of conformity;
- Where an AI system involves the processing of personal data, a statement that that AI system complies with the GDPR;
- References to any relevant harmonized standards used or any other common specification in relation to which conformity is declared;
- Where applicable, the name and identification number of the notified body, a description of the conformity assessment procedure performed and identification of the certificate issued; and
- Place and date of issue of the declaration, name and function of the person who signed it as well as an indication for, and on behalf of whom, that person signed, signature

Where high-risk AI systems are subject to other Union harmonization legislation (which also requires an EU declaration of conformity) a single EU declaration of conformity will be drawn up in respect of all Union legislations applicable to the high-risk AI system. The declaration will contain all the information required for identification of the Union harmonization legislation to which the declaration relates.

By drawing up the EU declaration of conformity, the organization, as a provider, will assume responsibility for compliance with the requirements. The organization will keep the EU declaration of conformity up-to-date as appropriate.



AI Governance Policy

CE Marking of Conformity

Defines requirements for CE marking of compliant AI systems.

CE Marking of Conformity

The organization will affix the CE marking visibly, legibly, and indelibly for high-risk AI systems. Where this is not possible (or warranted based on the nature of the high-risk AI system), the CE marking will be affixed to the packaging (or to the accompanying documentation) as appropriate. The CE marking will be subject to general principles set out in Article 30 of Regulation (EC) No 765/2008. For high-risk AI systems provided digitally, a digital CE marking will be used, only if it can be easily accessed via the interface from which the AI system is accessed (or via an easily accessible machine-readable code or other electronic means).

Where applicable, the CE marking will be followed by the identification number of the notified body responsible for the conformity assessment procedures. The identification number of the notified body will be affixed by the body itself or, under its instructions, by the provider or by its authorized representative. The identification number will also be indicated in any promotional material, which mentions the high-risk AI system fulfils the requirements for CE marking. Where high-risk AI systems are subject to other Union law (which also provides for the affixing of the CE marking) the CE marking will indicate the high-risk AI system also fulfils the requirements of the other law.

Before placing on the market (or putting into service) a high-risk AI system, the organization (or, where applicable, the authorized representative), except for critical infrastructure systems, will register the system in the EU database.

Before placing on the market (or putting into service) an AI system for which the provider has concluded it is not high-risk in application, the organization (as a provider) or, where applicable, the authorized representative will register themselves and the system in the EU database. Before putting into service (or using a high-risk AI system) listed in Annex III, with the exception of critical infrastructure systems, deployers who are public authorities, agencies or bodies (or persons acting on their behalf) will register themselves, select the system and register its use in the EU database.

For high-risk AI systems (including biometrics, law enforcement, and migration, asylum, and border control management), the registration will be done in a secure non-public section of the EU database and include only certain information, as applicable:
Critical infrastructure high-risk AI systems will be registered at the national level.



AI Governance Policy

Cooperation with Competent Authorities

Establishes obligations to cooperate with regulatory authorities.

Cooperation with Competent Authorities

The organization, as a provider of high-risk AI systems, upon request by a national competent authority, will provide all the information and documentation necessary to demonstrate the conformity of the high-risk AI system with the requirements of this policy and related regulations. The information/documentation will be provided in an official Union language determined by the Member State concerned.

Upon a reasoned request from a national competent authority, the organization will also give the authority access to the logs automatically generated by the high-risk AI system, to the extent such logs are under the organization's control. Information obtained by a national competent authority will be treated in compliance with confidentiality obligations established by regulations.



AI Governance Policy

Product Manufacturer Obligations

Defines AI-related responsibilities of product manufacturers.

Product Manufacturer Obligations

Where a high-risk AI system is placed on the market (or put into service together with other manufacture products) related to products to which the legal acts listed in Annex II, section A, apply, the manufacturer of the product will take responsibility over compliance with the AI system and related regulation. As far as the AI system is concerned, the same obligations imposed by the regulations on a provider are assumed by the manufacturer.



AI Governance Policy

Importer Obligations

Specifies compliance duties for AI system importers.

Importer Obligations

Before placing a high-risk AI system on the market, the organization, as an importer of such system, will ensure:

- the appropriate conformity assessment procedure has been carried out by the provider of the AI system;
- the provider has drawn up the technical documentation accordingly;
- the system bears the required conformity marking and is accompanied by the required documentation and instructions of use; and
- the provider has appointed an authorized representative.

Where an importer considers (or has reason to consider) a high-risk AI system is NOT in conformity with regulations, the importer will not place the system on the market until the AI system has been brought into conformity. Where the high-risk AI system presents a risk, the importer will inform the provider of the AI system and the market surveillance authorities to that effect.

Importers will indicate their name, registered trade name (or registered trade mark), and the address at which they can be contacted on the high-risk AI system (or, where that is not possible,) on its packaging (or its accompanying documentation), as applicable.

Importers will ensure, while a high-risk AI system is under their responsibility, (where applicable), storage or transport conditions do not jeopardize its compliance with requirements. Importers will keep, for a period ending ten (10) years after the AI system has been placed on the market (or put into service), a copy of the certificate issued by the notified body, where applicable, of the instructions for use and of the EU declaration of conformity.

Importers will provide national competent authorities, upon a reasoned request, with all necessary information and documentation to demonstrate the conformity of a high-risk AI system with requirements in a language, which can be easily understood by the national competent authority. The importer will ensure technical documentation can be made available to authorities.

The importer will also cooperate with those authorities on any action national competent authority takes in relation to the system to reduce and mitigate the risks posed by the high-risk AI system.

Any importer will be considered a provider for regulatory purposes and will be subject to the obligations of the provider, in any of the following circumstances:

- the importer places their name (or trademark) on a high-risk AI system already placed on

the market (or put into service), without prejudice to contractual arrangements stipulating the obligations are allocated otherwise;

- the importer makes a substantial modification to a high-risk AI system having already been placed on the market (or has already been put into service) and in a way it remains a high-risk AI system;
- the importer modifies the intended purpose of an AI system (including a general purpose AI system) which has not been classified as high-risk and has already been placed on the market (or put into service) in such manner the AI system becomes a high risk AI system.

Where these circumstances occur, the provider initially placing the high-risk AI system on the market (or putting it into service) will no longer be considered a provider for regulatory purposes. Where these circumstances occur, the provider initially placing the AI system on the market (or putting it into service) will no longer be considered a provider a provider for regulatory purposes. This former provider will closely cooperate and will make available the necessary information and provide the reasonably expected technical access and other assistance required for the fulfilment of the obligations set out in regulations, in particular regarding the compliance with the conformity assessment of high-risk AI systems.

This does not apply in the cases where the former provider has expressly excluded the change of its system into a high-risk system and therefore the obligation to hand over the documentation. For high-risk AI systems being a safety component of products to which the legal acts listed in Annex II, section A apply, the manufacturer of those products will be considered the provider of the high-risk AI system and will be subject to the obligations under either of the following scenarios: the high-risk AI system is placed on the market together with the product under the name (or trademark) of the product manufacturer and the high-risk AI system is put into service under the name (or trademark) of the product manufacturer after the product has been placed on the market.

The provider of a high-risk AI system and the third party supplying an AI system, tools, services, components, or processes used (or integrated) in a high-risk AI system will, by written agreement, specify the necessary information, capabilities, technical access and other assistance (based on the generally acknowledged state of the art) in order to enable the provider of the high-risk AI system to fully comply with the obligations set out in regulations. Note: This obligation does not apply to third parties making accessible to the public tools, services, processes, or AI components other than general-purpose AI models under a free and open license. (See the AI Office's model contractual terms for additional information.)

The above is without prejudice to the need to respect and protect intellectual property rights and confidential business information (or trade secrets) according to regulations.



AI Governance Policy

Distributors Obligations

Defines responsibilities of AI system distributors.

Distributors Obligations

Before making a high-risk AI system available on the market, distributors will verify the high-risk AI system bears the required CE conformity marking, it is accompanied by the required documentation and instruction of use, and the provider (and the importer of the system, as applicable), have complied with the obligations set out in the regulations.

Where a distributor considers (or has reason to consider) a high-risk AI system is NOT in conformity with requirements, the distributor will not make the high-risk AI system available on the market until the system has been brought into conformity with those requirements. Furthermore, where the system presents a risk, the distributor will inform the provider (or the importer of the system, as applicable), to that effect.

Distributors will ensure, while a high-risk AI system is under their responsibility, where applicable, storage or transport conditions do not jeopardize the compliance of the system with the requirements.

A distributor considering (or has reason to consider) a high-risk AI system, which it has made available on the market, is NOT in conformity with the requirements will take the corrective actions necessary to bring the system into conformity with those requirements, to withdraw it (or recall it), or will ensure the provider, the importer, or any relevant operator, as appropriate, takes those corrective actions. Where the high-risk AI system presents a risk, the distributor will immediately inform the national competent authorities of the Member States in which it has made the product available to that effect, giving details, in particular, of the non-compliance and of any corrective actions taken.

Upon a reasoned request from a national competent authority, distributors of high-risk AI systems will provide that authority with all the information and documentation necessary to demonstrate the conformity of a high-risk system with the requirements. Distributors will also cooperate with that national competent authority on any action taken by the authority.

Distributors will cooperate with national competent authorities on any action those authorities take in relation to an AI system, of which they are the distributor, in particular to reduce or mitigate the risk posed by the high-risk AI system.

Any distributor will be considered a provider for regulatory purposes and will be subject to the obligations of the provider, in any of the following circumstances:

- the distributor places their name (or trademark) on a high-risk AI system already placed on the market (or put into service), without prejudice to contractual arrangements stipulating the obligations are allocated otherwise;

- the distributor makes a substantial modification to a high-risk AI system having already been placed on the market (or has already been put into service) and in a way it remains a high-risk AI system;
- the distributor modifies the intended purpose of an AI system (including a general purpose AI system) which has not been classified as high-risk and has already been placed on the market (or put into service) in such manner the AI system becomes a high risk AI system.

Where these circumstances occur, the provider initially placing the high-risk AI system on the market (or putting it into service) will no longer be considered a provider for regulatory purposes. Where these circumstances occur, the provider initially placing the AI system on the market (or putting it into service) will no longer be considered a provider a provider for regulatory purposes. This former provider will closely cooperate and will make available the necessary information and provide the reasonably expected technical access and other assistance required for the fulfilment of the obligations set out in regulations, in particular regarding the compliance with the conformity assessment of high-risk AI systems.

This does not apply in the cases where the former provider has expressly excluded the change of its system into a high-risk system and therefore the obligation to hand over the documentation. For high-risk AI systems being a safety component of products to which the legal acts listed in Annex II, section A apply, the manufacturer of those products will be considered the provider of the high-risk AI system and will be subject to the obligations under either of the following scenarios: the high-risk AI system is placed on the market together with the product under the name (or trademark) of the product manufacturer and the high-risk AI system is put into service under the name (or trademark) of the product manufacturer after the product has been placed on the market.

The provider of a high-risk AI system and the third party supplying an AI system, tools, services, components, or processes used (or integrated) in a high-risk AI system will, by written agreement, specify the necessary information, capabilities, technical access and other assistance (based on the generally acknowledged state of the art) in order to enable the provider of the high-risk AI system to fully comply with the obligations set out in regulations. Note: This obligation does not apply to third parties making accessible to the public tools, services, processes, or AI components other than general-purpose AI models under a free and open license. (See the AI Office's model contractual terms for additional information.)

The above is without prejudice to the need to respect and protect intellectual property rights and confidential business information (or trade secrets) according to regulations.



AI Governance Policy

Deployers Obligations

Outlines obligations for entities deploying AI systems.

Deployers Obligations

Deployers of high-risk AI systems will take appropriate technical and organizational measures to ensure they use such systems according to instructions of use accompanying the systems.

To the extent deployers exercise control over the high-risk AI system, the deployer will ensure the natural persons assigned to ensure human oversight of the high-risk AI systems have the necessary competence, training and authority as well as the necessary support. These obligations are without prejudice to other deployer obligations under regulations and to the deployer's discretion in organizing its own resources and activities for the purpose of implementing the human oversight measures indicated by the provider.

To the extent the deployer exercises control over the input data, the deployer will ensure input data is relevant and sufficiently representative in view of the intended purpose of the high-risk AI system.

Deployers will monitor the operation of the high-risk AI system on the basis of the instructions of use and when relevant, inform providers accordingly.

When deployers have reasons to consider the use in accordance with the instructions of use may result in the AI system presenting a risk, the deployer will, without undue delay, inform the provider (or distributor) and relevant market surveillance authority and suspend the use of the system.

The deployer will also immediately inform the provider and the importer (or distributor) and relevant market surveillance authorities when they have identified any serious incident.

Deployers of high-risk AI systems will keep the logs automatically generated by the high-risk AI system to the extent such logs are under their control for a period appropriate to the intended purpose of the high-risk AI system, of at least six (6) months, unless provided otherwise in applicable regulations, in particular in Union law on the protection of personal data.

Prior to putting into service or use a high-risk AI system at the workplace, deployers who are employers will inform workers representatives and the affected workers they will be subject to the system. This information shall be provided, where applicable, according to rules and procedures laid down in regulations and practice on information of workers and their representatives.

Deployers of high-risk AI systems being public authorities (or Union institutions, bodies, offices and agencies) will comply with the registration obligations. If the deployer finds the system they envisage to use has not been registered in the EU database, the deployer will not use the system and will inform the provider or the distributor.

Where applicable, deployers of high-risk AI systems will use the information provided to users to comply with their obligation to carry out a data protection impact assessment under the GDPR. In the framework of an investigation for the targeted search of a person convicted (or suspected) of having committed a criminal offense, the deployer of an AI system for post-remote biometric identification will request an authorization, prior, or without undue delay and no later than forty-eight (48) hours, by a judicial authority (or an administrative authority) whose decision is binding and subject to judicial review, for the use of the system, (except when the system is used for the initial identification of a potential suspect based on objective and verifiable facts directly linked to the offence).

Each use will be limited to what is strictly necessary for the investigation of a specific criminal offense. If the requested authorization provide is rejected, the use of the post remote biometric identification system linked to authorization will be stopped with immediate effect and the personal data linked to the use of the system for which the authorization was requested will be deleted.

In any case, such AI system for post remote biometric identification will not be used for law enforcement purposes in an untargeted way, without any link to a criminal offense, a criminal proceeding, a genuine and present (or genuine and foreseeable) threat of a criminal offense or the search for a specific missing person.

It will be ensured no decision producing an adverse legal effect on a person may be taken by the law enforcement authorities solely based on the output of these post remote biometric identification systems. Regulations regarding the processing of biometric data under GDPR are still in effect.

Regardless of the purpose or deployer, each use of these systems will be documented in the relevant police file and will be made available to the relevant market surveillance authority and the national data protection authority upon request, excluding the disclosure of sensitive operational data related to law enforcement.

Deployers will, in addition, submit annual reports to the relevant market surveillance and national data protection authorities on the uses of post-remote biometric identification systems, excluding the disclosure of sensitive operational data related to law enforcement.

The reports can be aggregated to cover several deployments in one operation. Note: Member States may introduce, in accordance with Union law, more restrictive laws on the use of post remote biometric identification systems.

Deployers of high-risk AI systems referred to in Annex III making decisions (or assist in making decisions) related to natural persons will inform the natural persons they are subject to the use of the high-risk AI system.

Deployers will cooperate with the relevant national competent authorities on any action those authorities take in relation with the high-risk system in order to implement the EU AI Act.

Any deployers will be considered a provider for regulatory purposes and will be subject to the obligations of the provider, in any of the following circumstances:

- the deployers places their name (or trademark) on a high-risk AI system already placed on the market (or put into service), without prejudice to contractual arrangements stipulating the obligations are allocated otherwise;
- the deployers makes a substantial modification to a high-risk AI system having already been placed on the market (or has already been put into service) and in a way it remains a high-risk AI system;
- the deployers modifies the intended purpose of an AI system (including a general purpose AI system) which has not been classified as high-risk and has already been placed on the market (or put into service) in such manner the AI system becomes a high risk AI system.

Where these circumstances occur, the provider initially placing the high-risk AI system on the market (or putting it into service) will no longer be considered a provider for regulatory purposes. Where these circumstances occur, the provider initially placing the AI system on the market (or putting it into service) will no longer be considered a provider a provider for regulatory purposes. This former provider will closely cooperate and will make available the necessary information and provide the reasonably expected technical access and other assistance required for the fulfilment of the obligations set out in regulations, in particular regarding the compliance with the conformity assessment of high-risk AI systems.

This does not apply in the cases where the former provider has expressly excluded the change of its system into a high-risk system and therefore the obligation to hand over the documentation. For high-risk AI systems being a safety component of products to which the legal acts listed in Annex II, section A apply, the manufacturer of those products will be considered the provider of the high-risk AI system and will be subject to the obligations under either of the following scenarios: the high-risk AI system is placed on the market together with the product under the name (or trademark) of the product manufacturer and the high-risk AI system is put into service under the name (or trademark) of the product manufacturer after the product has been placed on the market.

The provider of a high-risk AI system and the third party supplying an AI system, tools, services, components, or processes used (or integrated) in a high-risk AI system will, by written agreement, specify the necessary information, capabilities, technical access and other assistance (based on the generally acknowledged state of the art) in order to enable the provider of the high-risk AI system to fully comply with the obligations set out in regulations. Note: This obligation does not apply to third parties making accessible to the public tools, services, processes, or AI components other than general-purpose AI models under a free and open license. (See the AI Office's model contractual terms for additional information.)

The above is without prejudice to the need to respect and protect intellectual property rights and confidential business information (or trade secrets) according to regulations.

Right to Explanation of Individual Decision-Making

Any affected person subject to a decision which is taken by the organization, as a deployer, on the basis of the output from an high-risk AI system (with the exception of critical infrastructure systems), and which produces legal effects or similarly significantly affects him/her in a way considered to adversely impact their health, safety and fundamental rights will have the right to request from the organization clear and meaningful explanations on the role of the AI system in the decision-making procedure as well as the main elements of the decision taken. *Note: This right does not apply to the use of AI systems for which exceptions from, or restrictions to, the obligation follow from Union or national law in compliance with Union law.*



AI Governance Policy

Fundamental Rights Impact Assessment

Assesses AI impacts on fundamental rights and freedoms.

Fundamental Rights Impact Assessment

Prior to deploying a high-risk AI system (with the exception of AI systems intended to be used in critical infrastructure), deployers that are bodies governed by public law (or private operators providing public services) and operators deploying high-risk systems (for the evaluation of creditworthiness or used for risk assessment/pricing for life/health insurance), will perform an assessment of the impact on fundamental rights that the use of the system may produce. For this purpose, deployers will perform an assessment consisting of the following:

- a description of the deployer's processes in which the high-risk AI system will be used in line with its intended purpose;
- a description of the period of time and frequency in which each high-risk AI system is intended to be used;
- the categories of natural persons and groups likely to be affected by its use in the specific context;
- the specific risks of harm likely to impact the categories of persons (or group of persons) identified, taking into account the information given by the provider in the instruction of use;
- a description of the implementation of human oversight measures, according to the instructions of use; and
- the measures to be taken in case of the materialization of these risks, including their arrangements for internal governance and complaint mechanisms.

This obligation applies to the first use of the high-risk AI system. The deployer may, in similar cases, rely on previously conducted fundamental rights impact assessments or existing impact assessments carried out by provider. If, during the use of the high-risk AI system, the deployer considers any of the factors listed to have changed (or no longer up to date), the deployer will take the necessary steps to update the information.

Once the impact assessment has been performed, the deployer will notify the market surveillance authority of the results of the assessment, submitting the filled template as a part of the notification.

If any of these obligations are already met through the data protection impact assessment conducted under the GDPR, the fundamental rights impact assessment will be conducted in conjunction with the data protection impact assessment.

Note: *The AI Office shall develop a template for a questionnaire (including through an automated tool) to facilitate deployers to implement the obligations in a simplified manner.*



AI Governance Policy

Other Third-Party Obligations

Defines obligations for other actors in the AI value chain.

Other Third-Party Obligations

Any other third-party will be considered a provider for regulatory purposes and will be subject to the obligations of the provider, in any of the following circumstances:

- the other third-party places their name (or trademark) on a high-risk AI system already placed on the market (or put into service), without prejudice to contractual arrangements stipulating the obligations are allocated otherwise;
- the other third-party makes a substantial modification to a high-risk AI system having already been placed on the market (or has already been put into service) and in a way it remains a high-risk AI system;
- the other third-party modifies the intended purpose of an AI system (including a general purpose AI system) which has not been classified as high-risk and has already been placed on the market (or put into service) in such manner the AI system becomes a high risk AI system.

Where these circumstances occur, the provider initially placing the high-risk AI system on the market (or putting it into service) will no longer be considered a provider for regulatory purposes. Where these circumstances occur, the provider initially placing the AI system on the market (or putting it into service) will no longer be considered a provider a provider for regulatory purposes. This former provider will closely cooperate and will make available the necessary information and provide the reasonably expected technical access and other assistance required for the fulfilment of the obligations set out in regulations, in particular regarding the compliance with the conformity assessment of high-risk AI systems.

This does not apply in the cases where the former provider has expressly excluded the change of its system into a high-risk system and therefore the obligation to hand over the documentation. For high-risk AI systems being a safety component of products to which the legal acts listed in Annex II, section A apply, the manufacturer of those products will be considered the provider of the high-risk AI system and will be subject to the obligations under either of the following scenarios: the high-risk AI system is placed on the market together with the product under the name (or trademark) of the product manufacturer and the high-risk AI system is put into service under the name (or trademark) of the product manufacturer after the product has been placed on the market.

The provider of a high-risk AI system and the third party supplying an AI system, tools, services, components, or processes used (or integrated) in a high-risk AI system will, by written agreement, specify the necessary information, capabilities, technical access and other assistance (based on the generally acknowledged state of the art) in order to enable the provider of the high-risk AI system to fully comply with the obligations set out in regulations. Note: This obligation does not apply to third parties making accessible to the public tools, services, processes, or AI components

other than general-purpose AI models under a free and open license. (See the AI Office's model contractual terms for additional information.)

The above is without prejudice to the need to respect and protect intellectual property rights and confidential business information (or trade secrets) according to regulations.



AI Governance Policy

Post-Market Monitoring

Establishes ongoing monitoring after AI system deployment.

Post-Market Monitoring

The organization will establish and document a post-market monitoring system in a manner proportionate to the nature of the artificial intelligence technologies and the risks of the high-risk AI system.

The post-market monitoring system will actively and systematically collect, document, and analyze relevant data provided by users or collected through other sources on the performance of high-risk AI systems throughout its lifetime as well as allow the organization to evaluate the continuous compliance of AI systems with requirements. Where relevant, post-market monitoring will include an analysis of the interaction with other AI systems. Note: This obligation does not cover sensitive operational data of deployers which are law enforcement authorities.

The post-market monitoring system will be based on a post-market monitoring plan. The post-market monitoring plan will be part of the technical documentation.

For high-risk AI systems covered by the legal acts referred to in Annex II, Section A, where a post-market monitoring system and plan is already established under the legislation, the elements described will be integrated into the system and plan as appropriate.

Reporting of Concerns

The organization will define and implement a process to report concerns about the organization's role with respect to an AI system throughout its life cycle. The reporting mechanisms will provide the following functions:

- options for confidentiality, anonymity (or both) to provide for effective protection from reprisals for both the persons concerned with reporting and investigation (e.g. by allowing reports to be made anonymously and confidentially);
- available and promoted to employed/contracted persons;
- staffed with qualified persons and stipulates appropriate investigation/resolution powers for these qualified persons;
- provides for mechanisms to report and to escalate to management in a timely manner;
- provides reports and maintain confidentiality/anonymity respecting general business confidentiality considerations; and
- provides response mechanisms within an appropriate time frame.

Reporting of Serious Incidents

The organization will report any serious incident to the market surveillance authorities of the Member States where that incident occurred. The period of reporting will take account of the severity of the serious incident. Notification will be made immediately after the organization, as a provider, has established a causal link between the AI system and the serious incident (or the reasonable likelihood of such a link) and in any event, not later than fifteen (15) days after the organization (or where applicable, the deployer), becomes aware of the serious incident. In the event of a widespread infringement (or a serious incident) the report will be provided immediately, and not later than two (2) days after the organization, as a provider, (or where applicable, the deployer) becomes aware of the incident.

In the event of death of a person, the report will be provided immediately after the organization, as a provider, (or the deployer) has established (or as soon as it suspects) a causal relationship between the high-risk AI system and the serious incident, but not later than ten (10) days after the date on which the organization (or where applicable, the deployer) becomes aware of the serious incident.

Where necessary to ensure timely reporting, the organization, as a provider (or where applicable, the deployer) may submit an initial report as incomplete followed up by a complete report. Following the reporting of a serious incident, the organization, as a provider, will (without delay) perform the necessary investigations in relation to the serious incident and the AI system concerned. This will include a risk assessment of the incident and corrective action. The organization will co-operate with the competent authorities and where relevant with the notified body concerned during the investigations. The organization will not perform any investigation which involves altering the AI system concerned in a way which may affect any subsequent evaluation of the causes of the incident, prior to informing the competent authorities of such action.

For high-risk AI systems referred to in Annex III placed on the market (or put into service) by the organization, as a provider, subject to Union legislative instruments laying down reporting obligations equivalent to those set out in the EU AI Act, the notification of serious incidents will be limited.

For high-risk AI systems, which are safety components of devices, (or are themselves devices) covered by Regulation (EU) 2017/745 and Regulation (EU) 2017/746 the notification of serious incidents will be limited and be made to the national competent authority chosen for this purpose by the Member States where the incident occurred.



AI Governance Policy

Access to Data and Documentation

Defines access rights to AI-related data and documentation.

Access to Data and Documentation

National public authorities (or bodies), which supervise (or enforce) the respect of obligations under Union law protecting fundamental rights (including the right to non-discrimination) in relation to the use of high-risk AI systems will have the power to request (and access) any documentation created (or maintained) under the regulations (in accessible language and format) when access to the documentation is necessary for the fulfilment of the competences under their mandate within the limits of their jurisdiction. The relevant public authority or body will inform the market surveillance authority of the Member State concerned of any such request.

Where the documentation is insufficient to ascertain whether a breach of obligations under Union law intended to protect fundamental rights has occurred, the public authority (or body) may make a reasoned request to the market surveillance authority to organize testing of the high-risk AI system through technical means. The market surveillance authority will organize the testing with the close involvement of the requesting public authority (or body) within reasonable time following the request.

Any information and documentation obtained by the national public authorities (or bodies) will be treated in compliance with confidentiality obligations.



Secure Your Compliance.
Secure the Contract.

Get Audit Ready for your next big deal.

We don't just provide templates; our team of experienced experts and auditors brings the deep regulatory knowledge needed to turn these pages into a passing audit score and your next big contract.

www.elevateconsult.com cmmc@elevateconsult.com +1 888 601 5351