

RESEARCH ARTICLE



AI Act in BPMN: More Than Just Words—A Perspective for Providers of AI Systems

Marinella Quaranta^{1,2,3,*} and Ilaria Angela Amantea¹

¹Department of Computer Science, University of Turin, Italy

²Department of Legal Studies, University of Bologna, Italy

³Vrije Universiteit Brussel, Belgium

Abstract: This article presents a practical and visual Business Process Modeling and Notation (BPMN)-based guide to support the assessment of artificial intelligence (AI) systems under the EU AI Act, with the aim of identifying their risk classification and the corresponding compliance obligations. The proposed approach structures the assessment process as a guided decision-making pathway, composed of a sequence of questions derived from the regulatory framework. Through this step-by-step process, providers of AI systems can systematically determine whether their AI systems fall within prohibited, high-risk, limited-risk, or minimal-risk categories and identify the applicable sets of obligations. The use of BPMN enables a transparent and intuitive representation of the assessment logic, making the methodology accessible to a wide range of AI systems' providers, including those without legal expertise. At the same time, it ensures traceability of each decision point and facilitates the identification of critical elements within the compliance process. A key feature of the proposed framework is its adaptability: the BPMN structure can be easily updated as new guidelines, technical standards, and implementing measures are issued under the AI Act. Furthermore, the methodology has been developed within an institutional setting, specifically the AI Factory IT4LIA, and the EuSAiR project, ensuring alignment with ongoing regulatory developments and practical relevance for real-world applications.

Keywords: AI Act, compliance, BPMN, practical guide

1. Introduction

The AI Act¹ is the European framework that aims at giving definitions, regulate general requirements of artificial intelligence (AI) systems, and establish the cornerstones of what is allowed, prohibited, and acceptable in AI practices.

However, while the AI Act seeks to foster the reliable and responsible deployment of AI, it is still uncertain to what extent these regulatory requirements are actually being put into practice by those involved in AI development and design [1–3].

In this evolving regulatory landscape, the European Commission has introduced supportive mechanisms such as AI Regulatory Sandboxes and AI Factories to foster harmonization, clarify interpretative uncertainties, and refine compliance practices. Nevertheless, substantial uncertainty persists for companies. Businesses

currently face a challenging position: they are legally required to comply with the AI Act, yet many technical and operational aspects remain unclear or incomplete, even at the legislative and supervisory levels. This contribution seeks to address that gap by providing a practical and operational tool for organizations.

Through the use of Business Process Modeling (BPM) methodology, we have developed a structured decision-making flow in Business Process Modeling and Notation (BPMN) standard language. The scheme is designed to guide companies in assessing the risk classification of their AI systems and identifying the corresponding obligations under the regulatory framework as it currently stands. BPMN constitutes an accessible and intuitively interpretable modeling language, requiring no specialized legal background to understand its structure. Its modular design also ensures extensibility, allowing straightforward updates as forthcoming implementing measures and technical regulations are adopted.

This work complements a previously published handbook: “AI Compliance Frameworks: A Handbook For Compliance Tools based on AI Acts’ Requirements” [4] that provides detailed guidance on how to fulfill the AI Act’s compliance requirements across the AI lifecycle. Together, these two instruments offer a comprehensive compliance-support paradigm. It has been developed in the Italian AI Factory IT4LIA project².

*Corresponding author: Marinella Quaranta, Department of Computer Science, University of Turin, Italy, Department of Legal Studies, University of Bologna, Italy and Vrije Universiteit Brussel, Belgium. Email: marinella.quaranta@unito.it

¹Regulation (EU) 2024/1689 of the European Parliament and of the Council of June 13, 2024, laying down harmonized rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139, and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797, and (EU) 2020/1828 (Artificial Intelligence Act).

²<https://it4lia-aiactory.eu/it/>

Also, this work aims to contribute to the field of legal visualization for regulatory compliance [5–7]. BPMN promotes the transparent interpretation and representation of regulatory provisions, including highly complex legislation. More recently, the European Union has been moving toward a broader vision of regulatory simplification³. This structured decision-making flow is not directly meant to harmonize all the international frameworks on risk assessment, whereas its goal is to give a clear structure for existing EU regulation on AI compliance. Therefore, this structure only includes the AI Act and the Guidelines published so far. Finally, the scope of this work only includes the “provider” in its definition of Article 3(3)⁴. The work acknowledges that the AI Act also addresses obligations for other entities in the AI production chain (deployers, distributors, etc.). However, they are outside the scope of this framework as it is deemed to be a guide exclusively for providers.

This article is divided in the following sections: Section 2 presents the legal context and related works, providing information on some existing tools and on the BPM methodology along with the BPMN Standard Language, Section 3 is the core of the article, the AI Act in BPMN is presented and explained through various processes and sub-processes, Section 4 supports the BPMN structure by giving the reader further information on obligations included in the scheme. Section 5 presents two case studies to test and validate the presented framework, and Section 6 is the concluding section.

2. Legal Background and Related Work

2.1. Legal background

The uncertainty linked to the generality of norms set by the AI Act led to the publication of further documents, giving information and details about specific topics (e.g., Guidelines on AI system⁵, Guidelines on prohibited AI practices⁶, etc.). The AI Act itself set deadlines for the publication of Commission Guidelines (Article 72, Article 6, and others). However, part of these deadlines has not been respected, and the Commission issued a new document to postpone them and proposed further solutions: the Digital Omnibus⁷.

To support innovation within its broad regulatory framework, the AI Act introduces AI Regulatory Sandboxes, explicitly provided for in Article 57: a “space for experimentation” where innovators can test compliance with the Regulation and other applicable legal standards. Within this framework, authorities provide guidance, supervision, and support, while monitoring potential risks to health, safety, and fundamental rights and requiring appropriate mitigation measures where necessary. The objective of these AI Regulatory Sandboxes is to promote

regulatory learning and facilitate market access, particularly for small and medium-sized enterprises (SMEs) and start-ups developing innovative AI systems. Moreover, within the framework of the European strategy for Artificial Intelligence, AI Factories are also emerging as advanced infrastructures aimed at transforming Europe’s innovation ecosystem. An AI Factory is a high-performance computing (HPC) infrastructure equipped with advanced computational resources, high-quality datasets, and state-of-the-art software tools, designed to support the large-scale development and experimentation of AI models. It is not merely a matter of computing power; rather, AI Factories function as integrated platforms combining computational capacity, data access, development environments, and qualified human support. The European Commission promotes these infrastructures as “innovation hubs” where companies, start-ups, research centers, and public administrations can co-develop AI solutions, test algorithms, access sovereign data, and build competencies within an open and interoperable environment.

These AI Factories operate within a broader ecosystem, integrating with European Digital Innovation Hubs (EDIHs), European Digital Infrastructure Consortia, and national policies. The work presented in this article has been conducted within the framework of the Italian AI Factory IT4LIA in connection with the Italian AI Regulatory Sandbox⁸. The initiative offers integrated services spanning the entire AI lifecycle, including access to curated data repositories, scalable AI development environments, pre-trained foundation models, testing and validation facilities, and compliance support aligned with the AI Act. IT4LIA also provides test-before-invest services, proof-of-concept funding for start-ups and SMEs, and trustworthiness assessments through bias detection, security audits, and regulatory guidance. Dedicated open-source AI models will be developed for strategic sectors such as healthcare, manufacturing, agrifood, climate, finance, and smart cities, facilitating domain-specific innovation. Also, IT4LIA promotes knowledge transfer through training programs, workshops, hackathons, internships, and an AI networking platform that fosters collaboration between academia, industry, and public institutions.

In the context of the Italian Sandboxes, the EuSAIR project is an initiative included in the Digital Europe program. Its goal is connected to the AI Factory: it aims at supporting and harmonizing the implementation of AI Regulatory Sandboxes across all 27 EU Member States, as mandated by the AI Act. The EuSAIR project and two case studies selected as part of our collaboration are presented and discussed in detail in Section 5.

2.2. Related work

Considering that the AI Act, by its nature as a European Regulation, is inherently general, that it represents the first legislative instrument to impose detailed and far-reaching limitations on AI systems, and that several implementing measures are still pending or delayed, it becomes evident that the current compliance landscape is complex, fragmented, uncertain, and lacking full transparency, mostly to non-specialists. This is especially true for companies that develop AI systems: although they possess strong expertise in their respective technical or industrial domains, they often lack specialized legal knowledge necessary to interpret an evolving and complex regulatory framework.

Beyond being regulatory mechanisms, AI Regulatory Sandboxes and AI Factories may also be considered as the first

³<https://www.consilium.europa.eu/en/policies/simplification/>

⁴“Provider” means a natural or legal person, public authority, agency, or other body that develops an AI system or a general-purpose AI model or that has an AI system or a general-purpose AI model developed and places it on the market or puts the AI system into service under its own name or trademark, whether for payment or free of charge.

⁵Commission Guidelines on the definition of an artificial intelligence system established by Regulation (EU) 2024/1689 (AI Act).

⁶Commission Guidelines on prohibited artificial intelligence practices established by Regulation (EU) 2024/1689 (AI Act).

⁷Regulation (EU) 2026/1744 of the European Parliament and of the Council of 8 July 2026 amending Regulations (EU) 2024/1689 and (EU) 2018/1139 as regards the simplification of the implementation of harmonised rules on artificial intelligence, OJ L, 2026/1744, 24 July 2026.

⁸<https://eusair-project.eu/consortium/unibo/>

technical instruments made available to support implementation and compliance.

Indeed, the European Commission presents them not merely as conceptual or policy-oriented instruments but as practical, operational platforms designed to work in direct collaboration with companies. Their purpose extends beyond regulatory reflection: they actively support businesses in complying with requirements, testing solutions in controlled environments, and addressing concrete implementation challenges. In this sense, AI Factories function as hands-on support mechanisms, facilitating regulatory clarity, fostering innovation, and enabling structured interaction between regulators and market actors within a supervised and cooperative framework.

Furthermore, some more technical compliance instruments are beginning to emerge. Foremost among them, and currently the official initiative, is the “EU AI Act Compliance Checker”⁹, introduced by the European Commission.

This interactive tool leads stakeholders through a sequence of yes/no questions to determine which obligations under the AI Act are relevant to them. In addition, the checker customizes its guidance according to the user’s role (such as provider, deployer, distributor, etc.), so that each stakeholder is clearly informed about their specific responsibilities.

While such an interactive compliance checker represents a useful preliminary tool, its explanatory capacity remains limited. Systems based on a sequence of binary (yes/no) questions can assist various stakeholders in the AI ecosystem, such as providers, deployers, distributors, and importers, in identifying the likely risk classification of an AI system and the corresponding obligations under the AI Act. In this sense, they provide an accessible entry point for understanding the regulatory framework and may support an initial assessment of compliance requirements, particularly for actors with limited legal or technical expertise. However, this approach often lacks transparency regarding the reasoning process underlying the outcome. Even when users retrace the sequence of questions, it may remain unclear which specific legal provisions or criteria certain obligations have been assigned to a given use case.

Our framework, while pursuing a similar objective, namely, the classification of AI systems under the AI Act, aims at offering a more structured and transparent framework. By explicitly representing the decision paths that lead to specific classifications and obligations, our model provides greater traceability of the reasoning process and may therefore serve as a clearer basis for further regulatory assessment and compliance analysis.

Another solution presented is a question-based decision tree framework to accelerate the understanding of the classification of an AI system [8]. Rather than exposing users to the full decision tree, the system guides them step by step, helping them classify AI systems based on risk levels. Still, this approach has limitations: users do not see the overall structure, which may reduce deeper understanding of the regulation, and the lack of a publicly available interactive tool limits practical adoption. As of now, the structure of the decision tree underlying the methodology has not been published or made publicly available.

Other scholars are proposing the use of ontologies to classify the risk of AI systems [9, 10]. The ontology-based approach is used to structure the EU AI Act using SPARQL to query and reason over its concepts, including risk levels, actors, and

obligations. By formalizing AI Act concepts in a machine-readable ontology, the approach enables systematic classification, automated compliance checks, and reasoning, while its modular design allows updates as guidelines evolve, and SPARQL supports precise, consistent queries. However, limitations include the difficulty of encoding legal nuance and context-dependent obligations, the effort required to update interdependent elements, and challenges in translating evolving interpretations into formal queries.

In parallel, a growing number of private actors are developing compliance checkers. Some are open access such as the German start-up Trail.ML¹⁰ or the interpretative table “EU AI Act Compliance Matrix”¹¹ of the International Association of Privacy Professionals. There are also several private and closed-access tools designed to support compliance with relevant provisions of the regulation. Examples of commercial solutions are PwC AI Compliance¹², Capgemini’s AI compliance infographic¹³, FairNow¹⁴, Diligent¹⁵, Lumenova¹⁶, and Modulos¹⁷. However, these tools are also largely structured as checklist-based instruments and, in the absence of a formally established certification authority and complete regulatory details, cannot issue legally valid conformity certificates.

It is important to clarify that any claim to provide a formal “AI Act compliance certification” at this stage is, strictly speaking, premature. No officially designated conformity assessment bodies are yet fully operational for all relevant categories, nor has the regulatory framework been completed through the adoption of the necessary implementing measures and technical standards. Consequently, any market offering purporting to deliver binding certification should be understood primarily as a commercial or advisory service rather than a legally recognized conformity assessment.

The added value of our approach lies in its institutional embedding and methodological structure. The tool is being developed within the framework of IT4LIA, as part of a structured flow that interfaces with the Italian AI Regulatory Sandbox. In both initiatives, the National Cybersecurity Authority (ACN) acts as a partner, and it is also designated to serve as the national supervisory authority for AI.

Unlike purely private compliance checkers, the AI Regulatory Sandbox has the capacity to issue reports with legal relevance. At present, such a report represents the closest available instrument to a formal compliance certification.

Our contribution should therefore be understood as a structured compliance-support tool: it guides organizations through a transparent decision-making pathway, is institutionalized within an official framework, and produces traceable and reviewable outputs, thereby allowing critical points to be clearly identified, reassessed, and updated as the regulatory landscape evolves.

¹⁰<https://www.trail-ml.com/>

¹¹<https://iapp.org/resources/article/eu-ai-act-compliance-matrix>

¹²<https://www.pwc.com/cz/en/sluzby/umela-intelligence-ai/ai-act-ai-compliance-tool.html>

¹³<https://www.capgemini.com/solutions/eu-ai-act-compliance/>

¹⁴<https://fairnow.ai/platform/>

¹⁵<https://www.diligent.com/>

¹⁶<https://www.lumenova.ai/platform/>

¹⁷<https://www.modulos.ai/modulos-ai-governance-platform/>

⁹<https://ai-act-service-desk.ec.europa.eu/en/eu-ai-act-compliance-checker>

2.3. BPM methodology and BPMN standard language

Business Process Management (BPM) [11] is widely recognized as a strategic approach for optimizing processes, even in complex environments like healthcare [12, 13] or courts [14].

In formal analytical terms, a statute may be conceptualized as a complex process to be followed. Legislative provisions typically establish a structured set of specific obligations, procedural requirements, and conditional pathways that must be observed by the relevant actors. Rather than prescribing a purely linear sequence of actions, legal norms often entail nonlinear procedural architectures, incorporating decision points, alternative routes, exceptions, and differentiated obligations depending on the factual circumstances under examination. Accordingly, the application of a legal framework can be understood as the execution of a regulated process, characterized by branching logic and context-dependent procedural variations.

A central focus of BPM is modeling and engineering processes so as to enable verification and validation by domain and system experts. This implies the development of formal visual representations, such as process maps or flowcharts, which depict sequences of activities and decision points. In this work, processes are represented using the Business Process Modeling and Notation (BPMN) standard [15]. BPM methodology and BPMN standard have already been applied to create maps of norms [16–18].

The strength of BPMN is to employ a set of standardized graphic elements: events (circles), activities or sub-processes (rectangles), gateways for decision or synchronization points (diamonds), sequence flows (arrows), and pools to allocate responsibilities across organizational units.

A business process model constitutes a coherent and self-contained arrangement of tasks aimed at achieving a specific organizational objective. It specifies the activities performed (“what”), their sequencing (“when”), the responsible actors (“who”), and the data involved (“how”). Process traces represent concrete executions of the model, derived either from design specifications or reconstructed through event logs.

Therefore, a process appears as in Figure 1. The process begins with the initial event, represented by the circle, followed by Activity 1. A control flow then leads to the decision point represented by the question “Q?”. If the answer is “NO”, the process proceeds directly to Activity 2. If the answer is “YES”, the process proceeds sequentially through Activity 3 and Activity 4. Finally, the process terminates with the end event, represented by the final circle. In this way, it is possible to obtain a process trace that refers to the sequence of tasks and events linked by the

connectors defined by the model. Referring to the process in Figure 1, trace examples can be:

- TRACE 1: Start-1-Q?-2-End
- TRACE 2: Start-1-Q?-3-4-End.

Thus, not only can a structured decision-making pathway be obtained, but full transparency can also be ensured at each individual stage of the process.

Moreover, compliance management represents a central dimension of BPM [19]. Organizations increasingly rely on BPM to ensure that operational workflows align with regulatory requirements, giving rise to structured Compliance Management Frameworks [20]. Given the evolving and often general nature of current regulatory provisions, BPM provides a structured methodology for mapping, assessing, and verifying compliance between operational workflows and applicable legal standards.

BPMN therefore not only enables the structured guidance of decision-making and the allocation of obligation packages based on specific characteristics or exceptions but also ensures full transparency at each individual stage of the decision-making pathway. Moreover, it ensures ease of use for stakeholders across diverse business domains, regardless of their specific area of expertise.

3. AI Act in BPMN

The AI Act is a complex framework that imposes obligations on all parties involved in the AI production chain, including providers, distributors, importers, and deployers, among others. This paper specifically focuses on the obligations of the provider.

In this section, the BPMN representation guides the assessment of an AI system under the AI Act. At the conclusion of the process, the outputs are both the risk-level classification of the AI system under consideration and the complete set of applicable obligation packages. These obligation packages are subsequently detailed in the final table (Table 1).

Figure 2 is the high-level process.

It is divided into five sub-processes: “*Checking the Scope*,” “*Checking the System*,” “*Checking the Model*,” “*Checking the Context*,” and “*Checking the Interaction with People*.”

Each sub-process follows a progressive logic that mirrors the sequence of the AI Act itself. This progression is designed to ensure that each step logically builds upon the previous one. The aim is to avoid duplication and harmonize obligations while focusing on increasingly specific regulatory criteria. The sub-processes

Figure 1
Example of a process

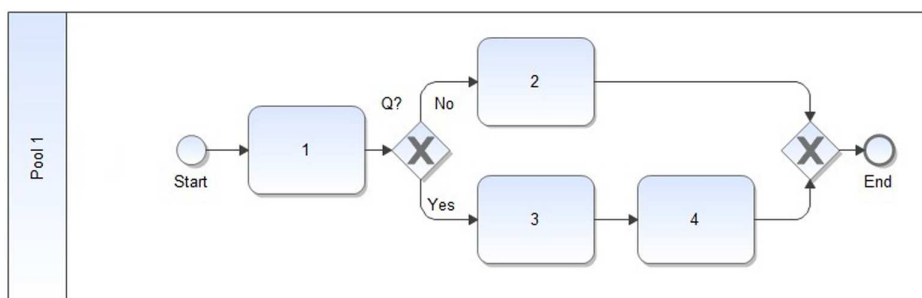


Figure 2
High-level process in BPMN for AI system assessment according to the AI Act Regulation



escalate in complexity, and the tasks are organized to get to a final “amount of obligations.”

Each sub-process determines relevant features, namely:

- “*Checking the Scope*” sub-process: this phase determines whether the system falls within or outside the scope of the AI Act with regard to specific uses. Its primary function is to exclude uses that are expressly outside the Regulation’s applicability.
- “*Checking the System*” sub-process: this phase assesses whether the system qualifies as an AI system under the AI Act. It serves to exclude systems that are too simple to meet the legal definition of AI.

Together, “*Checking the Scope*” and “*Checking the System*” sub-processes establish whether the technology under examination constitutes an AI system subject to the AI Act or not. From this point onward, the framework assigns the risk classifications and the corresponding packages of obligations, which may be cumulative where multiple conditions apply.

- “*Checking the Model*” sub-process: this phase verifies whether the system qualifies as a General-Purpose AI (GPAI) model. If so, the GPAI label is assigned, and the related obligations are identified.
- “*Checking the Context*” sub-process: this phase analyzes the intended context and purpose of use. It examines the core provisions of the AI Act, particularly Chapter III. This phase is fundamental, as it determines whether the system qualifies as a prohibited practice or as a high-risk system, thereby dictating the applicable obligations. Chapter III is central to the AI

Act Regulation because it defines differentiated obligations for high-risk systems and specifies relevant exceptions. All these obligations and exceptions are guided to be assessed in this sub-process.

- “*Checking the Interaction with People*” sub-process: this final stage determines whether the AI system falls within the 0-risk or minimal-risk category and assigns corresponding transparency or residual obligations.

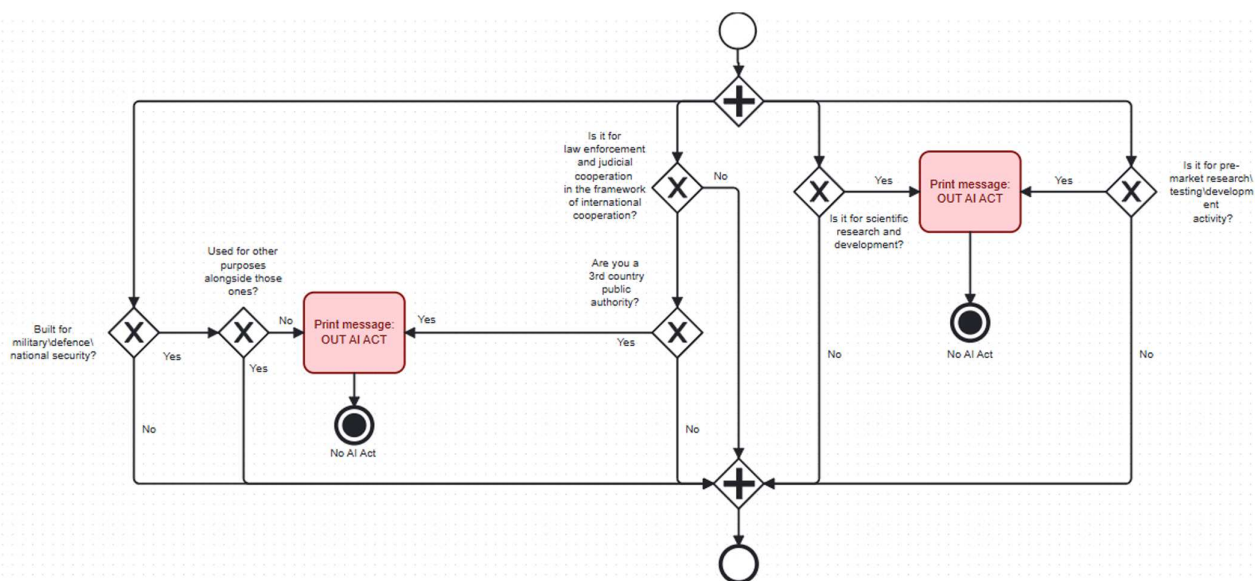
It is important to emphasize that the proposed structure is designed to allocate and accumulate obligations rather than to operate through mutually exclusive classifications. Our approach is precautionary and oriented toward ensuring the highest level of safeguards and guarantees: that is, an AI system is not framed as being merely GPAI or solely high risk; rather, it may simultaneously trigger multiple categories of obligations. The architecture of the framework explicitly reflects this additive logic, ensuring that all applicable obligations are identified and combined where relevant.

In each sub-process, the tasks in red represent output messages for the provider. They indicate either the assigned risk level or the exit from the assessment process.

Figure 3 represents the sub-process *Checking the Scope*. This sub-process aims at determining whether the system falls within or outside the scope of the AI Act.

To fall within the scope of the AI Act and proceed with the assessment, it is important that all four questions are answered negatively. Namely, the system must not be intended for use in military, defense, or national security contexts, nor for purposes of international cooperation or agreements for law enforcement and

Figure 3
Details of the first sub-process “*Checking the Scope*”



judicial cooperation, nor for scientific research and development, and nor for pre-market research, testing, or development activity. If even one answer is positive, the process ends, as the system falls outside the scope of the regulation. In this case, the message “OUT OF AI ACT” is delivered to the provider, and the assessment process is terminated (end of process “No AI Act”). Only in the “international cooperation/agreement” case, two Yeses lead to the process termination (end of process “No AI Act”).

Both this first sub-process (“Checking the Scope”) and the second sub-process (“Checking the System”) define whether a system falls under the scope of the AI Act or not. We decided to split the sub-processes for two reasons:

- The sub-process “Checking the Scope”:
- The sub-process Checking the System”:
 - Refers to Article 2 of the AI Act,
 - Aims at excluding some types of uses.
 - Refers to the Guidelines for the Definition of an AI system¹⁸,
 - Focuses on the complexity of the system in order to exclude some systems from the definition of AI (see below).

¹⁸Communication to the Commission Approval of the content of the draft Communication from the Commission – Commission Guidelines on the definition of an artificial intelligence system established by Regulation (EU) 2024/1689 (AI Act).

If the system falls within the scope of the AI Act, the process continues with the second sub-process: “Checking the System,” shown in Figure 4. It aims at assessing whether the considered system is complex enough to fall into the definition of an AI system or not. Categories to be analyzed in this section are described in the system guidelines. The sections of the system guidelines taken into consideration are (42–45) systems for improving mathematical optimization, (46–47) basic data processing, (48) systems based on classical heuristics, and (49–51) simple prediction systems.

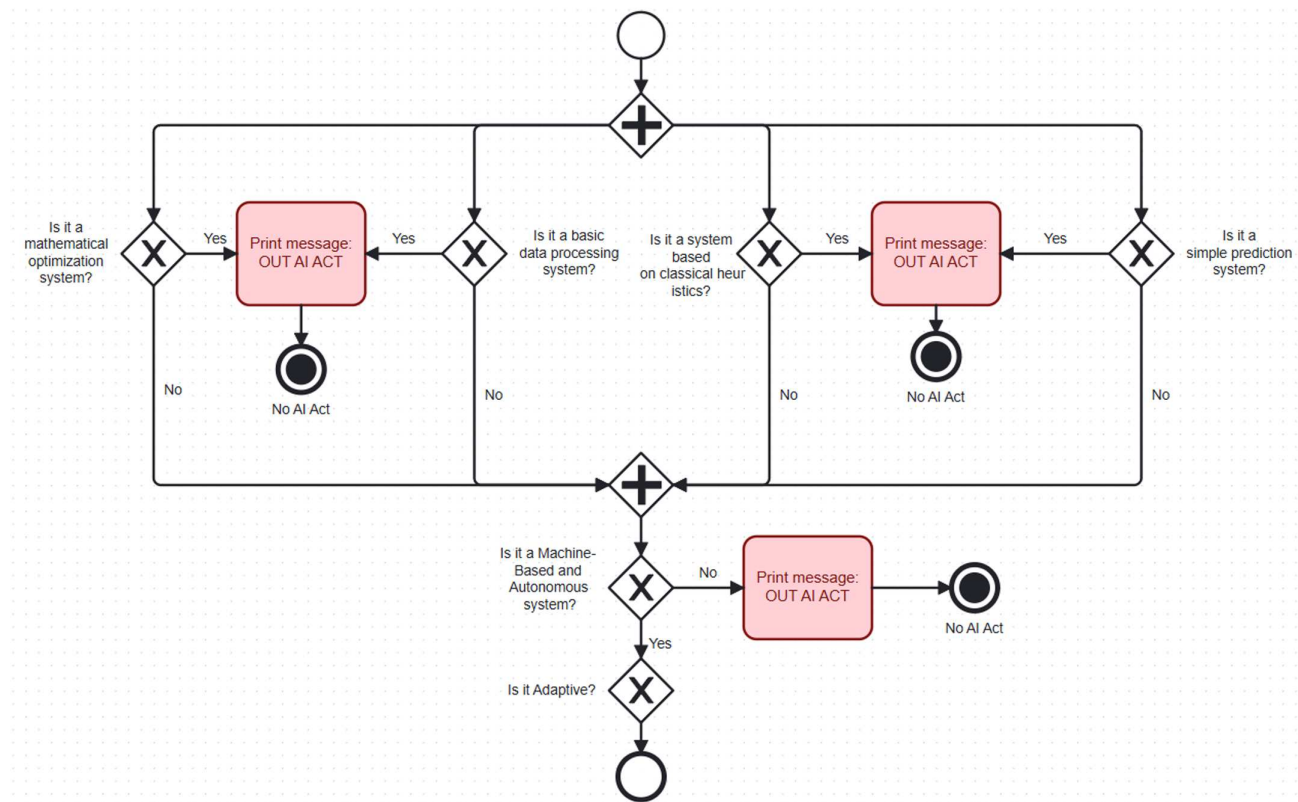
In this case, it is important to check four different cases simultaneously, and only if none of these cases apply can we proceed with the assessment. If one answer is positive, the process ends. In such a case, the system is too simple to be considered an AI system; thus, it will not fall under the AI regulation. The message “OUT OF AI ACT” is delivered to the provider, and the assessment process is terminated (end of process “No AI Act”).

After these specific questions, it is important to check whether the system is machine-based, autonomous, and adaptive. Without these characteristics, it cannot be considered an AI system.

We decided to isolate the final question as:

- The first four questions aim at excluding specific types of systems.
- The last question aims at checking the cumulative characteristics required for a system to be defined as an AI system.
- There are still doubts about the interpretations of machine-based, autonomous, and adaptive concepts related to the

Figure 4
Details of the second sub-process “Checking the System”



definition of AI. The isolation of such a question foresees the possibility of creating a new branch of the tree, starting from this gateway.

At this point, it has been established that the system under assessment falls within the scope of the AI Act and qualifies as an AI system. From this stage onward, the analysis will focus on identifying the obligations to be complied with.

First, we need to check the model, sub-process “*Checking the Model*,” shown in Figure 5. This sub-process aims at verifying whether the AI is a GPAI or not. To be classified as GPAI, the system has to serve a wide range of tasks, work on a large amount of data, and have the cumulative amount of computation used for its training that is greater than 10^{23} floating-point operations (FLOP). Once established that the system is a GPAI, a message is printed for the analysts (“*Print message: GPAI*”). Then, we need to establish which obligations the providers have to comply with:

- If the system is not open-access and free license, the provider has to comply with the obligations labeled “*GPAI obligation*” in Table 1.
- If the system is open-access and free license but does not provide substantial info related to the system, the provider has to comply with the obligations labeled “*GPAI open and free*” in Table 1.

In any case, it is necessary to verify whether the GPAI model has high-impact capabilities and whether the cumulative amount of computation used for its training is greater than 10^{25} FLOP. If this is the case, the AI system is a “*GPAI WITH SUBSTANTIAL RISK*,” the provider has to comply with (or also with) the obligations labeled “*GPAI obl with risk*” in Table 1. We want to stress that subsequent obligations are cumulative; for example, a provider of a system not open-access and free license with more than 10^{25} FLOPs has to comply with “*GPAI obligation*” and with the “*GPAI obl with risk*.”

At this stage of the assessment, it is necessary to determine whether the AI system may be placed on the European market or whether it falls within the category of prohibited practices (Figure 6)¹⁹.

Five main practices are banned:

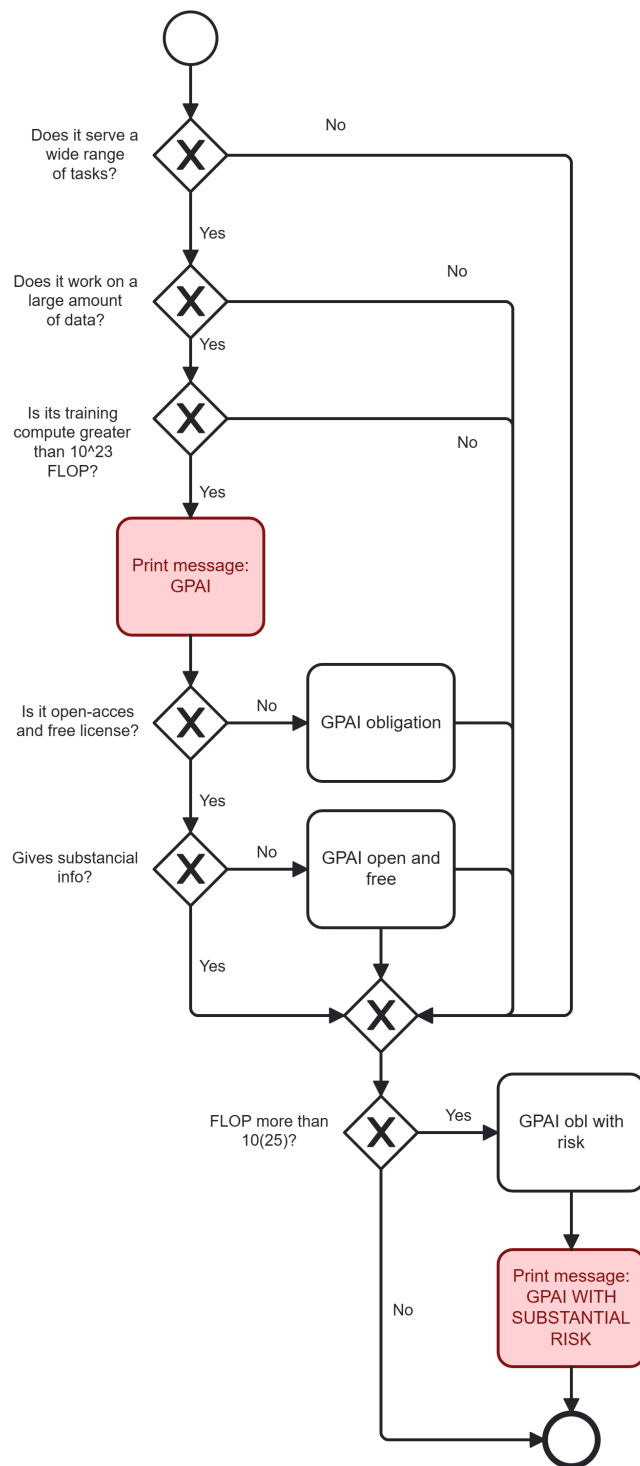
- Criminal offense prediction
- Creation of Images of Child Sexual Abuse Material or non-consensual intimate media (added by the Omnibus)
- Biometrics used to infer personal information
- Manipulation
- Emotion inference used in education or workplace

These categories are shown in Figure 7.

Some of them have exceptions. Exceptions may reclassify a prohibited (banned) system as high risk, thereby moving it into the second part of the process (“*Checking for High-Risk*”). It nevertheless remains essential to verify all possible cases and their corresponding exceptions. The process terminates only if the system ultimately falls within the banned category. A prohibited system is, in fact, excluded from the European market without any possibility of choice (“*Print message: BANNED*”).

Conversely, the checks will continue sequentially even where an exception applies, in order to prevent the system from bypassing the prohibition by falling within a prohibited category subject to an exception, while also falling within another prohibited

Figure 5
Details of the third sub-process “*Checking the Model*”



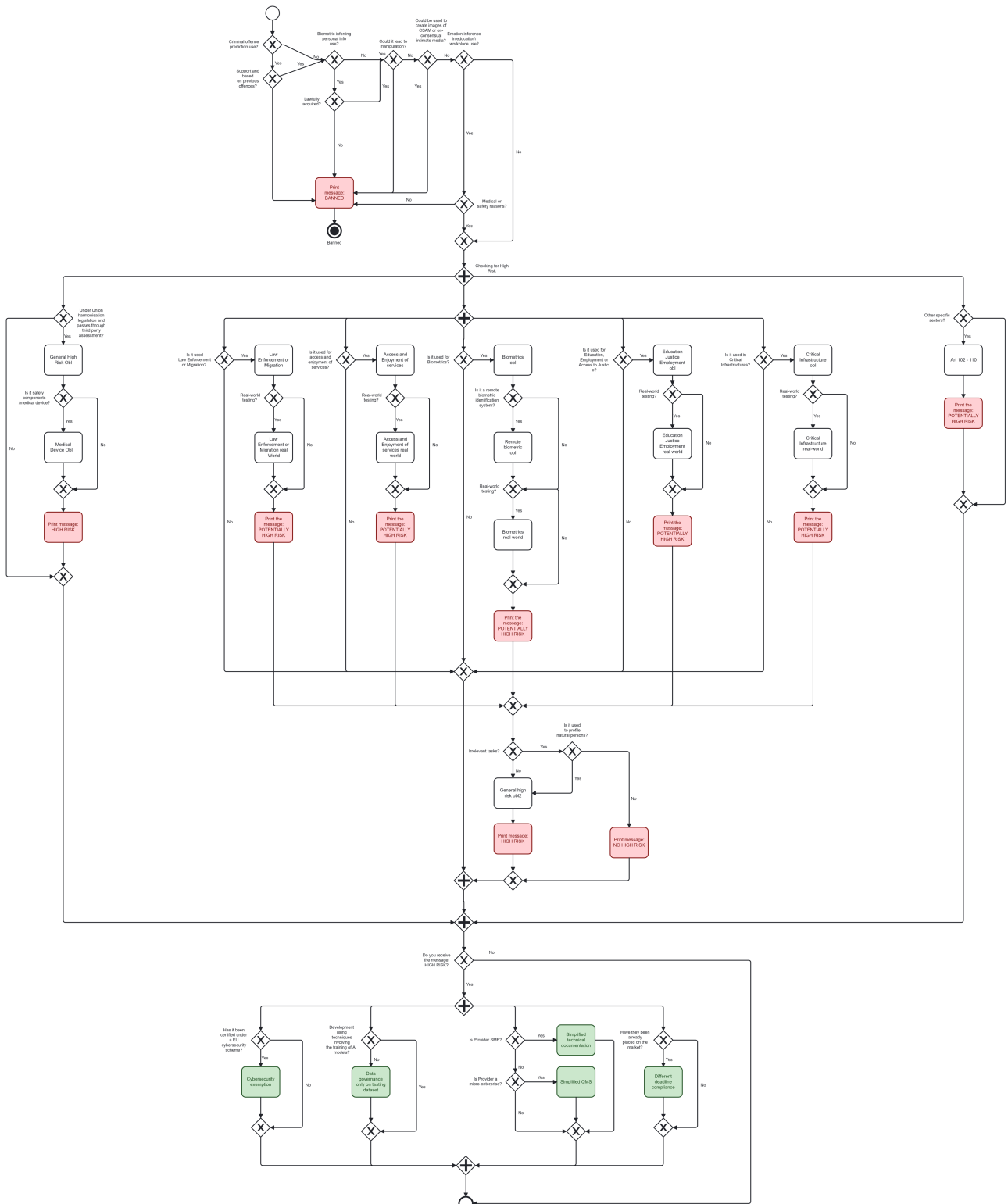
category for which no exception is provided. If the AI system does not fall within the category of prohibited practices, the assessment proceeds to determine whether it qualifies as a high-risk AI system.

An AI system must comply with the obligations applicable to high-risk systems if it falls within one of the following three categories:

- It falls under the Union harmonization legislation, and it is subject to third-party assessment – *Annex I of the AI Act*.

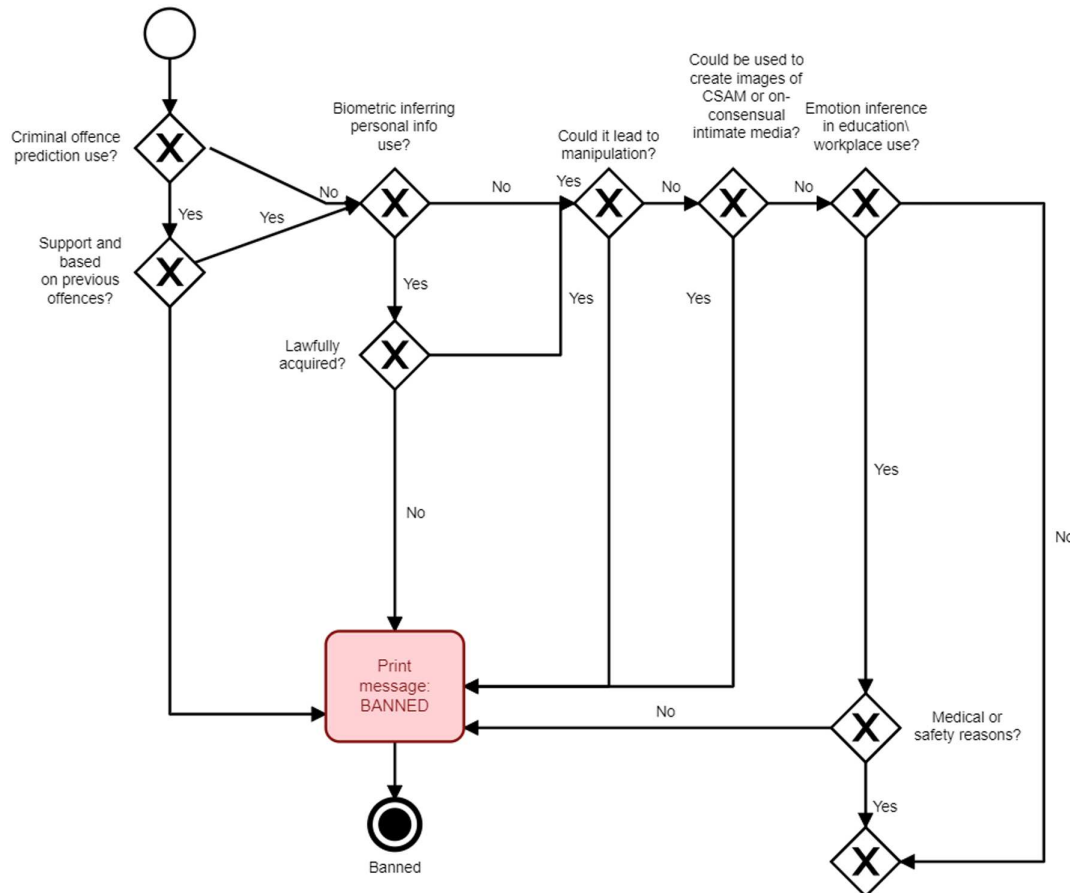
¹⁹Available at <https://github.com/Marinelele/Project-AI-Act-BPMN/blob/main/Context1.png>

Figure 6
Details of the fourth sub-process “Checking the Context”.



Note: This is the overall sub-process. Each branch is shown in Figures 7, 8, 9, 10, 11, 12, and 13.

Figure 7
Details of the fourth sub-process. The banned branch



- It falls within one of the uses provided in *Annex III of the AI Act*.
- It falls within one of the regulatory frameworks *amended by Articles 102–110 of the AI Act* (“*Other specific sectors?*”).

These categories are shown in Figure 8.

It is important to emphasize that an AI system falling within one of the first two categories is explicitly classified by the AI Act as a “high-risk system” and must therefore comply with the obligations set out in Chapter III “HIGH-RISK AI SYSTEMS.”

With regard to the third category, Articles 102–110 introduce amendments to specific pieces of EU legislation; these amendments prescribe that systems concerned in the amended legal frameworks shall comply with the obligations laid down in Chapter III (see Figure 9).

Such systems might not be explicitly labeled as “high-risk systems.” However, by the wording of the AI Act, it seems that these systems have to abide by rules for high-risk systems (Chapter III). Thus, the wording seems to imply that once an AI system is embedded in the amended regulations and directives (300/2008, 167/2013, etc.), it should at least follow all the norms prescribed for high-risk systems.

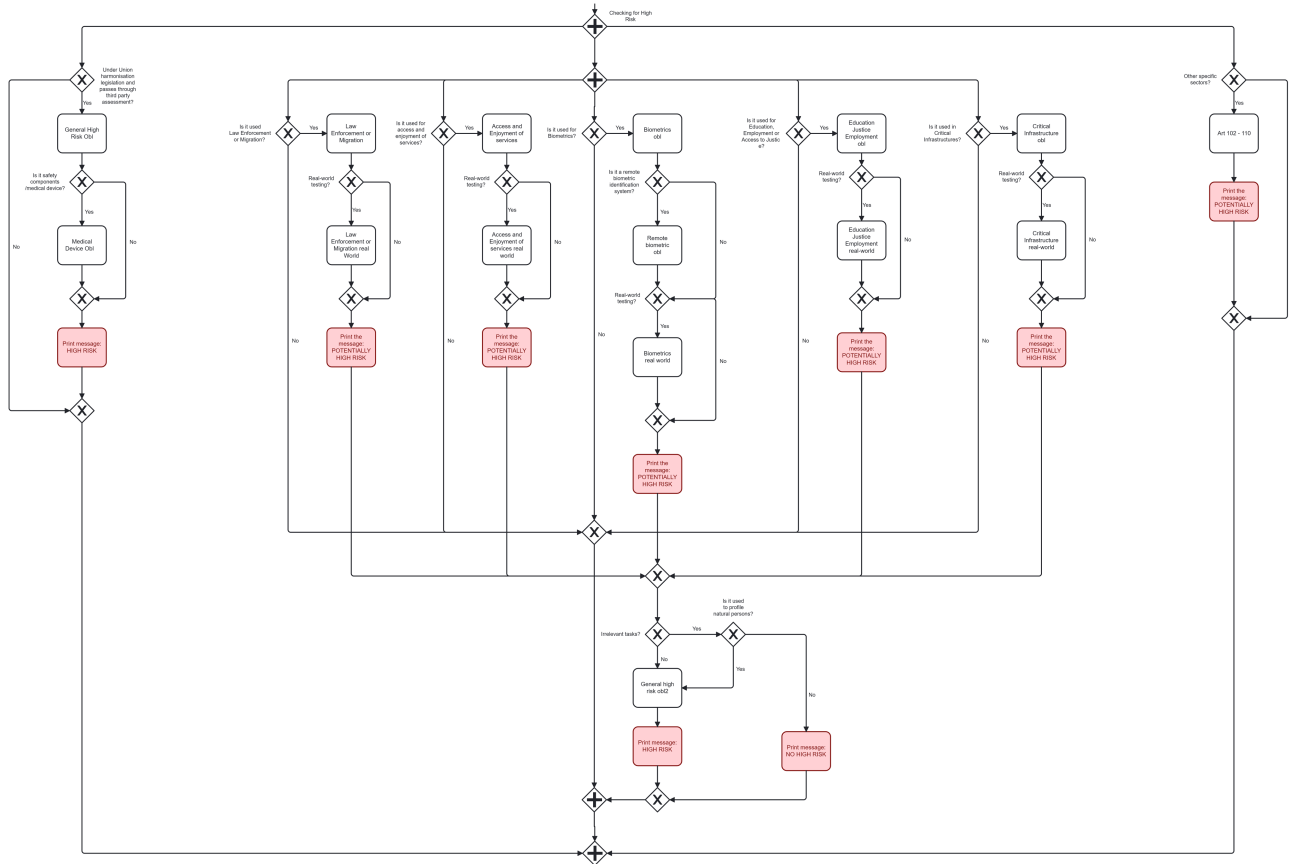
In the current situation, this distinction in terminology may not produce significant practical differences in the assessment of the obligations. However, it could become relevant in the event of future regulatory developments specifically addressing AI systems formally labeled as high risk. In such a scenario, it may be debated whether a system required to comply with high-risk

obligations should, by definition, be considered a high-risk system or whether the legislator’s decision not to include it within the categories listed in Annex I or Annex III reflects a deliberate intent to exclude it from the formal classification (and potentially from future regulatory measures) applicable to high-risk systems. Given that our BPM framework is structured to guide the identification and allocation of the relevant obligations, this particular third category has likewise been incorporated at this stage of the control process. In this work, we decided to create a separate branch for this issue to accommodate future modifications and evolving interpretations of the AI Act.

In the high risks section, if the AI system falls within one of the subjects listed in Annex I (“*Under Union harmonization legislation and passes through third-party assessment?*”), it must comply with the obligations labeled “*General High Risk Obl.*” set out in Table 1. In addition, where the AI system is “*safety components or a medical device,*” the obligations “*Medical Device Obl.*” are added to the obligations “*General High Risk Obl.*” (see Figure 10). Finally, the provider receives the message that the AI system considered is a high-risk AI system (“*Print message: HIGH RISK*”).

Moreover, while Annex I provides a list of specific subjects, Annex III provides some general fields of use and related exceptions. The fields are law enforcement, migration, access and enjoyment of services, biometrics, education, employment, access to justice, and critical infrastructures. If the AI system falls within one of these categories, the applicable obligations are “*Law Enforcement or Migration,*” “*Access and enjoyment of*

Figure 8
Details of the fourth sub-process. The high-risk branch



services,” “Biometrics obl,” “Remote biometric obl,” “Education Justice Employment obl,” and “Critical Infrastructure obl,” corresponding to the same labels in Table 1. If it falls within more than one category, the corresponding obligations are cumulative.

Finally, where the AI system falls within one of these categories and is deployed in real-world conditions, the additional obligation labeled “real world” is also to be applied (see Figure 11).

The corresponding labels with details are in Table 1. If the AI system under consideration falls within one or more of these categories, it is potentially classified as a high-risk system (“Print message: POTENTIALLY HIGH RISK”). The final two elements to be assessed, both of which may constitute exceptions and therefore prevent the system from being classified as high risk, are whether the system is used for an “Irrelevant task” and whether it is not “used for profiling natural persons” (see Figure 12).

If the system is used for an irrelevant task²⁰ and not to profile natural persons²¹, the system is not classified as high risk (“Print message: NO HIGH RISK”) and must comply with the obligations labeled “Residual obl” in Table 1. Otherwise, the system is confirmed as a high-risk AI system (“Print message: HIGH RISK”), and the obligations labeled “General high risk obl2” in Table 1 have to be added.

Finally, if the system falls into the amendments of Articles 102–110 (“Other specific sectors”), the provider must comply with the obligations labeled “Harmonizing obl” in Table 1 and receives the message “Print message: respect HIGH RISK OBL.”

At this point, if the system is classified as a high-risk system (so if “you receive the message: high risk”), we need to check whether there are exemptions or not:

- If the AI system “has already been certified under an EU cybersecurity schema,” the exemptions labeled as “Cybersecurity exemptions” in Table 1 may be applied.
- If the AI system is “developed using techniques involving the training of AI models,” the exemptions labeled as “Data governance only on testing dataset obl” in Table 1 may be applied.
- If the “Provider is an SME” or “a micro-enterprise,” the exemptions labeled, respectively, as “Simplified technical documentation” or “Simplified QMS” in Table 1 may be applied.
- If the AI system is “already placed on the market,” the exemptions labeled as “Different deadline compliance” in Table 1 may be applied.

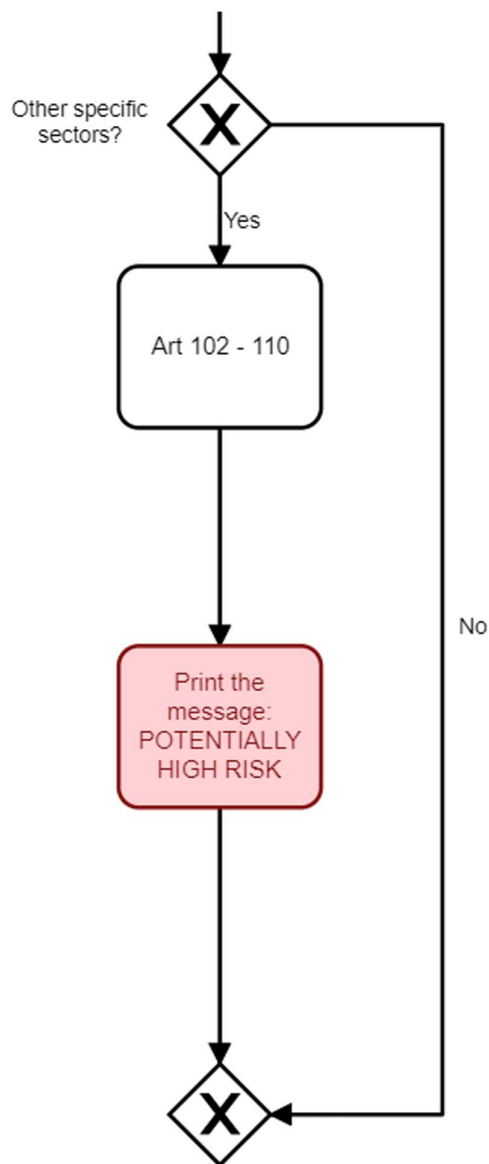
In order to highlight exemptions, their task is green (see Figure 13).

These can be regarded as harmonization measures. The exemptions serve the function of exempting providers from abiding by specific obligations; for example, if they already have a

²⁰Irrelevant task is an expression used in this paper to sum the examples defined in Article 6(3) as grounds to consider an AI system not high risk.

²¹Profiling natural persons can be considered the exception to the exception, as defined by the last line of Article 6(3).

Figure 9
Details of the fourth sub-process. The Article 102–110 branch



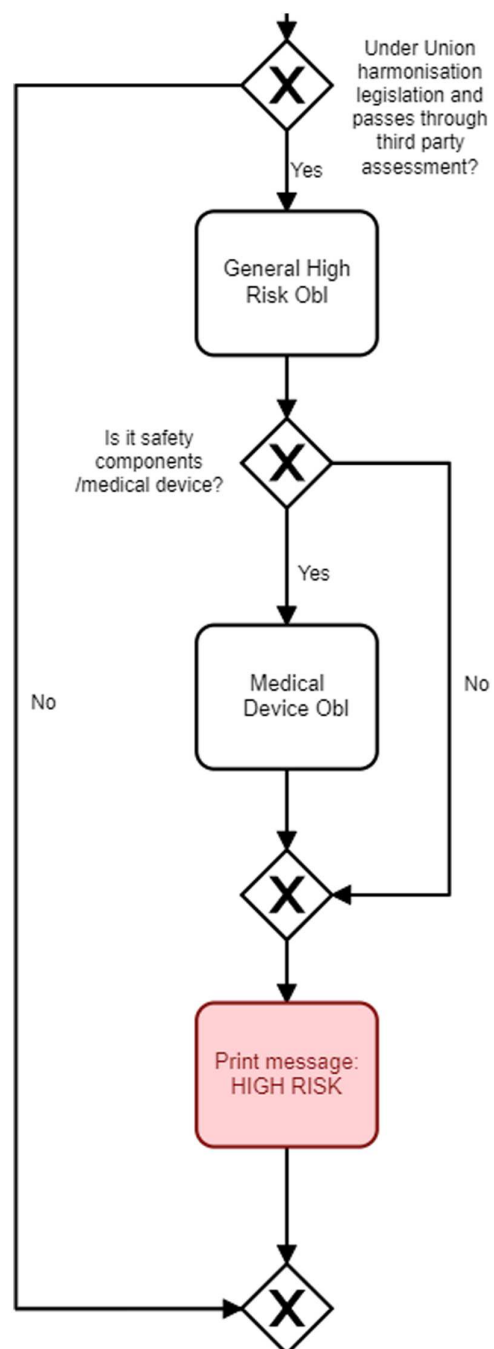
cybersecurity scheme issued pursuant to Regulation (EU) 2019/881, compliance with part of the data governance obligations²² is presumed.

The exemptions concerning cybersecurity, SME and micro-enterprises, and AI systems already placed on the market formally refer explicitly to systems classified as high risk. However, within our framework, these exemptions are also applied to “*Other specific sectors*” (of articles from 102 to 110) even if they are not explicitly classified as high-risk systems.

Finally, the assessment ends with the last sub-process: “*Checking the Interaction with People*,” shown in detail in Figure 14.

This final set of checks is intended to determine whether the AI system is classified as limited risk or falls within the residual category commonly referred to as minimal risk. If the AI system “*interacts with people*,” it has to comply with the obligations labeled

Figure 10
Details of the fourth branch. The Union harmonization legislation branch



“*Info obl*” in Table 1 and is classified as a limited-risk AI system (“*Print message: LIMITED RISK*”).

If there are no interactions with people, it is still necessary to verify whether the AI system “*generates synthetic content*.” If synthetic content generation is involved, the assessment proceeds by checking if the system is used for just editing, merely reproducing existing content, merely recording data, and just producing ephemeral data not to be stored or disseminated. The last

²²Article 10(4).

Figure 11
Details of the fourth sub-process. The Annex III branch

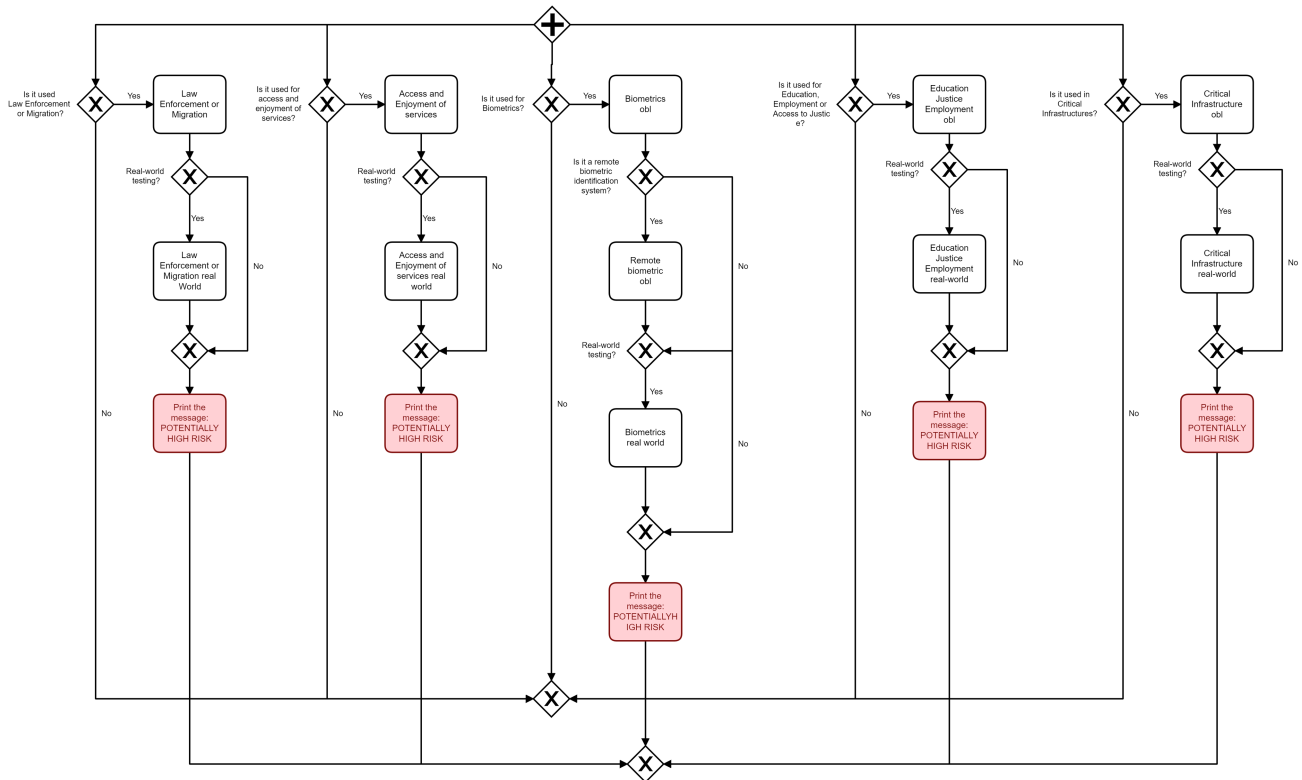


Figure 12
Details of the fourth process. Exceptions

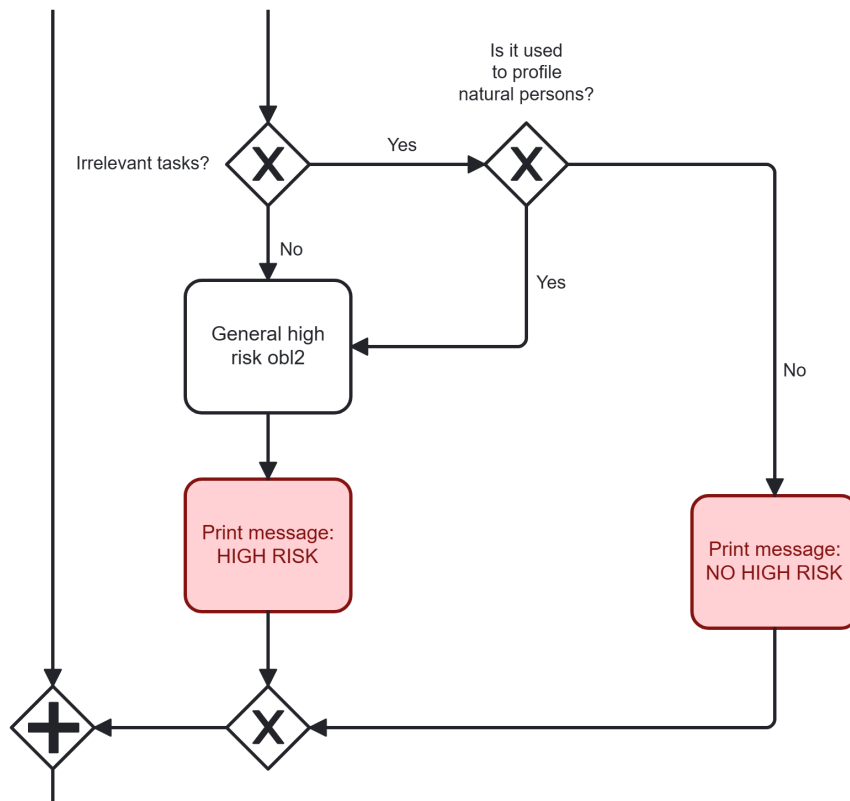
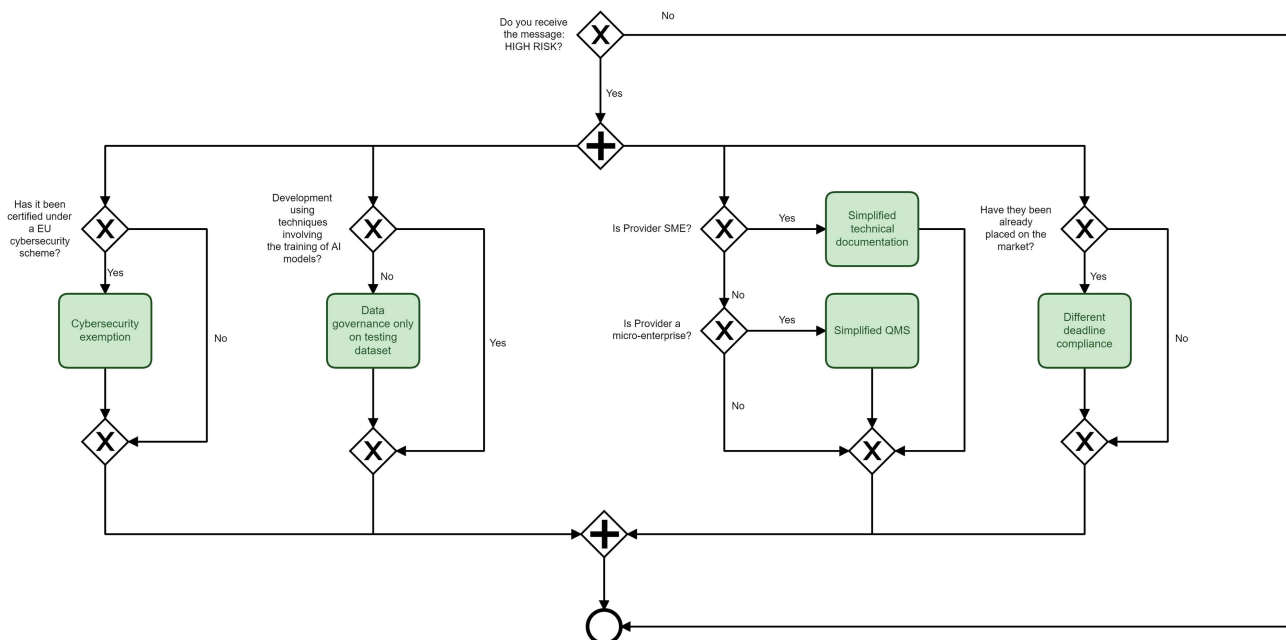


Figure 13
Details of the fourth process. The exemptions



three categories of exceptions belong to the Guidelines on the Transparency of AI systems²³.

If the system generates content and none of the exceptions applies, it qualifies as limited-risk AI (“*Print message: LIMITED RISK*”) and must comply with the obligations labeled “*Marking ob*” in Table 1.

If the system neither interacts with people nor generates synthetic content, or if it generates synthetic content solely for the abovementioned purposes, it falls within the residual category of minimal-risk AI systems (“*Print message: MINIMAL RISK*”). In this case, there are no mandatory obligations, only voluntary codes of conduct to which entities may choose to adhere. This obligation can be found in Table 1 under the label “*Voluntary codes of conduct*.”

4. Obligations

This section describes the scope of the obligations addressed in this methodology and explains the criteria for their selection.

Although the current version of the methodology is primarily grounded in the requirements of the AI Act, it has been intentionally designed to remain adaptable. In particular, it anticipates the integration of forthcoming technical standards. Such standards are essential in translating the Act’s high-level legal requirements into operational and technical practices. As several of these standards are still under development, they will be incorporated once published, thereby strengthening the methodology’s ability to provide precise compliance pathways and implementable technical solutions.

In Table 1, 25 Articles of the AI Act are cited. For the whole framework, a higher number has been taken into consideration. As an example, Article 2 has not been cited; however, it was the foundation for the initial process. There is a difference between the total number of Articles used to build the decision tree and those presented in this table. The goal of this table is a clear visualization

of obligations to follow, which is why the table only covers articles that are useful for such a purpose.

These provisions cover a broad range of topics under the AI Act and are specifically chosen to construct a coherent compliance pathway for AI systems. Provisions relating primarily to Commission governance structures and external supervisory mechanisms have been excluded, as they do not directly impose operational obligations on AI system providers or deployers.

In the 25 Articles, some refer to GPAI obligations (53–55), most of them refer to high-risk obligations (8–22, 49, 60, 61, 71, 72, 73), one refers to limited-risk obligations (50), and finally one to minimal-risk norms (95).

The high-risk classification bears the highest number of obligations. Besides those in Chapter III of the AI Act, some obligations either refer to specific conditions (e.g., testing in real-world conditions (60–61)), to exemptions due to harmonization (e.g., cybersecurity exemption (42)), or to exemptions due to the enterprise’s small dimensions (e.g., Simplified QMS (63)). Finally, as mentioned in Section 3, we chose to put the harmonizing obligations (102–110) as part of the high-risk context, as the amending obligations seem to suggest that an AI in those fields should be compliant with Chapter III of the AI Act²⁴.

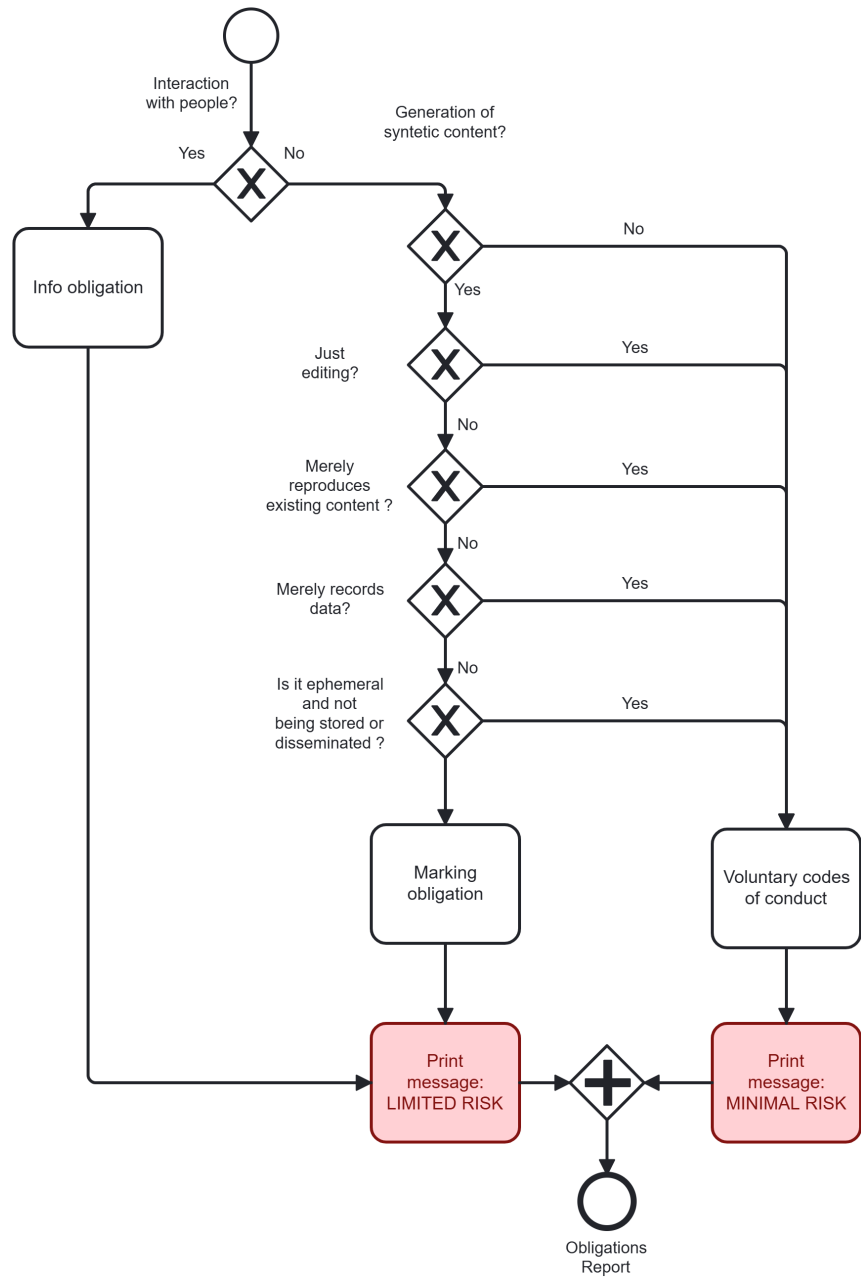
In addition, 8 Annexes have been taken into account to ensure detailed technical and procedural alignment with the obligations established in the selected Articles; for example, Article 43(3) gives information about the internal conformity assessment, and Annex VI gives details on it. Thus, the Annexes provide essential specifications, including documentation requirements. The Annexes chosen are Annexes IV, V, VI, VII, VIII, IX, XI, and XII.

Furthermore, Recital 104 has been taken into consideration, as it provides interpretative clarification on GPAIs. In particular, it elaborates on the exception applicable to GPAI models: where the required technical documentation has already been made available, the provider is subject to fewer compliance obligations than would

²³Draft Guidelines on the implementation of the transparency obligations for certain AI systems under Article 50 of Regulation (EU) 2024/1689 (the “AI Act”).

²⁴Chapter III of the AI Act is “HIGH-RISK AI SYSTEMS.”

Figure 14
Details of the fifth sub-process “Checking the Interaction with People”



ordinarily apply to GPAI models. For example, the provider is obliged to provide a summary of the data used in the training stage.

4.1. Simplification of obligations

In Table 1, there are three columns: “Label,” “Obligations,” and “Articles”:

- The first column “Label” refers to the denomination that was given to a group of obligations, the result of a certain pathway based on the nature and features of an AI system. This solution was chosen to avoid putting a long text as a result of a certain path.

- The second column “Obligations” refers to the obligation or obligations that are linked to an Article, Annex, or Recital. One Article may bear more than one obligation; for example, Art 15 includes “Accuracy,” “Robustness,” and “Cybersecurity.”
- The third column “Articles” refers to the Articles, Annexes, or Recitals that contain such an obligation.

The obligations presented in the table are a simplification of the norms prescribed in the AI Act. Some obligations listed retained their original titles (e.g., Technical Documentation—Article 11), whereas others were named differently from their article title, so as to clarify the action required.

One example of the latter is Article 111, named “AI systems already placed on the market or put into service and GPAI models already placed on the market.”

Table 1
Details on label, obligations, and articles

Label	Obligations	Articles
GPAI obligation	Transparency of Info for General Purpose Provider information and Technical Documentation for General Purpose Copyright Compliance Training Content Summary Updating Obligation	Art 53(1)(a–d), Art 53(3–7), Annex XI, Annex XII, <i>General Purpose AI, Code of Practice (soon)</i>
GPAI open and free	Summary Training Data Content Obligation to Comply with EU Copyright Law	Recital 104
GPAI obl with risk	Model Evaluation Risk Assessment for General Purpose Incident Reporting for General Purpose Cybersecurity Protection for General Purpose	Art 55(1)(a–d), Annex XI <i>General Purpose AI, Code of Practice (soon)</i>
General High Risk Obl	Compliance Obligation Risk Management System Data and Data Governance Technical Documentation Record Keeping Information to Deployer Human Oversight Accuracy Robustness Cybersecurity Quality Management System Documentation Keeping Corrective Actions Cooperation with Authorities Appointing Representative Post-market Monitoring EU Declaration of Conformity Application to Notified Body of Conformity Assessment	Art 8–22, Art 72, Annex IV, Annex V, Annex VII
Medical Device obl	Reporting of Serious Incidents (NOTE 2)	Art 73(10)
Law Enforcement or Migration	Internal Conformity Registration	Art 49(4), Annex VI
Law Enforcement or Migration real world	Real World Obligations Informed Consent for Real World test Internal Conformity Registration	Art 60, Art 61, Art 49(4), Annex IX, Annex VIII
Access and enjoyment of services obl	Internal Conformity Registration Post-market Monitoring (NOTE 2)	Art 49(1), Art 72(4), Annex VI, Annex VIII
Access and enjoyment of services real world	Internal Conformity Registration Post-market Monitoring (NOTE 2) Real World Obligations Informed Consent for Real World Test	Art 49(1), Art 60, Art 61, Art 72(4), Annex VIII, Annex VII

(Continued)

Table 1
(Continued)

Label	Obligations	Articles
Biometrics obl	Generic Registration	Art 49(3), Annex VIII
Remote biometric obl	Specific Registration	Art 49(4), Annex VII
Biometric real world	Real World Obligations Informed Consent for Real World Test EU Database	Art 49(4), Art 60, Art 61, Art 71(4), Annex IX
Education Justice Employment obl	Internal Conformity Registration	Art 43(3), Art 49, Annex VI
Education Justice Employment Real World obl	Real World Obligations Informed Consent for Real World Test Internal Conformity Registration	Art 49, Art 60, Art 61, Annex VIII, Annex IX
Critical Infrastructure obl	National Registration	Art 43(3), Art 49(5), Annex VI
Critical Infrastructure Real world	Real World Obligations Informed Consent for Real World Test National Registration	Art 49(5), Art 60, Art 61, Art 71, Annex IX, Annex VI
Harmonizing obl	Amending Obligations	Art 102–110
General high risk obl2	Compliance Obligation Risk Management System Data and Data Governance Technical Documentation Record Keeping Information to Deployer Human Oversight Accuracy Robustness Cybersecurity Quality Management System Documentation Keeping Corrective Actions Cooperation with Authorities Appointing Representative	Art 8–22
Residual obl	Registration	Art 49(2), Annex VIII
Cybersecurity exemption	Cybersecurity (NOTE)	Art 42 and Art 10(4)
Data governance only on testing data sets	Data and Data Governance (NOTE 3)	Art 10(6)
Simplified technical documentation	Technical Documentation (NOTE 2)	Art 11(1), Annex IV
Simplified QMS management from the Commission	Simplified for Micro-enterprise Compliance Obligation Risk Management System Data and Data Governance Technical Documentation Record Keeping Information to Deployer Human Oversight Accuracy	Art 63(1), Art 9–15, Art 72, Art 73

(Continued)

Table 1
(Continued)

Label	Obligations	Articles
	Robustness Cybersecurity Post-market Monitoring Reporting of Serious Incidents	
Different Compliance Deadline	Different Compliance Deadline	Art 111
Marking Obligation	Transparency – Obligation for Providers and Deployers	Art 50(2)
Voluntary codes of conduct	Voluntary Codes of Conduct	Art 95
Info Obligation	Transparency – Obligation for Providers and Deployers	Art 50(1)

It reads as follows:

“[...]AI systems which are components of the large-scale IT systems [...] that have been placed on the market or put into service before 2 August 2027 shall be brought into compliance with this Regulation by 31 December 2030.”

In the table, this obligation has been reported as “*Different Compliance Deadline*,” as the purpose of these titles is to make obligations clearer, explicit, and easier to spot.

The obligations presented in the table are also complemented by the details in the work “AI Compliance Framework: A Handbook For Compliance Tools based on AI Acts’ Requirements” [4]. In the handbook, each obligation is divided into three phases: Compliance (Project phase), Conformance (Development and Deployment), and Audit (Monitoring). This division is designed to underline in detail practical actions to be taken in various phases of the creation of an AI system.

For example, the Risk Management System prescribed by Article 9 has been divided in this way²⁵:

- Compliance: check for the Risk Management plan and the related documentation.
- Conformance: check for risk controls, including testing of mitigations and validation of risk thresholds.
- Audit: check that the Risk Management plan and measures are always in force, maintained, and updated.

This Compliance Framework is an additional layer to the decision tree framework. The next section presents two case studies referring to the classification of systems and the allocation of obligations. The conjunction with the Compliance Framework will be analyzed in another paper.

5. Two Italian Case Studies in Finance

In Italy, a company aiming at developing an AI system is encouraged to join initiatives by the Italian Sandbox. Generally, these initiatives are in collaboration with other Sandboxes all over Europe and stakeholders from public and private environments.

One example of this initiative is EUSAiR, the EU Regulatory Sandboxes for the AI project. The EUSAiR project is a two-year European Union initiative funded under the Digital Europe program. Its primary mission is to support and harmonize the implementation of AI Regulatory Sandboxes across all 27 EU Member States, as mandated by the EU AI Act. These sandboxes

are controlled, real-world testing environments where developers can safely train, test, and validate innovative AI systems under the oversight of national authorities before bringing them to market. The EUSAiR project tasks involve developing a uniform set of guidelines, tools, and benchmarks across the EU to streamline how regulatory sandboxes operate, empowering innovators by providing SMEs and start-ups with priority access to these sandboxes, significantly reducing their compliance costs and technical barriers to entry, enhancing authority oversight by equipping national competent authorities with operational standards, risk-assessment guidelines, and a “Train-for-Trainers” kit to improve their oversight of emerging AI technologies. Also, they analyze real-life experiences and collect best practices from sandbox participants to provide data-driven feedback directly to European policymakers. Their strategy involves pilot testing within HPC infrastructures, such as Europe’s supercomputing centers (e.g., LUMI and Cineca), to observe how AI solutions scale in practice. They also collaborate with key European infrastructure like EDIHs and Testing and Experimentation Facilities to ensure a unified ecosystem.

In this context, two models²⁶ will be examined:

- A synthetic data generator
- A rating model

In these two case studies, the EUSAiR legal team collaborated in reviewing and validating the legal analysis conducted within the BPMN framework, which is designed to support providers throughout the assessment of their AI systems. The case studies provide qualitative evidence based on the practical application of the framework and collaboration with legal experts.

The provider of these two models is an Italian company specialized in areas such as risk management, wealth and asset management, financial advisory, sustainability, and AI-driven analytics.

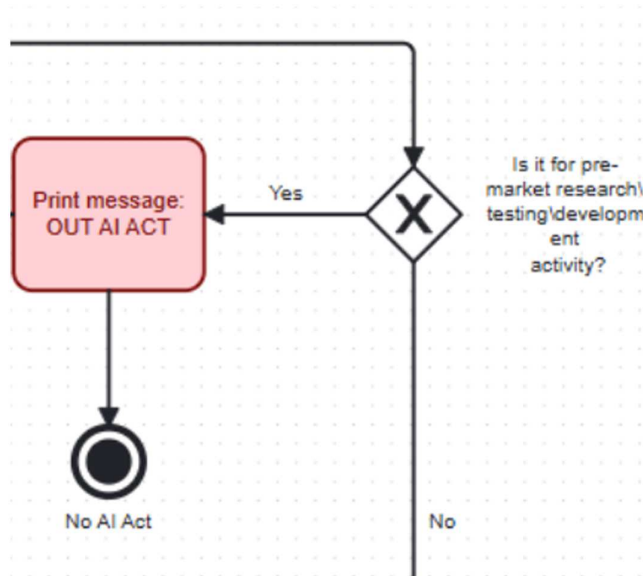
5.1. First model: The synthetic data generator

In this context, the first model generates information related to bank current accounts. It is described as a generative AI model

²⁵This is just a sum of what is described in the Handbook; the complete description can be found in section 3.14 of the Handbook.

²⁶In this paper, we refer to **model** as a complete learned function that maps inputs to outputs, typically trained on data to perform a specific task such as classification. In contrast, we refer to **module** as a smaller component that performs a specific sub-function and is often combined with other modules to form a complete model.

Figure 15
Detail view of previous Figure 3



generating data for scenario analysis, forecasting, etc. One of the most significant examples of its use is the testing of Anti-Money Laundering (AML) policies.

5.1.1. Scope

Starting from the scope, this model is not in any of the categories that might exclude it from the scope of the AI Act.

It is not implied for military, defense, or national security purposes, nor for law enforcement and judicial cooperation or research and development. In the pre-market research/testing or development activities category, one specification is due. It is used for testing activities in the banking field. However, its goal is to be sold to bank clients, not for an internal pre-market activity (reference in Figure 15). The provider (company) will sell it to clients (banks).

5.1.2. System

Moving the analysis to the system section, this model is described as a generative AI model to create synthetic banking data for data augmentation, scenario analysis, forecasting, and the creation of privacy-preserving data.

Due to these characteristics, the system can be deemed to be inside the AI Act scope, as it is not used for mathematical optimization, basic data processing, or based on classical heuristics, nor is it a simple prediction system (reference in Figure 3).

This model is machine-based. According to the Guidelines on the definition of AI systems: “AI systems must be computationally driven and based on machine operations,” this model runs on the computational means of the company to carry out the operations of generating content (reference in Figure 16). On the same note, the model can be deemed as autonomous, as it only needs a human entry to start the operations, and its parameters are thoroughly assessed through set metrics. It does not embed self-learning capabilities. However, this characteristic is not fully

relevant to assess whether an AI system falls within the AI Act scope²⁷.

5.1.3. Model

Going to the model section, this system is claimed to be a generative AI model to create synthetic banking data. Generative AI models and GPAI systems are not synonyms, though.

In the AI Act, these two terms are associated in Recital 99 and Recital 105. However, specific kinds of generative models are considered to be GPAIs. In the two recitals, the AI Act refers to “Large Generative AI models.”

According to the Act, the three relevant characteristics are related to the range of tasks, the employment of a large amount of data, and training compute.

In the synthetic data generator, the range of tasks is limited to the generation of content for the banking field. Therefore, this sub-process would conclude after the first gateway (reference in Figure 17).

5.1.4. Context

In the context section, the first half of the sub-process is dedicated to banned practices. The synthetic data generator is not used for criminal prediction or for manipulation of individuals or groups; it does not create images or infer emotions in the work or education spaces.

Data of customers used to generate synthetic content are anonymized when they are entry data; there is no reference to personal characteristics or behaviors. Therefore, they would not fall into the banned biometric categorization systems (reference in Figure 18).

The second half of the sub-process is dedicated to high-risk uses. The model is not included in the union harmonization legislation, nor is it under regulations listed in Art 102–110. Thus, the focus is on the uses of Annex III.

The most fitting category in Annex III is “Access to and enjoyment of essential private services and essential public services and benefits,” as in letter B, there is a reference to credit score and the determination of persons’ access to financial resources or essential services (reference in Figure 19).

The influence of synthetic data databases for scenario analysis and money laundering is debatable. For the sake of completeness, the AI Act mentions the use of AI systems to detect money laundering. Yet, these models are considered to be outside the high-risk category in case these investigations are conducted by law enforcement units for the purpose of prevention, detection, investigation, and prosecution of criminal offenses. This is not the case. Additionally, in Annex III, number 5, letter B excludes systems used for the purpose of detecting financial fraud.

Given that this model is not directly used for the purpose of detecting financial fraud itself, whereas it is used to create synthetic transactions and current accounts, two scenarios are feasible.

²⁷The Guidelines state “The use of the term ‘may’ in relation to this element of the definition indicates that a system may, but does not necessarily have to, possess adaptiveness or self-learning capabilities after deployment to constitute an AI system.” Therefore, it will probably be a ground of interpretation for AI systems with other characteristics and will be discussed at the EU Jurisprudence level. For these reasons, it is included in the assessment process.

Figure 16
Detail view of the second part of Figure 3

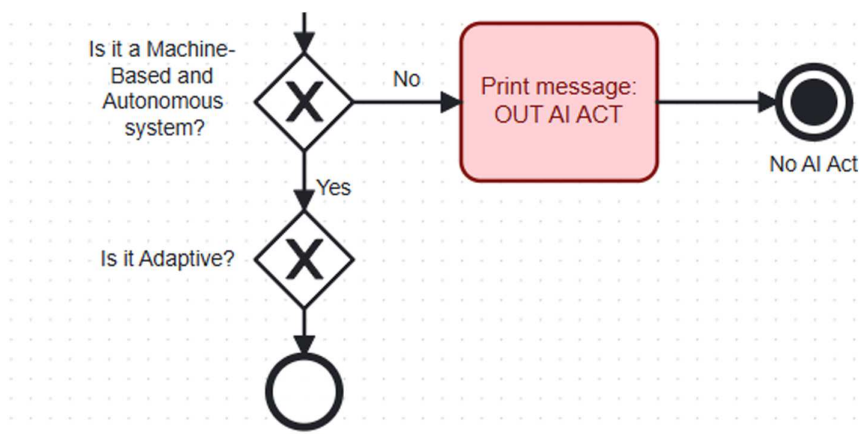


Figure 17
Detail view of previous Figure 5

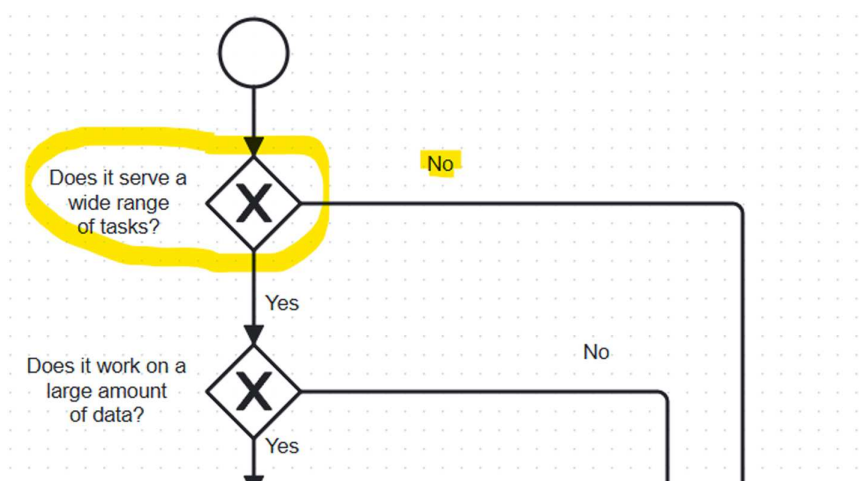
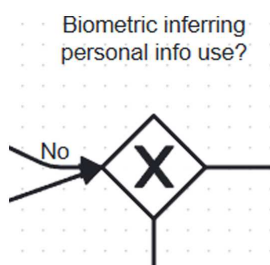


Figure 18
Detail view of previous Figure 6



the rationale of conducting analysis in the banking field. Testing AML policies represents only one illustrative application of the synthetic data generated by the model, which may be employed across a broad range of analytical and predictive tasks within the banking domain. Nonetheless, this scenario likely leads to exclusion from the high-risk class under Art 6(3), as it could be considered preprocessing and thus part of irrelevant tasks (see Figure 20). Also, as it creates synthetic info and accounts, it would not be considered as profiling natural persons.

Different in rationale, both scenarios lead to not categorizing this model as a high-risk system under the AI Act.

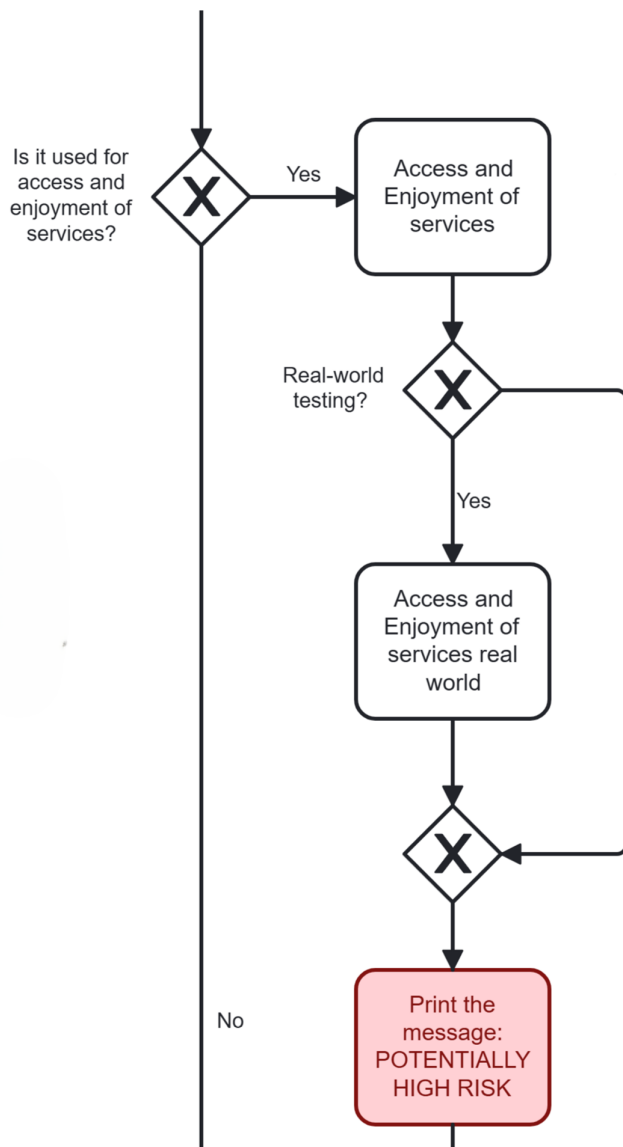
- **Scenario 1: linked to an Anti-Money Laundering model/module.** In such a case, it could fall under the exception of Annex III 6(B); thus, it would be excluded from the high-risk category.
- **Scenario 2: not linked to an Anti-Money Laundering model/module.** In such a case, it could be deemed as not falling under the exception but still falling in the category of credit scoring due to

5.1.5. Interaction with people

In the final step, the assessment checks whether the model needs to abide by transparency obligations.

Generation of synthetic content is a characteristic listed in Art 50(2). The rationale of this Article is to prevent users from being exposed to synthetic content without their knowledge. Still,

Figure 19
Detail view on access and enjoyment of Figure 6



this content is not for bank customers or users in general; it is material to be fed for analytic purposes.

On top of that, this content could be deemed as merely reproducing existing content (see Figure 21), as synthetic generation combines and interpolates patterns already present in the training data, and no new information is generated.

5.2. Second model: The rating model

The second model is a rating model: it extracts data from bank account information, categorizes it, and rates it. It operates through methods of machine learning and is used for means of credit scoring.

5.2.1. Scope

This model is not in any of the categories that might exclude it from the scope of the AI Act.

As in the previous case, it could be considered as “pre-market activities” (reference in Figure 15). However, the provider

(company) will sell it to clients (banks), and they will use it for pre-market activities. Thus, it is a product to be sold.

5.2.2. System

This model can be divided into two modules: the data categorization module and the account creation module.

- The first module has a rule-based framework and embeds a family of logistic regression models.
- The second module is a modeling framework that employs methods of machine learning.

These two modules compose the two-stage design pattern of the rating model. The first module has a rule-based framework, and according to the Guidelines on AI definition, rule-based models can be found to be outside the AI Act scope, especially if they fall under the category of systems based on classical heuristics (see Figure 22).

Also, the first module embeds methods of logistic regression, which is cited in the Guidelines as an exception to the AI scope, specifically in the category of systems for improving mathematical optimization (see Figure 23).

Conversely, the second module employs methods of machine learning, which are included in the definition of an AI system. Nonetheless, these two modules are part of a complex model; thus, they will be considered together. This remark was made to highlight that logistic regression is currently very popular in machine learning methods, despite it being cited in the Guidelines²⁸.

5.2.3. Model

According to the Act, the relevant characteristics are the range of tasks, the employment of a large amount of data, and training compute. In the rating model, the range of tasks is limited to the analysis of bank data. Therefore, this sub-process would conclude after the first gateway (reference in Figure 17).

5.2.4. Context

The model is not used for criminal prediction or manipulation of individuals or groups; it does not create images or infer emotions in the work or education spaces. In the data entries of customers, there is no reference to personal characteristics or behaviors. Therefore, it would not fall under the banned biometric categorization systems.

The model is not included in the union harmonization legislation, nor is it under regulations listed in Art 102–110.

It would fall under Annex III 5(B) for the purposes of credit scoring (reference in Figure 19). The assessment then proceeds by determining whether the potentially high-risk systems fall within the exceptions set out in Article 6(3).

Here, we display two scenarios:

- **Scenario 1: modules considered apart.** In case modules are judged individually, and logistic regression is inside the scope of the AI Act, it is feasible that the first module will fall outside the high-risk category for the narrow purposes of preprocessing data.
- **Scenario 2: modules. together. strong.** In case modules are judged as a whole complex system, the final purpose will remain credit scoring; thus, the model will fall into the high-risk category.

²⁸Guidelines on the definition of an artificial intelligence system (42).

Figure 20
Detail view on the 6(3) exception in Figure 6

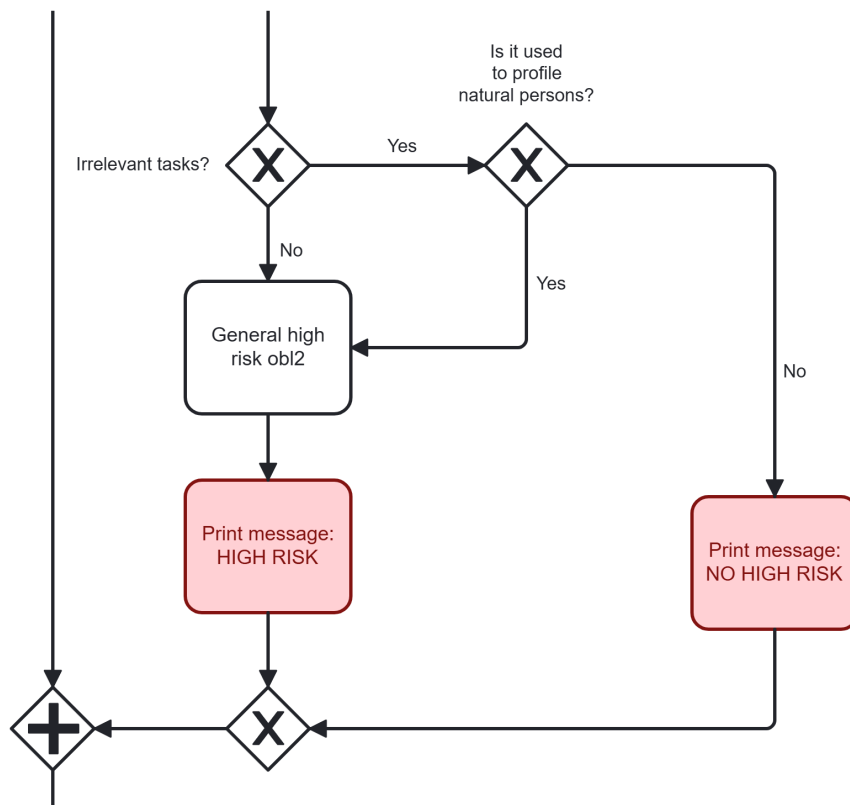


Figure 21
Detail view on exception of Figure 14

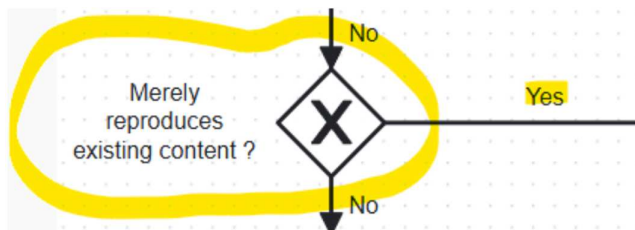


Figure 22
Detail view on exception of Figure 4

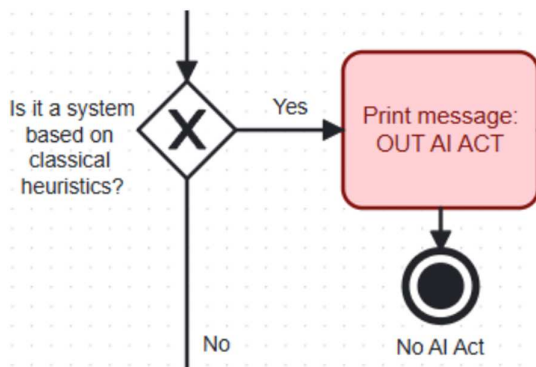
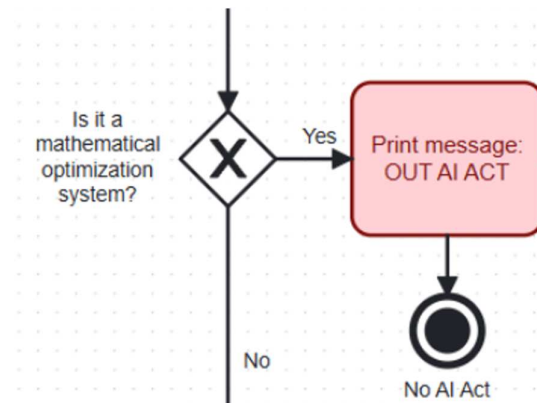


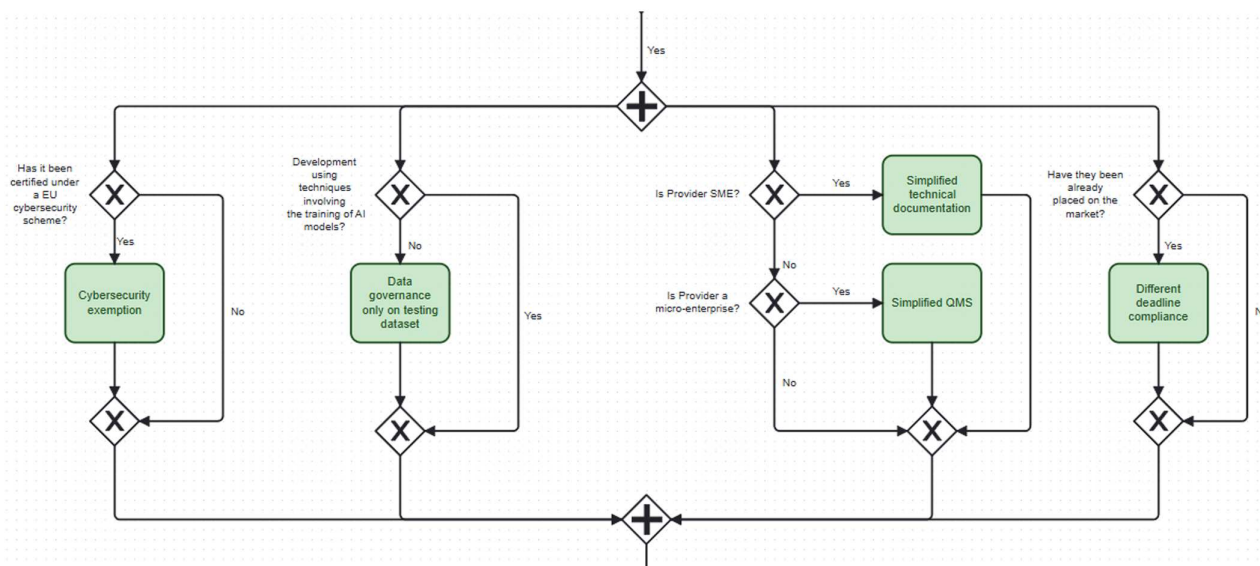
Figure 23
Detail view on exception of Figure 4



In Scenario 1, high-risk obligations on the provider are lifted. In Scenario 2, to properly define which obligations the provider might have to abide by, some information is to be added: the company has not already adopted a cybersecurity scheme, and the developer has used training models and AI techniques to build the model. Additionally, the company is not an SME or a micro-enterprise, and the model has not been placed on the market (see Figure 24).

Thus, the provider has full obligations for its high-risk AI model (see Table 1 “General high risk obl2”).

Figure 24
Detail view on exception of Figure 6



5.2.5. Interaction with people

This model does not interact with people nor generate synthetic content; thus, the provider should not abide by transparency obligations.

6. Limitations and Conclusion

This paper proposes a structured methodology designed to support AI providers in addressing the regulatory requirements introduced by the EU AI Act. The approach translates the provisions of the regulation into an operational framework that assists organizations in identifying the classification of their AI systems and the corresponding obligations established by the regulation. The methodology is based on a BPMN structure that enables legal professionals and organizations to follow a clear and traceable compliance pathway, systematically connecting legal provisions with concrete assessment steps.

Further guidance is given by the integration of the BPMN model with a complementary handbook, published separately, which provides interpretation and contextual explanations of the relevant provisions of the AI Act. Together, these two elements provide providers and controllers with a structured overview of the classification logic of the regulation, the obligations associated with different categories of AI systems, and the timeline for their implementation.

The authors acknowledge that this framework presents two main limitations:

- **Uncertainty handling:** the framework is based on the norms of the AI Act. Part of these norms has changed due to fast-evolving technology. For this reason, the AI Act and its interpretation through Guidelines and implementing documents are still uncertain and could lead to uncertainty in assessment (e.g., in Figure 3, the questions on the Autonomous and Adaptive system). However, the framework has a secondary aim of isolating ambiguous definitions from which, when there are further documents, further branches with further questions and options can be added. The authors have built this framework with this matter in mind, intending to isolate

problematic interpretative knots (such as the Art 102–110 in Figure 9) to untie in the future. The interpretative approach adopted by the author towards the AI Act may be characterized as a garantist approach to regulatory uncertainty: where the applicable legal qualification is uncertain, preference is given to the stricter interpretation, particularly where this results in the imposition of additional obligations, brings a system within the material scope of the AI Act, or leads to classification within a higher-risk category. This approach is grounded in the regulatory rationale underpinning the AI Act, which was initially conceived primarily around the need to ensure effective oversight, control, and monitoring of AI systems presenting heightened risks to individuals and society. While such an interpretative approach may, in certain circumstances, hinder the pursuit of a more expedited path toward compliance, it prioritizes a more precautionary approach to the governance of AI systems and models. Nevertheless, future interpretative developments will provide greater clarity and, where appropriate, permit a less restrictive approach to classification and compliance.

- **Non-automated system:** this framework is not automatic. Currently, it works as a guideline for providers to navigate the classification and requirements of the AI Act, but no element of automation is included. However, we see this as the first step of a bigger project. Our next step involves the formalization of this BPMN in a machine-readable format.

For such reasons, future work will focus on expanding specific branches of the BPMN model to incorporate additional regulatory guidance and sector-specific requirements, for example, Article 6 of the AI Act, which refers to the development of a comprehensive list of practical examples distinguishing high-risk and non-high-risk AI systems. These guidelines are at their initial stages of drafting; once they are published in their final form, they can be integrated into the existing BPMN sub-process dedicated to the “Context” assessment, which already addresses Article 6 and its interaction with relevant New Legislative Framework (NLF) product safety legislation. The sub-process is already present within the BPMN and can therefore be refined. In this way, the methodology can progressively integrate new

regulatory clarifications by refining specific decision nodes and sub-processes without requiring structural changes to the overall BPMN architecture.

It is important to highlight again that any tool claim to provide a formal “AI Act compliance certification” at this stage would be premature. Conformity assessment bodies are not yet fully operational across all relevant categories, and the regulatory framework still requires the adoption of implementing measures and technical standards. Consequently, services currently offered on the market that claim to provide binding certification should be understood primarily as commercial or advisory compliance-support tools rather than legally recognized conformity assessments.

The added value of the proposed approach lies in its methodological structure and institutional embedding. This methodology is being developed within the framework of the Italian AI Factory IT4LIA and interfaces with the Italian AI Regulatory Sandbox, where the ACN participates as a partner and is designated to act as the national supervisory authority for AI. Within this ecosystem, Sandbox reports may currently represent the closest available mechanism to a formal compliance validation. The proposed system should therefore be understood as a structured compliance-support tool that guides organizations through a transparent decision-making pathway and produces traceable and reviewable outputs, allowing critical points to be clearly identified and reassessed as the regulatory landscape continues to evolve.

Ethical Statement

This study does not contain any studies with human or animal subjects performed by any of the authors.

Conflicts of Interest

Ilaria Angela Amantea is the Peer Review Board Member for *Journal of Computational Law and Legal Technology* and was not involved in the editorial review or the decision to publish this article. The authors declare that they have no conflicts of interest to this work.

Data Availability Statement

The data that support the findings of this study are openly available in the repository at <https://github.com/Marinelele/Project-AI-Act-BPMN/blob/main/Context1.png>

Author Contribution Statement

Marinella Quaranta: Conceptualization, Methodology, Validation, Formal analysis, Investigation, Resources, Data curation, Writing – original draft, Writing – review & editing, Visualization, Supervision, Project administration. **Ilaria Angela Amantea:** Methodology, Validation, Formal analysis, Investigation, Writing – review & editing, Supervision, Project administration.

References

- [1] Ebers, M., Hoch, V. R., Rosenkranz, F., Ruschemeier, H., & Steinrötter, B. (2021). The European Commission's proposal for an Artificial Intelligence Act—A critical assessment by members of the Robotics and AI Law Society (RAILS). *J - Multidisciplinary Scientific Journal*, 4(4), 589–603. <https://doi.org/10.3390/j4040043>
- [2] Ebers, M. (2021). Standardizing AI — The case of the European Commission's proposal for an Artificial Intelligence Act. *SSRN Preprint*. <http://dx.doi.org/10.2139/ssrn.3900378>
- [3] Veale, M., & Borgesius, F. Z. (2021). Demystifying the draft EU Artificial Intelligence Act — Analysing the good, the bad, and the unclear elements of the proposed approach. *Computer Law Review International*, 22(4), 97–112. <https://doi.org/10.9785/crl-2021-220402>
- [4] Quaranta, M., Amantea, I., Governatori, G., Molinari, M., Billi, M., Galli, F., & Rotolo, A. (2025). AI compliance frameworks: A handbook for compliance tools based on AI Acts' requirements. *SSRN Preprint*.
- [5] McLachlan, S., & Webley, L. C. (2021). Visualisation of law and legal process: An opportunity missed. *Information Visualization*, 20(2-3), 192–204. <https://doi.org/10.1177/14738716211012608>
- [6] Gotsch, C., & Puchan, J. (2024). Harmonising innovation and governance: A lifecycle model for high-risk AI systems under the European AI Act. In *AKWI Jahrestagung 2024*, 79–95. <https://doi.org/10.18420/AKWI2024-006>
- [7] Costa, M. Z., Guizzardi, G., & Almeida, J. P. A. (2022). On capturing legal knowledge in ontology and process models combined. In *35th International Conference on Legal Knowledge and Information Systems*, 267–272.
- [8] Hanif, H., Constantino, J., Sekwenz, M. T., Van Eeten, M., Ubacht, J., Wagner, B., & Zhauniarovich, Y. (2024). Navigating the EU AI Act maze using a decision-tree approach. *ACM Journal on Responsible Computing*, 1(3), 1–16. <https://doi.org/10.1145/3677174>
- [9] Golpayegani, D., Pandit, H. J., & Lewis, D. (2022). AIRO: An ontology for representing AI risks based on the proposed EU AI Act and ISO risk management standards. In *Proceedings of the 18th International Conference on Semantic Systems*, 51–65. <https://doi.org/10.3233/SSW220008>
- [10] Hernandez, J., Golpayegani, D., & Lewis, D. (2024). Ontology-based approach for mapping concepts and requirements from regulations and standards: The case of the EU AI Act and international standards. In *Legal Knowledge and Information Systems*, 295–300. <https://doi.org/10.3233/FAIA241258>
- [11] Dumas, M., La Rosa, M., Mendling, J., & Reijers, H. A. (2018). Introduction to business process management. In M. Dumas, M. La. Rosa, J. Mendling, & H. A. Reijers (Eds.), *Fundamentals of business process management* (pp. 1–33). Springer. https://doi.org/10.1007/978-3-662-56509-4_1
- [12] Amantea, I. A., & Quaranta, M. (2024). Eco-sustainability and efficiency of healthcare complex systems. In *Proceedings of the 14th International Conference on Simulation and Modeling Methodologies, Technologies and Applications, SIMULTECH, 1*, 423–430. <https://doi.org/10.5220/0012857400003758>
- [13] Jin, Y. H., Tan, L. M., Khan, K. S., Deng, T., Huang, C., Han, F., . . . , & Wang, X. H. (2021). Determinants of successful guideline implementation: A national cross-sectional survey. *BMC Medical Informatics and Decision Making*, 21, 19. <https://doi.org/10.1186/s12911-020-01382-w>
- [14] Amantea, I. A., Quaranta, M., Molinari, M., Peduto, C., & Demarchi, F. (2024). E-dossiers for optimising data transmission in Italian civil area: Telematisation at the Tribunal of Cuneo. *Journal of Simulation Engineering*, 4, 10–1.
- [15] Allweyer, T. (2016). *BPMN 2.0: Introduction to the standard for business process modeling*. Germany: BoD—Books on Demand.
- [16] Agostinelli, S., De Luzi, F., Maggi, F. M., Marrella, A., & Volpi, A. (2025). Design patterns for GDPR-aware

- process modeling in BPMN. *Information Systems*, 137, 102646. <https://doi.org/10.1016/j.is.2025.102646>
- [17] Ciaghi, A., Weldemariam, K., Villafiorita, A., & Kessler, F. (2011). Law modeling with ontological support and BPMN: A case study. In *Proceedings of the Fifth International Conference on Digital Society*, 29–34.
- [18] Gatta, R., Vallati, M., Fernandez-Llatas, C., Martinez-Millana, A., Orini, S., Sacchi, L., ..., & Castellano, M. (2019). Clinical guidelines: A crossroad of many research areas. Challenges and opportunities in process mining for healthcare. In *International Conference on Business Process Management*, 545–556. https://doi.org/10.1007/978-3-030-37453-2_44
- [19] Van der Aalst, W. M. (2013). Business process management: A comprehensive survey. *International Scholarly Research Notices*, 2013(1), 507984. <https://doi.org/10.1155/2013/507984>
- [20] Governatori, G. (2015). The rigorous approach to process compliance. In *2015 IEEE 19th International Enterprise Distributed Object Computing Workshop*, 33–40. <https://doi.org/10.1109/EDOCW.2015.28>

How to Cite: Quaranta, M., & Amantea, I. A. (2026). AI Act in BPMN: More Than Just Words—A Perspective for Providers of AI Systems. *Journal of Computational Law and Legal Technology*. <https://doi.org/10.47852/bonviewJCLLT62029734>