

REGULATION OF APPLIED ARTIFICIAL INTELLIGENCE IN BIOMEDICAL ENGINEERING AS A HIGH-RISK ARTIFICIAL INTELLIGENCE SYSTEM IN THE EU AI ACT

Srdjan Djordjević¹  Nikola Ivković¹  Nikola Milosavljević¹ 

¹University of Kragujevac Faculty of Law, Kragujevac, Serbia

Artificial intelligence (AI) represents a global phenomenon changing all spheres of human life. Biomedical engineering is no exception, as many AI systems are applied to biomedical engineering inventions. The European Union has enacted the new EU AI Act, one of the world's first laws on AI. The main topic of this research is to examine what changes this new regulation brings to AI development in the field of biomedical engineering. An AI system applied in biomedical engineering is often considered a high-risk AI system, which means that AI developers are bound by a set of requirements and obligations to achieve a trustworthy, human-centric AI system. The authors analyze the impact and appropriateness of these requirements for developing AI systems in biomedical engineering using the legal dogmatic method, as well as by analyzing the secondary sources in the literature. The authors aim to present the current situation in AI regulation and make suggestions for further development.

Keywords: artificial intelligence, AI law, AI in biomedical engineering, EU AI Act, AI development

Submitted: January 23, 2026 **Revised:** May 25, 2026

Accepted: June 3, 2026

Published online: June 23, 2026

Copyright: © 2026, Author(s). This is an open access article published under the terms of the Creative Commons Attribution 4.0 International License. (<http://creativecommons.org/licenses/by/4.0/>).

Correspondence to:

Nikola Milosavljević
University of Kragujevac Faculty of Law
Kragujevac, Serbia
E-mail: nmilosavljevic@jura.kg.ac.rs

INTRODUCTION

In mid-2024, the European Union introduced a new Regulation 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 (hereinafter EU AI Act), representing one of the first comprehensive regulatory frameworks for AI globally. The primary objective of the EU AI Act is to protect citizens by ensuring that only human-centric and trustworthy AI is deployed within the EU (1,2). The Act primarily addresses high-risk AI systems, dedicating most of its provisions to their regulation. As we will see, AI systems incorporated into medical devices are classified as high-risk, imposing substantial obligations on developers and specific requirements for medical devices. It is therefore essential to examine these requirements and the associated burdens on EU-based AI developers, particularly given the competitive landscape with Chinese and American counterparts (3). Key questions include how these obligations affect the development of biomedical engineering and whether alternative regulatory approaches could better support AI innovation in this field.

To address these questions, this study examines the EU AI Act, related documents, and relevant literature. The analysis is structured as follows: first, potential applications of AI in biomedical engineering are surveyed to examine whether they would qualify as high-risk AI. Next, the principal requirements for high-risk AI under the EU AI Act are analyzed through the lens of ethical AI principles. Finally, the study provides an assessment and critique of these requirements, with particular attention to the necessity and significance of developing AI applications in biomedical engineering.

AI APPLIED IN BIOMEDICAL ENGINEERING AS A HIGH-RISK AI SYSTEM

AI in biomedical engineering

Biomedical engineering encompasses various subfields that advance healthcare and improve human well-being (4). To determine whether AI systems used in biomedical engineering qualify as high-risk, we first need to analyze the types of AI used. There are three main areas in which AI is applied.

The first area is experimentation and drug development. Experimentation, although based on extensive knowledge, still relies on a trial-and-error method. Consequently, many materials are wasted in the search for the right substance.

AI systems enable predictive modeling for drug testing, and recurrent neural networks are already used to predict chemical reactions in order to develop the desired drug or substance (5). For example, INS018_055, the first anti-fibrotic small-molecule inhibitor with promising anti-tumor relevance, was developed using the AlphaFold AI program from DeepMind, the drug discovery platform *Pharma.AI*, and the generative chemistry platform *Chemistry 42* (6). In this way, AI systems drastically reduce the cost and time of research and development. This means that the production price of the medications might be lower (7). Also, in the field of DNA research, predictions are useful for planning nucleic acid sequence assembly, cell culture, protein structure prediction, and related tasks (8).

The second area of application is diagnostics. Convolutional neural networks (7) are already applied in CT and MRI scanning, electroencephalography (EEG), and other fields of imaging diagnostics used for cancer, cardiovascular disorders, and neurological conditions (4). Great accomplishments have been achieved in the diagnosis of hepatocellular carcinoma, myelofibrosis, prostate cancer, and hypersensitivity pneumonitis (5). In many cases, the success rate of AI systems exceeded that of humans. This means these systems will be of great assistance once they are integrated into physicians' everyday work. The use of AI systems in diagnostics, alongside human control, will most likely minimize the risk of error. This is because machines have strong analytical capabilities that can identify key indicators, while human controllers have real-world experience that can prevent potential false diagnoses by the machine. It is fair to say that some of these systems are already in use, although not declared as AI (9). Another application is precision medicine. With AI-powered devices, parameters such as heart rate, blood pressure, and pulse can be continuously monitored. This has immense potential for early diagnostics and appropriate therapy initiation (4). Programs such as *xRapid-lab*, *xRapid-Malaria*, *PCR.AI*, and *PIVOT* enable faster, more accurate detection of diseases under a microscope (7).

Finally, the third area of AI application in biomedical engineering is prosthetics and healing. AI systems are integrated into implants and exoskeletons, which could drastically improve users' quality of life. The fact that AI has "learning" capabilities could prove useful, as it can learn user patterns to help guide movement (7).

Biomedical AI as a high-risk AI system

The EU AI Act (10) adopts a so-called risk-based approach, meaning that all AI systems are classified into four groups

(or, by some accounts, three groups) (11): forbidden AI practices, high-risk AI systems, low-risk AI systems, and/or minimum-risk AI systems. There is no official, legally binding definition of high-risk AI systems in the EU AI Act. However, recital 7 provides that, in order to ensure a consistent and high level of protection of public interests in the fields of health, safety, and fundamental rights, common rules for high-risk AI systems should be established. This sentence, although not an official definition, contains important factors for classifying a system as a high-risk AI system (namely, danger to health, safety, and fundamental rights).

On the other hand, Article 6 of the EU AI Act contains the actual classification criteria for high-risk AI systems, dividing them into two groups. The first group comprises AI systems applied in specific areas of human life, which make them so risk-prone that they cannot be excluded as high-risk AI under any circumstances. AI systems are considered to fall in this group if two conditions are cumulatively fulfilled: a) the AI system is intended to be used as a safety component of a product, or the AI system is itself a product, covered by the Union harmonization legislation listed in Annex I; and b) the product whose safety component pursuant to point (a) is the AI system, or the AI system itself as a product, is required to undergo a third-party conformity assessment prior to being placed on the market or put into service pursuant to the Union harmonization legislation listed in Annex I.

Annex I provides a compilation of EU legislative acts regulating various products and services subject to special conditions for internal market placement, including toys, machinery, etc. The other group consists of AI systems that are enumerated in Annex III of the regulation. In certain cases, some of the AI systems from this group can be treated as a low-risk AI system if they do 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, and if conditions from Article 6, paragraphs 4-8 are met.

The list contained in Annex I of the EU AI act enlists, among others, Regulation (EU) 2017/745 of the European Parliament and of the Council of 5 April 2017 on medical devices, amending Directive 2001/83/EC, Regulation (EC) No 178/2002 and Regulation (EC) No 1223/2009 and repealing Council Directives 90/385/EEC and 93/42/EEC (hereinafter MDR) and Regulation (EU) 2017/746 of the European Parliament and of the Council of 5 April 2017 on *in vitro* diagnostic medical devices and repealing Directive 98/79/EC and Commission Decision 2010/227/EU (hereinafter IVDR). The MDR and IVDR are the regulations

governing the field of medical devices. In the MDR (to which the IVDR also refers for this definition), it is explicitly defined that a “medical device means any instrument, apparatus, appliance, software, implant, reagent, material or other article intended by the manufacturer to be used, alone or in combination, for human beings for one or more of the following specific medical purposes...investigation, replacement or modification of the anatomy or of a physiological or pathological process or state...” This simplifies the classification, since it is clear that AI systems applied in biomedical devices are high-risk if they are intended to serve as a safety component of a product, or if the AI system itself is a product of biomedical engineering. Since these medical devices undergo third-party conformity assessment, AI systems used in biomedical engineering must meet both requirements to be classified as high-risk AI. Since they fall into the first group, there is no possibility for these systems to be exempted from the high-risk AI systems categorization (11–13). However, it is possible that some systems not regarded as high-risk because they do not fall within the definition of a medical device may still pose significant risks in practice.

The question that remains open is whether every system applied to a biomedical engineering product can be categorized as AI. The EU AI Act defines an AI system as a machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments (10). The main challenge will be to determine the level of autonomy required for such systems to be considered AI (14). However, most systems applied in medical devices are likely to qualify as AI systems, and hence as high-risk AI systems under the EU AI Act. This means there is a set of requirements these systems must comply with, which we discuss in the following section.

PRINCIPLES AND REQUIREMENTS FOR HIGH-RISK AI SYSTEMS UNDER THE EU AI ACT

Requirements for high-risk AI systems are set out in Recital 1 in order to achieve human-centric, trustworthy AI, which is the ultimate goal of the EU AI Act. These requirements were derived from seven ethical principles developed by the High-Level Expert Group on Artificial Intelligence in its Ethics Guidelines for Trustworthy AI (HLEG Ethics Guidelines), which were subsequently incorporated into Recital 27 of the EU AI Act. Since they form the core of the high-risk AI

regulation, as well as the EU AI Act as a whole, we will further analyze the requirements for high-risk AI systems through these principles:

Human agency and oversight

This principle means that AI systems must be developed and used as tools that serve people, respect human dignity and personal autonomy, and function in a way that can be appropriately controlled and overseen by humans (10). It is a continuation of the “human-in-the-loop” principle already applied in the GDPR (15). The purpose of this principle is to promote safe and trustworthy AI, as it is believed that only AI controlled by humans can be considered human-centric and trustworthy (16). It aims to guarantee that AI will never gain supremacy over the human race, i.e., that “black-box” systems cannot replace human decision-making, which would not be in accordance with European values and fundamental rights (15). Due to its importance, it is placed before other principles. This principle also means that humans must be involved in AI decision-making. Given the present distrust in AI systems, the final decision must always be made by humans, which is also important for the accountability principle (16). The principles operationalized in Article 14 of the EU AI Act require human oversight. This requirement aims to prevent or minimize risks to health, safety, or fundamental rights, which is the general rationale for all high-risk AI system requirements. Repeating this aim in the human oversight requirement underscores its importance and the fact that it is the main protection against AI misuse. This requirement is to be achieved through two types of measures: those built into AI systems by the provider before they are placed on the market, and those implemented by the deployer (10). Built-in measures could, for example, be technical measures to facilitate the interpretation of outputs, while measures implemented by deployers could include guidelines and processes deployed within their organizations (17). It is required that a natural person assigned to the oversight of AI is able to: understand the capacities and limitations of the system, remain aware of tendencies toward relying or over-relying on the output produced, correctly interpret the output, decide not to use the high-risk AI system or to disregard, override or reverse its output, and intervene in the operation of the high-risk AI system or interrupt the system through a ‘stop’ button (10). According to these sub-requirements, it is essentially requested that the human controller has significant control over the AI system, which was deemed paramount in theory (18). This is achieved only when a controller understands the AI system's functioning,

which is particularly important given the “black box” problem in AI, i.e., the fact that humans lack insight into the AI “thinking” process and do not know how AI creates its output (15,18,19). It is also important that a human can override a machine’s decision and have a “killer switch” – a stop button that can at any moment abort the AI system process in case of emergency. In this way, it is believed that humans will retain sufficient control to prevent excessive AI system autonomy. However, the level of human responsibility and training is only provided in principle, and these aspects should be regulated in more detail (18). Interpreted literally, this principle would establish human dominance over the AI. Humans clearly hold a higher position. However, humans are not inherently superior to machines, since, as is well known, we also make mistakes. Therefore, we believe that humans should not be merely the controllers of the machine; human oversight cannot be a panacea for all the risks AI poses. The relationship between humans and machines should be one of mutual cooperation, in which each party performs the task it is best at and is controlled by the other to achieve a minimum level of error. Acknowledging that neither humans nor machines are perfect, in our opinion, is the path that could lead us to further scientific and technical development. On the other hand, rigid demands for human dominance over machines could have a chilling, or even negative, effect on AI development (15,16). Hence, these rules should be interpreted so that human oversight serves as a fallback mechanism in emergencies (20).

Technical robustness and safety

This principle means that AI systems must be developed and used in a way that allows robustness in the case of problems, provides resilience against attempts to alter the use or performance of the AI system so as to allow unlawful use by third parties, and minimizes unintended harm (10). The quality of an AI system must be at the highest level to minimize the potential dangers of misuse and backfiring. Through the requirements of accuracy, robustness, and cybersecurity, it is provided that high-risk AI systems shall be designed and developed to achieve an appropriate level of accuracy, robustness, and cybersecurity, and to perform consistently in those respects throughout their lifecycle. In theory, accuracy, robustness, and cybersecurity are three distinct requirements, different in nature, which makes interpretation much harder. It is also somewhat difficult to measure the “appropriate” level for all three requirements, given their distinct natures. Accuracy is not even defined in the EU AI Act, but it should not be regarded identically to

the accuracy principle in the GDPR; this does not mean that high-risk AI systems must never give false outputs, but rather that the system provides an appropriate level of correctness given the circumstances (20). Regarding robustness, it is provided that high-risk AI systems shall be as resilient as possible to errors, faults, or inconsistencies that may occur within the system or its environment. Some technical and organizational measures are enumerated, such as backup or fail-safe plans. Furthermore, systems that continue to learn must be developed in a way that eliminates or reduces the risk of biased outputs influencing input for future operations (feedback loops), which should be duly addressed (10). Robustness should be understood as resilience to every type of adversity, excluding cyberattacks (17). Finally, high-risk AI systems shall be resilient against attempts by unauthorized third parties to alter their use, outputs, or performance by exploiting system vulnerabilities, which guarantees cybersecurity (13). This requirement creates an obligation of consistent performance (20), which is also operationalized, among other things, by providing a risk management system. This system operates as a continuous iterative process planned and run throughout the entire lifecycle of a high-risk AI system, requiring regular systematic review and updating in the prescribed four steps (10). A risk management system is also designed to embrace the transparency principle, as risk management documentation allows us to partially observe the AI's work process (21).

Privacy and data governance

This principle means that AI systems must be developed and used in accordance with privacy and data protection rules while processing data that meets high standards in terms of quality and integrity (10). To protect human rights, it is necessary that personal data be protected in accordance with the GDPR. Furthermore, the faultless development and operation of an AI system depend heavily on providing it with adequate data. Hence, this requirement is not limited to personal data protection (20) but also prescribes that all other data must come from reliable sources and be processed in a way that does not discriminate against anyone. This is necessary because early applications of AI systems showed they were prone to bias and extremely unreliable when trained on inadequate datasets. It is provided that training, validation, and testing data sets shall be subject to data governance and management practices appropriate for the intended purpose of the high-risk AI system, and that these testing datasets shall be relevant, sufficiently representative, and, to the maximum extent

possible, free of errors and complete in view of the intended purpose (10). This basically means that the data input into a high-risk AI system must also be of high quality (20), in accordance with its intended purpose, and should be assessed for reasonableness and proportionality (17). Finally, regarding personal data protection, it is explicitly stated that special categories of personal data may be processed only in accordance with the GDPR and specific conditions (10). Unlike other data, personal data must meet a higher level of accuracy and quality under the GDPR. In one survey conducted among AI developers, this requirement was labeled as the hardest to implement. In their opinion, the most difficult task is to keep track of the data AI uses as input and to monitor for eventual biases and exposed personal data (22).

Transparency

This principle means that AI systems must be developed and used in a way that allows appropriate traceability and explainability while making humans aware that they are communicating or interacting with an AI system. It also requires duly informing deployers of the capabilities and limitations of that AI system and affected persons about their rights (10). Transparency is considered the greatest challenge with AI systems; therefore, most regulatory requirements are derived from this principle (1). An AI system's operation must be sufficiently transparent to enable deployers to interpret its output and use it appropriately. Furthermore, systems must include instructions for use containing concise, complete, correct, and clear information that is relevant, accessible, and comprehensible to deployers. The opacity and complexity of AI systems (i.e., the "black box" problem) are regarded as the highest risks concerning AI applications. Consequently, it is essential that the information necessary to understand the system is provided prior to placing it on the market (30). The minimum information required for these instructions is set out in Article 13, Paragraph 3 of the EU AI Act. These requirements for the instructions for use are essential given that these systems are highly sophisticated devices. The completeness of the instructions must also be measured to assess compliance with the EU AI Act (17). Yet, the most critical operational requirement is record-keeping. High-risk AI systems must technically support the automatic recording of events (logs) throughout their lifetime. This is of paramount importance due to the "black box" problem. To ensure a level of traceability of the functioning of a high-risk AI system that is appropriate to the intended purpose of the system, logging capabilities must enable the recording of

risk-relevant events (10). This requirement, as well as its importance, was recognized even at the level of HLEG Ethics Guidelines (20). It is, however, limited by the minimization principle of personal data protection (17). Finally, this principle is traced in the notification obligation, which dictates that the provider of an AI system is obliged to inform the market and the AI authority of any significant risks.

Diversity, non-discrimination, and fairness

This principle means that AI systems must be developed and used in a way that includes diverse actors and promotes equal access, gender equality, and cultural diversity, while avoiding discriminatory impacts and unfair biases that are prohibited by Union or national law (10). Highlighting this principle was necessary due to the recurrent biases shown in AI systems' operations. Sometimes biases are "inherited" from input data, whereas at other times, there is no clear explanation for them. This principle is operationalized through human oversight and technical robustness requirements. On the one hand, it is necessary that a human can check and remove unfair biases if and when they occur, and on the other hand, it is necessary that the system be built in a manner where these biases and discriminations do not occur, or only occur in a minimal number of cases (18). This principle also means that the use of AI systems should be user-centric and available to the widest possible range of users, without prejudice based on race, ethnicity, language, religion, skin color, sexual orientation, disability, social status, or other factors (23).

Social and environmental well-being

This principle means that AI systems must be developed and used in a sustainable and environmentally friendly manner as well as in a way that benefits all human beings while monitoring and assessing the long-term impacts on the individual, society, and democracy (10). AI systems have already changed the way we live. However, their future effects could be even more drastic and dangerous for society. First, it is necessary that AI systems do not harm the environment. Using AI systems is technically demanding and can require significant energy consumption and the creation of waste. Therefore, the technical robustness requirement and other obligations require that AI systems be environmentally friendly. Second, these systems can enhance or deteriorate social skills. Systems must be built in such a way as to prevent destruction or dereliction of social life and generally to protect the physical and mental

well-being of their users (23). Therefore, this principle is also operationalized through the technical robustness requirement.

Accountability

Accountability is not explicitly mentioned in Recital 27 as one of the principles, yet it was enumerated in the HLEG Ethics Guidelines. Moreover, several norms make it clear that this principle is implemented throughout the regulation. It means that, under the EU AI Act, developers, deployers, and other entities must keep records and documents to demonstrate compliance. With this in mind, alongside the EU AI Act, two international standards for risk management and compliance have been published: ISO/IEC 42001 "Information technology-Artificial intelligence-Management system" and ISO/IEC 23894 "Artificial intelligence-Guidance on risk management". According to these standards, as well as the EU AI Act itself, compliance proceedings have to be conducted (21). The compliance requirement provides that high-risk AI systems shall comply with the requirements, taking into account their intended purpose and the generally accepted state of the art in AI and AI-related technologies. The risk management system referred to in Article 9 shall be taken into account when ensuring compliance with those requirements. Also, the technical documentation requirement operationalizes this principle. The technical documentation shall be drawn up in such a way as to demonstrate that the high-risk AI system complies with the requirements set out in the EU AI Act and to provide national competent authorities and notified bodies with the necessary information in a clear and comprehensive form to assess the compliance of the AI system with those requirements (10). However, technical documentation should not be developed in isolation but rather in an iterative process as part of the risk management system, so these two requirements are very closely linked (20). Technical documentation will be comprehensive, especially in the field of medicine, as the MDR and IVDR already require extensive documentation (12). Finally, the accountability principle is also evident in the quality management obligation and document-keeping obligation. Notably, the quality management system obligation will be somewhat easier for developers in the medical field, as they already have a quality management system mandated by the MDR and IVDR, and the EU AI Act quality management system can be integrated into the existing one (12).

ASSESSMENT OF THE REGULATORY FRAMEWORK OF THE EU AI ACT

The EU AI Act's regulatory framework is both praised and criticized. On the one hand, it should be pointed out that the EU AI Act is one of the world's first comprehensive, legally binding pieces of legislation on AI (12). It is also important that the adopted regulatory framework is grounded in moral principles to place AI applications under the rule of law and morality. In this way, there is a clear attempt to prevent the use of AI that would be harmful or even dangerous to humanity (24).

However, the EU AI Act is often criticized for its approach to regulating AI development. First of all, the framework applied in the EU AI Act is essentially modeled after the GDPR (12,25). It is questionable whether data protection methods are appropriate for AI regulation. The fact is that AI is a far more complex phenomenon with a greater impact than simple data processing (14). Autonomous risk management and compliance procedures are necessary for personal data protection, as they are the only way for data processing by companies and governments to become transparent. Although transparency is important in AI, the issue here is very different. The challenge with AI transparency lies in the "black-box" problem: we do not know how AI reaches a conclusion or makes a decision, because the rules of human logic do not always apply to it. Hence, the activity log requirement is very important, as it helps illuminate how AI is "thinking," which, in turn, makes allocation of responsibility easier. However, it is questionable how much of this process can actually be revealed and how much of it can be meaningfully interpreted by humans. On the other hand, it is not clear why compliance proceedings are important for transparency in the AI industry. A quality management system, alongside extensive documentation, does not guarantee that the ultimate product will be flawless, nor will it solve the "black-box" problem. Furthermore, these requirements are quite expensive. According to some studies, they could cost between 193,000 and 330,000 euros (11). Although achieving this requirement is not considered the most challenging, a quality management system is the most expensive component for developers (22). Theoretical literature indicates that the documentation required for compliance often fails to meet the technical detail requirements of authorities, and that standardization is needed because information provided by different developers varies widely (1). Further, the subjects of these obligations are not only large corporations and government bodies but also startups and small and medium-sized

enterprises (SMEs), which often lack the manpower or funds to engage in complex compliance proceedings. Special regulatory measures are indeed in place to support startups and SMEs. However, these measures do not exempt them from compliance proceedings; they merely make the process somewhat less expensive (13). To protect the interests of AI system developers and remain in the race for AI development, the so-called "regulatory sandbox" method was introduced in the EU AI Act. Sandboxes are usually found in the financial sector. When new products are to be introduced, providers are allowed for some time to test them on the market in a controlled environment without being subjected to full regulatory enforcement. This way, a regulatory body can also learn about the product, and the developer can test it in real life. Therefore, it is important to implement them properly (26). Experts praised sandboxes as a good solution for AI development (2). However, sandboxes under the EU AI Act are provided only as an obligation for Member States, since the Union does not have authority over all fields of AI application in which sandboxes could be applied. This is the first flaw of the system, since there are no detailed rules for sandboxes, and control by the Union can only be exercised through reports. Second, sandboxes primarily address compliance obligations that can be burdensome during the development phase. Furthermore, a developer is still obliged to submit an exit report for activities performed in a sandbox. Third, sandboxes do not exclude the developer's liability, and competent authorities in member states could still stop work in the sandbox at any time. It is clear that the advantages provided by sandboxes are very limited and offer only uncertain guarantees to developers, who, in return, must reveal their trade secrets. Consequently, most developers will likely not participate in sandboxes (26). In the field of biomedical engineering, it is even questionable whether it would be rational to approve sandboxes, as health and safety could be at risk. The exception would be sandboxes for data governance obligations under the EU AI Act. This is the main reason the EU AI Act is often discussed in theory as being more accommodating to large AI companies (24,26). Conformity assessments are common for physical products and therefore include sampling, testing, inspection, and evaluation, which are challenging for intangible software (25). There are also problems with the risk-based approach. The risk-based approach, as well as the risk management system, rests on making compromises, since some levels of risk can be accepted and others cannot. However, for fundamental human rights, especially in health and safety, it is hard to argue that there

are acceptable risks (14). Finally, there is no clear criterion on why some systems are considered high-risk, while others are not (27). It is pointed out that classification is not based on empirical evidence and that research that was conducted prior to the EU AI Act adoption may be questionable, suggesting the classification was possibly conducted on a political basis (14,28). Bearing all this in mind, it is safe to say that the enumeration method for defining high-risk AI systems is quite inflexible and not future-proof, although that was one of the primary goals of adopting the EU AI Act (11,29).

In scientific areas such as biomedical engineering, the EU AI Act could introduce further complications. First of all, systems covered by Annex I cannot be arbitrarily added or removed, which means that AI systems classified as medical devices will automatically be considered high-risk, regardless of their specific clinical impact. This could discourage developers in this field; even when they develop a system that does not pose a high risk to health, safety, and fundamental rights, they must still comply with all the aforementioned obligations. Second, providers of biomedical services and equipment already have strict compliance obligations under the MDR and IVDR, which are now expanded by the EU AI Act. It is indeed provided that these conformity assessments should be merged and that requirements under multiple acts should be applied only once (10). However, it is argued in theory—because of this horizontal approach to regulation—that there are many inconsistencies across different regulations. This is especially true for the MDR. There are many different definitions of the same term in the EU AI Act and the MDR. Many requirements overlap, and there are obvious differences in their risk-based approaches: the MDR applies a risk-benefit balancing framework, whereas the EU AI Act does not. In theory, it is also noted that there are inconsistencies between compliance proceedings under the MDR and the EU AI Act, which would make merging the fact sheets more difficult (11–13). Fortunately, Article 43, paragraph 3, of the EU AI Act prescribes that compliance proceedings under the Act should be incorporated into existing proceedings, which should alleviate some pressure from developers. Additionally, the joint guidance document titled *"Interplay between the Medical Devices Regulation (MDR) & In vitro Diagnostic Medical Devices Regulation (IVDR) and the Artificial Intelligence Act (AIA)"*, adopted by the European Artificial Intelligence Board (AIB) and the Medical Device Coordination Group (MDCG), provides assistance to medical device developers for implementing EU AI Act rules into existing MDR and IVDR frameworks,

resolving some of the issues pointed out in theory (30). There is also hope that the Commission's proposal for the European Biotech Act, adopted in December 2025 (31), will make the obligations of developers clearer and somewhat more manageable.

We believe that the main issue is that a one-size-fits-all approach might not be the best solution, especially in biomedical engineering. In the literature, transparency is considered the most important requirement for achieving a trustworthy AI system (21). Yet, in biomedical engineering, reproducibility is far more important. Reproducibility means that the AI system always gives the same answer to the same question. For example, it is necessary that whenever it analyzes the same CT scan, it provides a consistent diagnostic output. Also, when predicting chemical reactions in experiments, the system must always generate the same prediction for the same inputs.

The reproducibility requirement is not even mentioned in the requirements for a high-risk AI system, although it could be understood as a form of an accuracy requirement. Further, it is very important that the AI system be well-trained and that inputs include a variety of data. When input data comes from a single group (e.g., people of a certain age or from certain areas), AI models can adjust too closely to that group's characteristics, leading to incorrect diagnoses or potentially dangerous applications in medicine. Hence, a variety of input data is a critical requirement. Finally, medical data is multidimensional, which means that to make conclusions in medicine, you need data of different types, such as images, numbers, and texts, which makes the creation of specific AI systems in medicine difficult (32). Therefore, we strongly believe that AI regulation must be tailored to the specific field in which it is applied, with rules specifically designed for that field. As we can see, that is also the intention of the European Commission, which has already adopted the proposal for the European Biotech Act.

The research conducted allows for several conclusions regarding the initial questions. The EU AI Act represents a significant first step in regulating AI, serving as a foundation for the further development and refinement of AI governance. Unlike other emerging fields, AI regulation has not been widely addressed at the national level within EU member states, making the process more challenging due to the lack of prior experience. Consequently, regulatory shortcomings are to be expected, and additional time and experience will be required to develop effective solutions. The EU AI Act provides a valuable starting point within the regulatory hierarchy. While AI is a complex social phenomenon that

cannot be comprehensively governed by a single law, an overarching regulation, similar to the GDPR, may serve as a constitutional framework for this legal domain. As sector-specific AI regulations are developed, the EU AI Act is expected to realize its full potential and resolve current inconsistencies. Limiting the EU AI Act to principled rules would enable the creation of tailored laws for each sector, reflecting the varying levels of risk and application across different fields.

Acknowledgements

This study was not supported by any sponsor or funder.

Author contributions

Conceptualization: N.I. and N.M.; Methodology: S.Dj., N.I., and N.M.; Investigation: S.Dj., N.I., and N.M.; Data curation: S.Dj., N.I., and N.M.; Formal analysis: S.Dj.; Writing – original draft: S.Dj., N.I., and N.M.; Writing – review & editing: S.Dj. and N.I. All authors have read and approved the published version of the manuscript.

REFERENCES

- Hupont I, Micheli M, Delipetrev B, Gomez E, Garrido JS. Documenting high-risk AI: a European regulatory perspective. *Computer*. 2023;56(5):18-27. [\[CrossRef\]](#)
- Stettinger G, Weissensteiner P, Khastgir S. Trustworthiness assurance assessment for high-risk AI-based systems. *IEEE Access*. 2024;12:22718-45. [\[CrossRef\]](#)
- Smuha NA. The Work of the High-Level Expert Group on AI as the Precursor of the AI Act. In: Ceyhun N, editor. *AI Governance and Liability in Europe – A primer*. Alphen aan den Rijn: Kluwer Law International; 2025. ISBN: 9789403528441.
- Akhtar ZB, Rawol AT. Innovative Approaches in Biomedical Engineering (BME) and Medical Science through Artificial Intelligence (AI) and Enhanced Computing. *Trends Telemed E-Health*. 2025;5(3):1-14. [\[CrossRef\]](#)
- Sargiotis D. Leveraging AI and Big Data for Advancements in Biomedical Research. *Biomed J Sci Tech Res*. 2024;57(4):49586-92. [\[CrossRef\]](#)
- da Silva RGL. The advancement of artificial intelligence in biomedical research and health innovation: challenges and opportunities in emerging economies. *Globalization and Health*. 2024;20:1-19. [\[CrossRef\]](#)
- Tripathi D, Hajra K, Mulukutla A, Shreshtha R, Maity D. Artificial Intelligence in Biomedical Engineering and Its Influence on Healthcare Structure: Current and Future Prospects. *Bioengineering*. 2025;12(2):163-81. [\[CrossRef\]](#)
- Yaman F, Adler A, Beal J. Opportunities and Challenges in Applying Artificial Intelligence to Bioengineering. In: Liò P, editor. *Automated reasoning for systems biology and medicine*. Cham: Springer; 2019. p. 425-52. [\[CrossRef\]](#)
- Sardanelli F, Castiglioni I, Colarieti A, Schiaffino S, Leo GD. Artificial intelligence (AI) in biomedical research: discussion on authors' declaration of AI in their articles title. *Eur Radiol Exp*. 2023;7(1):12-8. [\[CrossRef\]](#)
- European Parliament; Council of the European Union. Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). *Official Journal of the European Union*. [Internet]. 2024; OJ L, 2024/1689. Available from: eur-lex.europa.eu
- Gikay AA, Lau PL, Sengul C, Miron A, Malin B. High-Risk Artificial Intelligence Systems under the European Union's Artificial Intelligence Act: Systemic Flaws and Practical Challenges. *SSRN Electronic J*. 2023;1-22. [\[CrossRef\]](#)

Statement of Ethics

Not applicable.

Statement of Competing Interest

The authors declare no relevant conflicts of interest.

Statement of Data Availability

Not applicable.

Statement of Generative AI Use

No generative AI was used.

Publisher's Note: The statements, opinions, and data contained in AFMN Biomedicine articles are solely those of the individual author(s) and contributor(s) and do not necessarily represent the views of the publisher or the editor(s). The publisher and editor(s) disclaim responsibility for any harm or damage caused by the use of information or products mentioned in the publication.

12. Aboy M, Minssen T, Vayena E. Navigating the EU AI Act: implications for regulated digital medical products. *NPJ Digit Med.* 2024;7:1-6. [\[CrossRef\]](#)
13. Djeflal C, Mehl P, Müller V. The EU AI Act's Impacts on Digital Health. *Curr Dir Biomed Eng.* 2024;10(4):191-5. [\[CrossRef\]](#)
14. Ebers M. Truly Risk-based Regulation of Artificial Intelligence. *Eur J Risk Regulat.* 2024;1-20. [\[CrossRef\]](#)
15. Fink M. Human Oversight under Article 14 of the EU AI Act. *SSRN Electronic J.* 2025;1-15. [\[CrossRef\]](#)
16. Pehlivan CN, Forgó N, Valcke P. The EU Artificial Intelligence (AI) Act: A Commentary. Alphen aan den Rijn: Kluwer Law International; 2024. ISBN: 9789403541228 [\[CrossRef\]](#)
17. Feiler L, Forgó N, Nebel M. The EU AI Act: A Commentary. Surrey: Globe Law and Business Ltd; 2025. ISBN: 9781787429943.
18. Constantino J. Exploring Article 14 of the EU AI Proposal: Human in the Loop Challenges When Overseeing High-Risk AI Systems in Public Service Organisations. *Amsterdam Law Forum.* 2022;14(3):1-17. Available from: [reference-global.com](#)
19. De Cooman J. Humpty Dumpty and High-Risk AI Systems: The Ratione Materiae Dimension of the Proposal for an EU Artificial Intelligence Act. *Market Comput Law Rev.* 2022;6(1):49-88. [\[CrossRef\]](#)
20. Finck M. The EU Artificial Intelligence Act: A Commentary. Oxford: Oxford University Press; 2026. ISBN: 9780192889218.
21. Golpayegani D, Hupont I, Panigutti C, Pandit HJ, Schade S, O'Sullivan D, et al. AI Cards: Towards an Applied Framework for Machine-Readable AI and Risk Documentation Inspired by the EU AI Act. In: *Privacy Technologies and Policy. Lecture Notes in Computer Science.* Cham: Springer; 2024. p. 48-72. [\[CrossRef\]](#)
22. Wagner M, Song Q, Borg M, Engström E, Lysek M. AI Act High-Risk AI Compliance Challenge and Industry Impact: A Multiple Case Study. *SSRN Electronic J.* 2025;1-51. [\[CrossRef\]](#)
23. High-Level Expert Group on Artificial Intelligence. Ethics guidelines for trustworthy AI [Internet]. Brussels: European Commission; 2019. Available from: [digital-strategy.ec.europa.eu](#)
24. Kalodanis K, Rizomiliotis P, Feretzakis G, Papapavlou C, Anagnostopoulos D. High-Risk AI Systems—Lie Detection Application. *Future Internet.* 2025;17(1):26-49. [\[CrossRef\]](#)
25. Neuman KL, Cory-Wright D, Hespeler CB, White M. European Commission's Proposed Regulation on Artificial Intelligence: Conducting a Conformity Assessment for High-Risk AI—Say What? *RAIL: J Robotics Artif Intell.* 2022;5(2):135-144. [\[CrossRef\]](#)
26. Truby J, Brown RD, Ibrahim IA, Parellada OC. A Sandbox Approach to Regulating High-Risk Artificial Intelligence Applications. *Eur J Risk Regulat.* 2022;13(2):270-294 [\[CrossRef\]](#)
27. Grieman K, Early J. A Risk-based Approach to AI Regulation: System Categorisation and Explainable AI Practices. *SCRIPTed: A Journal of Law, Technology & Society.* 2023;20(1):56-88. [\[CrossRef\]](#)
28. Grozdanovski L, De Cooman J. Of Hypothesis and Facts: The Curious Origins of the EU's Regulation of High-Risk AI. *Eur J Law Reform.* 2022;24(1):123-134. [\[CrossRef\]](#)
29. Cabral TS. Rethinking the List-Based Approach to High-Risk Systems under the AI Act. *Nordic J Eur Law.* 2025;8(1):32-48. [\[CrossRef\]](#)
30. Joint Artificial Intelligence Board; Medical Device Coordination Group. Interplay between the Medical Devices Regulation (MDR) & In vitro Diagnostic Medical Devices Regulation (IVDR) and the Artificial Intelligence Act (AIA): document ref: MDCG 2025-6 [Internet]. Brussels: European Commission; 2025. Available from: [health.ec.europa.eu](#)
31. European Commission. Proposal for a Regulation on establishing a framework of measures for strengthening Union's biotechnology and biomanufacturing sectors... (European Biotech Act): COM(2025) 1022 final [Internet]. Brussels: European Commission; 2025. Available from: [health.ec.europa.eu](#)
32. Han H. Challenges of reproducible AI in biomedical data science. *BMC Med Genomics.* 2025;18(1):1-6. [\[CrossRef\]](#)