

JOANNA HOŁDA

MARIA CURIE-SKŁODOWSKA UNIVERSITY IN LUBLIN

JOANNA.HOLDA@MAIL.UMCS.PL

[HTTPS://ORCID.ORG/0000-0001-9334-8441](https://orcid.org/0000-0001-9334-8441)

Utopia or Reality? Will It Be Possible to Create an Effective Artificial Intelligence Standards System?

Abstract: Artificial intelligence (AI) is found in all areas of modern life, and its practical importance continues to grow. It creates opportunities, generates problems and risks. Its development is connected with ethical and formal issues related to legal regulations. The purpose of the article is to raise awareness of the achievements to date in the field of international AI standards and to reflect on the possibility of developing such cooperation. This article addresses the issue of universal standards for the security of AI, rather than all the ways and circumstances related to its application. It seems that in view of the dynamic development of AI, the challenge for today is the adoption of universal AI standards consistent with the principles of democracy, human rights and the rule of law, and the most effective would be the adoption of such a system in the dimension of global international cooperation (as the universal system of human rights was formed years ago). Also fundamental to this essay is the question of whether regulation of even selected aspects of AI is possible today on a global level – whether it is still a utopia or already a necessity arising from the need of today.

Keywords: artificial intelligence; international cooperation; law; standards; utopia

Introduction

The development of artificial intelligence (AI) poses challenges and risks that concern economic, political, and social or scientific issues, but also affect the lives of citizens. In view of the fact that AI is transboundary in nature, not tied to one place – the state, or the environment – it requires extensive cooperation between different actors – private and public (authorities at different levels, businesses, universities, specialists in different fields). AI systems are used in almost all areas of modern life, from manufacturing and medicine to the military, education and the activities of the media (especially new ones). Its development touches on ethical and formal issues related to legal regulations. Thus,

on the one hand, the challenge is the limits of AI, especially in an ethical context, which is the task of ethicists or philosophers; on the other hand, it involves how to shape legal regulations and standards (including non-legal ones), which can largely be referred to social science research. This article addresses the issue of universal AI standards, rather than all the ways and circumstances surrounding their application. The aim is to raise awareness of the existing body of work on international AI standards and to reflect on the possibility of developing such cooperation. It seems that, in view of the dynamic development of AI, the challenge for today is to adopt universal AI standards compatible with the principles of democracy, human rights and the rule of law; and the most effective approach would be the adoption of such a system in the context of global international cooperation (as the universal system of human rights was established years ago). Crucial for this essay is also the question of whether regulating even selected aspects of AI today is possible on a global level – whether this is still a utopia or already a necessity arising from the needs of today's times.

Literature review

Recent years have seen the dynamic development of AI and, consequently, the associated scientific and publishing activities (in various disciplines and fields). The increased popularity in the area in question has resulted in more and more new research problems being addressed. In the context of this essay, publications focusing on legal issues, international cooperation (also on the general principles and functions of international relations), or addressing human rights are relevant. The publication intensity, however, is not only scientific texts, but also in popular science and journalism – which indicates the great practical significance of AI and inspires a broad social discussion concerning its specifics, challenges, or threats. An undoubted value is the relatively easily and quickly accessible flow of information in the field under discussion. In preparing this essay, several positions were particularly important. One of the most interesting is a comprehensive study touching on legal, ethical and policy issues in the European Union, which was inspired by the process of introducing AI provisions in EU law (Braun & Harasimiuk, 2021). Also noteworthy are relatively new publications addressing AI issues in relation to international law (Lee, 2022), or discussing ethical issues in the context of international law (Abhivardhan, 2023). There are also interesting reflections on the relevance of AI in international law in the face of challenges to existing legal systems, particularly concerning smart weapons, cybersecurity and data protection (Liu, 2024).

Research on AI in international relations mainly focuses on four key themes: balance of power, disinformation, governance and ethics. However, in view of the (re)conceptualization of the technology, or the arms race, the established disciplinary framework is being broadened (Bode, 2024). One of the biggest challenges is its po-

sitioning within human rights standards. With the emergence of new technologies at least such as ChatGPT, new concerns about their application are emerging, as has also been recognized (Dulka, 2023). The literature also critically analyses how and to what extent AI can violate human rights and/or lead to socially harmful consequences and how to avoid this (Završnik & Simončič, 2023). It is also worth looking at popular science publications (in this case also available in Polish) discussing AI in various aspects including the impact on security and information (media/news) activities (Sumpter, 2019), or addressing issues at the interface of AI and democracy (O'Neil, 2017). In a broader, political context, Applebaum points out how new technologies help build authoritarian systems (Applebaum, 2024; cf. Harari, 2024).

The analysis of the available scientific literature allows us to assume that the multidimensionality of AI in the context of a possible common international (global) legal regulation is complicated by the continuous technological development generating new challenges and problems. Its multidimensionality concerns not only legislative issues, but also the creation of international cooperation. In view of this, a broader view from the perspective on economic and political interests is needed, and analyses of these topics are increasingly present in journalistic texts (readily and quickly available in the media), which adds value to deepening knowledge of the importance of AI and its impact on international relations in general. It is through journalism that factual analysis is conducted most swiftly, also providing inspiration for scientific research as well. Scientific literature is crucial in shaping the factual and methodological foundations, and current information on the political, economic or military situation is a natural complement.

Where we are with AI standards

When considering the possibility of systemic solutions in the form of international minimum safety standards for AI, it is worth considering where we currently are in such a process and whether it is possible – from a legislative, nature of international cooperation, technological or political point of view.

Talking about AI standards, I am thinking of a set of values to which they should refer to. The Council of Europe Framework Convention (Convention of AI, 2024), discussed below, can serve as a point of reference. The following should be considered key standards: protection of human rights; protection of democratic processes and respect for the rule of law; AI transparency; responsibility for adverse impacts on human rights, democracy and the rule of law; privacy and personal data protection; reliability of AI systems; supporting safe innovation; establishing a risk and impact management framework.

The leaders in AI should be considered to be the USA, China, the UK, Israel, Canada, France, India, Japan, Germany or Singapore, i.e. also the most important

world economies, located in different parts of the globe. To a greater or lesser extent, national standards or policies are being introduced there (e.g. the UK, Switzerland, Brazil, Canada, China, Japan, India, Singapore or Australia) regarding AI, although this is more often done through soft law rather than comprehensive legal regulation, the exception to this is the European Union and legislation introduced in member states, in this case France and Germany. This demonstrates the stature of the issue. Relevant to the creation of international AI standards is the executive order adopted by President Joe Biden in 2023 (Federal Register, 2023) concerning, among other things, security, the fight against deepfakes, rules for the creation, purchase and use of AI systems. It indicated the direction of US thinking on technological development. Even if this document is considered too restrictive and will need regular updates as technology changes, it is an important benchmark given the US experience to date as a leader in international cooperation. The question, however, is how President Donald Trump's administration will act in the area under discussion. The discussion on the development and security of AI is taking place in various forums around the world. This includes specialists, experts, scholars, but also politicians. Such discussions are also occurring in the media, which enhances public awareness.

Regulation of AI requires restraint and encouragement of innovation. Undoubtedly, common standards are important for developing international economic, technological, scientific, political or cultural cooperation. It is also consistent with the foundations of international cooperation, which aim, among other things, to address global problems and challenges, to develop innovation, to exchange experiences and good practices in various fields, and to promote peace and security. This is done through the adoption of laws, but before that by agreeing positions, exchanging knowledge and information and participating in international bodies and organizations. AI security issues are therefore within the scope of the international community and can be addressed through cooperation and the adoption of systemic solutions. Recent years indicate that such cooperation has been undertaken, which may forecast deepening work on global AI standards.

An important example is the AI standards document adopted in 2021 at the initiative of UNESCO (by more than 190 countries) (UNESCO, 2021). It is a non-binding normative instrument, providing a set of guidelines that member states will implement in a manner tailored to their own circumstances and capacities. The document sets out several objectives. One is to establish a universal canon of values, principles, and actions to guide states in developing AI laws, policies or other instruments, in a manner consistent with international law. In addition to this, it is intended to provide directions to various actors to address ethical issues at all stages of the life cycle of AI systems. The document aims to strengthen human rights, protect the interests of present and future generations, protect the environment, and respect cultural diversity at all stages of the AI system life cycle. It also aims to promote open dialogue and strive to reach and underdeveloped countries. It is worth pointing out at this point

UNESCO's other undisputed contributions – the International Research Centre on Artificial Intelligence (IRCAI) has been inaugurated under its auspices and is poised to become a global network of institutions and experts promoting cutting-edge AI research to help achieve the Sustainable Development Goals (Garcia, 2021). Another document strengthening AI standards was adopted at the AI Safety Summit 2023 event in November 2023 (Government of the UK, 2023) at Bletchley Park in Buckinghamshire (the symbolic site of the work to break the Enigma code). It was a UK initiative and the document (declaratory in nature, not universally binding international law) was adopted by nearly 30 countries belonging to different political systems and from different parts of the world: Saudi Arabia, Australia, Brazil, Chile, China, the Emirates, the Philippines, France, Germany, India, Indonesia, Ireland, Israel, Italy, Japan, Canada, Kenya, Korea, Nigeria, New Zealand (in 2024), Rwanda, Singapore, Spain, Switzerland, Turkey, Ukraine, the European Union, the United Kingdom and the USA. The aim is to strengthen global cooperation in the field of AI security and foster technological development. Overall, the main thrust of the declaration revolves around ethics, safety and risk mitigation of AI, fostering international cooperation, innovation and development. The document recognizes the potential, values, challenges and risks of AI thereby assuming the need to build a safe environment for AI. As the statement indicates, the global potential of AI can improve human well-being, peace and prosperity. Therefore, it should be designed, developed, deployed and used in a safe, human-centred, trustworthy and responsible manner.

In view of the fact that many threats are inherently international in nature, they are therefore best countered by building strong cooperation between states. Particular attention is being paid to cybersecurity, biotechnology and areas where pioneering AI systems may increase the risk of, for example, disinformation (i.e. in the media). In view of the likelihood of serious damage, it is important and urgent to take action to address it. International cooperation is therefore aimed at countering the risks associated with pioneering AI, identifying risks, building policies to ensure security. Important in terms of establishing where we are in the process of developing global, universal AI security standards is the statement, with a code of conduct towards AI, of the G7 countries at the Hiroshima Process in October 2023 (European Commission, 2023), on the timeline – adopted moments before the Bletchley initiative. The involvement of the world's richest countries, the economic leaders, is of much more than symbolic importance and warrants recognition of their support in the process of developing appropriate standards.

At the regional, European level, a law has emerged outlining similar objectives to those in the above-mentioned documents. The legislation in question is that of the European Union (Artificial Intelligence Act, 2024), which encompasses all 27 Member States, thus strengthening the EU *acquis* in the area of new technology law (in addition to legislation on, *inter alia*, data protection, markets, digital services, or the European Freedom of the Media Act, where the issue of technology, and the digital

environment is a priority and more or less directly affects AI). Although it covers the EU internal market, it nevertheless has a broader significance (alongside the Council of Europe Convention, it introduces the first legal solutions in the world in the area of AI), even if it is not perfect and will need to be revised in the future. The aim of the AI Act is to promote human-centered and trustworthy AI. Respect for fundamental rights is the essence of this law (Namysłowska et al., 2024). It is intended to unify the EU internal market, introducing a coherent framework for the development and use of AI systems, in line with EU values, as well as to stimulate innovation and employment, strengthening the EU's position as a leader in trustworthy AI.

The EU legislator aims to avoid fragmentation of the internal market by establishing consistent rules for AI operators, while protecting the public interest and rights of those using these systems. The AI Act covers suppliers, importers and other entities using AI-based systems operating in the European Economic Area. The Regulation will also apply to entities based outside the EEA as long as they make their services available to the EU. The Regulation distinguishes four risk categories for the use of AI-based systems: low risk category (i.e. systems that do not pose a significant threat), medium risk category (e.g. ChatGPT), high risk category (technologies that may impact the security and fundamental rights of the user), unacceptable risk category (systems that pose a huge security risk, e.g. systems designed for social assessment).

An important example of international cooperation in the area of AI standards development, and at the same time the world's first international agreement to cover this issue, is the Council of Europe Framework Convention. It can be considered a milestone in the creation of global, universal AI security standards. It touches on human rights and democracy issues (it overlaps with the tenets of the aforementioned AI Act; the EU is a party to this convention). It focuses on safe AI that respect human rights and democratic values, considering transparency and reliability of systems as important, thus, creating principles of trustworthy AI. The value of the document is the introduction of oversight mechanisms for AI. It is important to note that the document was also agreed and negotiated by non-European countries (e.g. USA, Canada, Japan, Israel, Australia, Vatican, South American countries), which gives it a strong legitimacy and a mandate.

The above examples are of course not exhaustive of all the activities of international bodies, politicians or organizations, but in my opinion they are the most important. The interest of representatives from communities in different parts of the globe demonstrates the need for universal AI standards. Of course, activities in the form of AI policies or strategies undertaken by individual countries are valuable, but due to the nature of AI, globalization and dynamic technological development, it is particularly important to build common, strong standards agreed and accepted by the global international community (which will also be implemented in national or, for example, international-regional systems). The existing *acquis* in the form of recom-

recommendations or declarations is a good foundation for deepening global cooperation. The European (EU and Council of Europe) acquis, on the other hand, provides an important reference point for future global regulations.

Perspectives for effective and universal AI standards in a global context

It is fundamental to assess whether the indicated phenomenon/problem is at all within the scope of interest of the international community and can be subject to legal regulation at this level. AI, as indicated above, not only has legislative potential, but the relevant processes have already begun, and this is a direct result of the existence of real, current needs felt from a global perspective. Relevant from the point of view of this essay is to determine what kind of standards system is effective. On a theoretical level, it is therefore worth attempting to define the concept of such “effectiveness”. An effective system is a legal system, i.e. a defined set of standards.

While the standards set out in recommendations or declarations, i.e. at the soft law level, are of great importance and can be an important point of reference for future regulation, they are far less effective because they are not legally binding, and compliance with them is based on the goodwill of states (even those that have “signed up” to them). The legal system is therefore the most effective. From my perspective, this means legislation that is adequate, proportionate and up-to-date. Putting in place mechanisms and tools to monitor and control compliance, the consequences of non-compliance and possible sanctions. It is equally important to establish the possibility of recourse to an appropriate judicial or administrative body. The rules that constitute an effective system are those that have been established as a result of consultations among various actors and circles – potentially interested parties – specialists, and experts in various fields, including ethicists, programmers, sociologists, business representatives, as well as politicians, who will ultimately implement them in their systems. There are many stakeholders here; different interests and values clash, economic and political alliances matter. An effective system is one that is well thought out, accepted and understood – in this case by the international community. A system created as a result of international – global – cooperation enhances security. In addition, by design, it is open to all countries in the world, regardless of their location, political culture or wealth – this provides additional reinforcement. That is to say, several elements are important – one of them being the regulation of the current state of affairs and the resulting needs, or the solution to an existing problem.

The timeliness of regulation of AI systems can be a challenge, which can undermine their effectiveness. Technological developments and innovations may “leapfrog” existing regulations and render them obsolete, even if the solutions adopted are relatively broad and general, precisely to cover newer types of systems emerging in the future. This also links to adequacy and proportionality (e.g. in relation to sanctions).

This will consequently require revision, update and amendment of the law. This is, of course, achievable, but any legislative process is time consuming, each time requiring reconciliation of positions and interests of those potentially involved. That is, potential legislation at the outset may require constant monitoring and possible revision. However, it is important to remember that the area under discussion is specific and new, so the occurrence of risks is inherent.

Regardless of the difficulties that may arise, the global societal (economic, cultural) interest is strong enough to make it worth taking risks and facing them on an ongoing basis. And even if regulations are not fully effective, their presence is welcome and positive. It is also important to be aware that AI security standards will take years to develop. International cooperation is based on voluntary participation and the goodwill of the associating states. Looking at the current state of international cooperation, it may be questionable whether states pledging to adopt safe AI standards will in fact do so. The weakness may be due to the lack of strong mechanisms to counter non-compliance with the standards. This can be seen in various situations where international law is violated (e.g. Russia's invasion of Ukraine and the ineffectiveness of the UN system are a glaring examples). Besides, we see a certain crisis of international cooperation in general, and it needs to be redefined also in view of emerging challenges. One of them is the global digital consensus (Stimson Center, 2022), which is directly linked to the development of AI. Another impediment to effective regulation is the crisis of international cooperation institutionally, primarily in the United Nations system (Hosli, 2021). This is relevant for building an effective global system based on an international organization (acting as initiator, process leader, oversight body), especially from the group of UN specialized organizations.

The situation is better in Europe; not without significance is the strong integration based on a common axiological system, sharing democratic values and the condition of human rights (although the situations in European countries, or international institutions and organizations, are not free from criticism). The importance for the effectiveness of the built system is therefore the shared values. Also of importance is the effect of scale, i.e. the number of participants in the process (states) – in Europe this seems easier due to the number of states potentially involved.

Conclusion

Answering the fundamental question of the essay, it should be assumed that the proposal to introduce universal regulations (standards) for AI security is not a utopia. It is a need arising from today's reality. However, we will probably still have to wait for global (universal, addressed to various entities of the international community regardless of location or wealth) legal regulations in this area, despite the fact that we have an increasingly wide range of agreed principles. Today, Europe is emerging

as the pioneer of AI regulation – the legacy of the Council of Europe and the EU demonstrates this. The legislative courage deserves praise. The systems are cooperating with each other, which means that a unified system will be created, and the EU's legislative activity allows us to assume that new legislation will also be created in the field of new technologies in the future. At the same time, European systems have a long tradition of good cooperation; an example being the approach to biotechnology, an area similarly sensitive to AI. However, the question is whether, in addition to being the undisputed regulatory leader today, Europe will also be a leader in the development of technology and AI? This could indeed influence the emergence of trustworthy and secure AI and set real standards for others.

Admittedly, European countries are among the leaders in AI, but some anxiety may be caused by the current geopolitical situation in which Europe finds itself. The proximity of war is causing governments to declare increased spending on the military sector, which may reflect on support for the development of new technologies in Europe. The main challenge for today, however, is the creation of a unified global system of universal AI security standards. As the experience of recent years has shown, this is achievable, albeit difficult and time-consuming. The output to date, also in the form of declarations or recommendations, proves the interest of the international community in the topic under discussion and gives hope for the construction of global standards in the future. However, it is worth exercising caution and moderate optimism, as both substantive and procedural law issues may arise in this process.

It is important to emphasise, however, that already emerging standards take into account democratic values, human rights and the rule of law, which points in the direction of future solutions. A universal AI security system could be developed under the auspices of the United Nations, or specialized organizations, which would be natural. The UNESCO output indicated above is an important point of reference here, although one may wonder whether, in the long run, this organization is the most appropriate and prepared in terms of content. Perhaps, for reasons of tasks, competence and experience, the World Trade Organization (WTO) or the World Intellectual Property Organization (WIPO) should take a leadership role.

The perceived crisis in international cooperation and criticism of the UN system is not insignificant. The fact is that there is currently no clear leader in the AI standards process – despite activity, no international organization or specific country has an obvious and unquestioned position here. The current political and economic situation is also a challenge. We are seeing changes in the existing rhythm, structures and international leadership. Alliances are changing, and what seemed constant, e.g. the strong cooperation between Europe and the US and the drive towards regulatory harmonization, is changing with the arrival of President Trump's new administration. There is also some uncertainty about China's behavior in the context of the potential introduction of global AI security standards. At the very least, the atmosphere of uncertainty, unpredictability and emerging crisis is affecting the overall atmosphere

and climate for the introduction of the law, causing stagnation rather than accelerating processes. It is also conceivable that some countries, or regions, will opt for much more business-friendly AI regulations, avoiding excessive restrictions. This could be the case for countries outside the EU (they are not bound by the rigor in which member states remain), or Asian countries, which could, as it were, “pitch” the idea of a global AI security standards regime, or push it back.

References

- Abhivardhan. (2023). *Artificial Intelligence Ethics and International Law. Practical Approaches to AI Governance*. BPB Publications.
- Applebaum, A. (2024). *Koncern autokracja. Dyktatorzy, którzy chcą rządzić światem*. Wyd. Agora.
- Bode, I. (2024). AI technologies and international relations: Do we need new analytical frameworks? *The RUSI Journal*, 169(5), 66–74. <https://doi.org/10.1080/03071847.2024.2392394>
- Braun, T., & Harasimiuk, D.E. (2021). *Regulating Artificial Intelligence. Binary Ethics and the Law*. Routledge.
- Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, Vilnius, September 5, 2024.
- Dulka, A. (2023). The use of artificial intelligence in international human rights law. *Stanford Technology Law Review*, 26(2), 316–366.
- European Commission. (2023). *G7 Leaders’ Statement on the Hiroshima AI Process*.
- Federal Register. (2023). *Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, Executive Order 14110*.
- Garcia, E.V. (2021). *UNESCO’s Recommendation on the Ethics of AI: why it matters and what to expect from it*. www.researchgate.net/publication/357074719_UNESCO's_Recommendation_on_the_Ethics_of_AI_why_it_matters_and_what_to_expect_from_it
- Government of the UK. (2023). *The Bletchley Declaration by Countries Attending the AI Safety Summit*.
- Harari, Y.N. (2024). *Nexus. Krótka historia informacji. Od epoki kamienia do sztucznej inteligencji*. Wyd. Literackie
- Hosli, M.O. (2021). The United Nations and challenges to multilateralism. In T. Garrett, M.O. Hosli, S. Niedecken, & N. Verbeek (Eds.), *The Future of Multilateralism Global Cooperation and International Organizations* (pp. 3–17). Rowman & Littlefield Publishers.
- Lee, J. (2022). *Artificial Intelligence and International Law*. Springer
- Liu, J. (2024). Artificial intelligence and international law: The impact of emerging technologies on the global legal system. *Economics, Law and Policy*, 7(2), 73–82. <https://doi.org/10.22158/el.p.v7n2p73>
- Namysłowska, M., Bieda, R., Budrewicz, P., Lubasz, D., Nowakowski, M., Pająk, R., Świerczyński, M., Więckowski, Z., Wochlik, I., & Wróblewski, M. (2024). On the Ethical, Legal and Social Implications of Artificial Intelligence Systems in the Member States of European Union. Observations against the Background of the Draft Regulation on Artificial Intelligence. *Przegląd Sejmowy*, 6, 61–84.
- O’Neil, C. (2017). *Bron matematycznej zagłady. Jak algorytmy zwiększają nierówności i zagrażają demokracji*. Wyd. Nauk. PWN.

- Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised 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).
- Stimson Center. (2022). *Rethinking Global Cooperation. Three New Frameworks for Collective Action in an Age of Uncertainty*.
- Sumpter, D. (2019). *Osaczeni przez liczby. O algorytmach, które kontrolują nasze życie. Od Facebooka i Googla po fake newsy i bańki filtrujące*. Wyd. Copernicus Center Press.
- UNESCO. (2021). *UNESCO Recommendation on the Ethics of Artificial Intelligence*.
- Završnik, A., & Simončič, K. (2023). *Artificial Intelligence, Social Harms and Human Rights*. Palgrave Macmillan.

