



CORRUPTION AND HUMAN RIGHTS

Research article

<https://doi.org/10.71559/ACSAR.1.1.2026.31>

Fundamental legal principles for anti-corruption artificial intelligence in the public sector

J.Strelcenoks 

PhD, Head of the Corruption Prevention and Combating Bureau of Latvia (2011–2016)

E-mail: jaroslavs.strelcenoks@gmail.com

<https://orcid.org/0009-0009-7209-6113>

Turība University, Latvia

Abstract

This article analyzes the transformative potential of artificial intelligence (AI) technologies in preventing corruption in the public sector. The research focuses on systematizing the key principles that should underpin the ethical and effective implementation of AI technologies in anti-corruption practices. The author emphasizes that AI serves as an effective tool, complementing human analysis and institutional mechanisms in preventing corruption. This assertion was emphasized at the recent Conference of the States Parties to the UN Convention against Corruption, the main outcome of which was a call to use AI technologies to strengthen international law enforcement cooperation in the fight against transborder corruption and to improve the accuracy, effectiveness, and objectivity of national efforts to assess corruption risks, as well as to realize the potential of AI and manage its risks, including through strengthened regulation and human oversight. The research is based on a comparative analysis of key principles for the use of AI technologies developed and proposed to the global community at the national level in a number of countries (e.g., the Vatican, China, the United States, and Singapore) and at the international level by organizations such as UNESCO, the Council of Europe, the European Parliament, and the OECD. The international community has supplemented the author's proposed principles based on his own research, which identified the most important legal principles for implementing AI technologies in the public sector. The research identified common principles that have proven effective: prioritizing data quality and reliability; a human-centered approach that complements, rather than replaces, expert analysis; transparency and explainability of AI algorithms to ensure trust and accountability; respect for human rights and protection of privacy; and the need to adapt legislation and develop digital competencies in the public sector.



Implementing AI based on these principles can significantly strengthen preventive anti-corruption controls.

Keywords: anticorruption, artificial intelligence, public sector, corruption prevention, legal principles

How to cite: Strelcenoks, J. (2026). Fundamental legal principles for anti-corruption artificial intelligence in the public sector. *Anti-Corruption Studies and Research, 1*(1), 66-81. <https://doi.org/10.71559/ACSAR.1.1.2026.31>

1. Introduction

With the rapid development of artificial intelligence (AI) technologies, powers and decision-making are increasingly being transferred from humans to AI for more effective resolution. This raises the question of whether AI can significantly reduce the risks of corruption and conflicts of interest in government and municipal institutions. This question was answered in Doha, Qatar, from December 15 to 19, 2025, where the 11th session of the Conference of the States Parties to the UN Convention against Corruption was held. The main outcome of the conference was the “Doha 2025 Declaration”, which calls for the use of AI and digital technologies to more effectively combat corruption. The declaration makes clear that responsibly using AI technologies is essential to shaping the integrity of tomorrow. The Doha 2025 Declaration also calls on States to use AI to strengthen international law enforcement cooperation in the fight against cross-border corruption and to improve the accuracy, effectiveness, and objectivity of national efforts to assess corruption risks, as well as to realize the potential of AI and manage its risks, including through enhanced regulation and human oversight.[1]

In this case, if AI is authorized to assess corruption risks, it will be necessary to determine the principles of the legal regulation mechanism for AI technologies. However, it is already necessary to emphasize that these principles should contribute to the regulation of social relations and the promotion of the interests of the state and society.

The issues of legal regulation of the legal and ethical principles of the interaction of a person, society and public sector institutions are a necessary condition for the guarantee of constitutional norms that strengthen the rights and freedom of the person as the highest value. AI technologies are rapidly moving towards interaction with people, society and the state. Consequently, when faced with AI, the constitutional norms guaranteed to the person must be preserved. However, the interests of the state and society must also be observed in order to prevent corruption and conflict of interest situations in the public sector.

Once upon a time, science fiction writer Isaac Asimov formulated the “laws of robotics” in 1942. Later, these rules were recognized as universal and formed the basis for the development of robotics and AI: (1) A robot may not harm a person or, through inaction, allow a person to come to harm; (2) A robot must obey all human orders, unless these orders contradict the first law; (3) A robot must ensure its own safety unless this conflicts with the first or second law.[2]

These rules are informal ethical standards for robotics developers. However, the more advanced cyber-physical and other AI systems become, the greater the need to develop specific legal principles and norms governing the use of AI.

According to the author, when AI takes on public sector management tasks and is based on big data, it is important to develop and adopt in advance the ethics of AI’s own decision-making and

operation. The ethics of programmers who create the algorithms according to which the AI system operates and learns itself.

2. Literature Review

The research conducted in the article is based on legal sources and scientific literature, the author's publications, as well as the author's research. In solving the tasks set, the works of both Latvian and foreign researchers were used, mainly in Latvian and English. In order to additionally analyze the situation in practice, other literary sources and publications, both in printed publications and on Internet resources, were also used. Also, the article uses mostly primary (initial) and secondary data, but in some cases also tertiary data.

The theoretical and methodological basis of the research is mainly formed by the opinions of foreign authors from interdisciplinary fields (*Azoulay, A., Demaidi, M.N., Nicolet-Serra, L., Overly, S., Heikkilä, M., Roberts, H., Cowls, J., Sterling, B., Wagner, B., William M., Zou, G., Zhang, H., and others*), foreign legal acts (both valid and invalid, which allow us to trace the development of the legal regulation of AI and conflict of interest), as well as practical research: the author's research on the existence and prospects of AI technologies in the public sector, guidelines, reports, statistical reports, etc. The doctoral thesis is the first research of this level on the legal regulation of AI in the anti-corruption sphere in Latvia in relation to the prevention of conflict of interest.

During the research period, the author has participated in the development of legal regulation on corruption and conflict of interest, which gives the author advantages in analyzing the goals of creating, amending, and eliminating anti-corruption and conflict of interest norms and the needs of specific periods, as well as the advantages of introducing AI technologies as a tool in the anti-corruption sphere for preventing conflict of interest.

The basis of the foreign theoretical base in this sector is mainly philosophical and technical research and publications on the legal framework of AI, basic requirements and basic principles of implementation, as well as ethical aspects. The most well-known approaches in the law of implementing AI technologies include: 1) the lag of legal framework from the development of AI (USA, China); 2) the promotion of AI development through legal framework (Germany) and 3) the slowing down of AI development or legal framework before the development of AI (France, England, EU).

Over the past five years, interest in AI and its legal regulation has been so great that publications number in the thousands. The author's article is in the field of law, which concerns the use of AI technologies in the anti-corruption sphere to prevent conflicts of interest. Philosophical views are also the basis for protecting the interests of public officials and society due to the risks of introducing AI technologies and possible harm to humanity, including public officials, their relatives and business partners.

In the article, the author is interested in strengthening AI technologies in national regulatory enactments and, with their help, ensuring that public administration prevents conflict of interest situations in the activities of public officials in the public sector.

3. Methods

The structure of the article is designed to achieve its goal and, while including various legal interests (legal regulation of AI technologies; prevention of conflict of interest; securing the interests of public officials; respecting the interests of society), conduct a more comprehensive scientific research.

Research methods used in the article and the theoretical and methodological basis of the research.

In order to achieve the goal and implement the tasks set in the work, the author relies on and uses the following scientific research methods: logically constructive method, problem analysis method, comparative analysis method (to research and analyze legal norms in foreign countries, deducing what is common and different, as well as examining the issues of regulatory effectiveness regarding AI technologies in various legal systems: the continental legal system (for example, in Latvia, Germany, France, Austria), the Anglo-Saxon legal system (for example, in the USA, Canada) and the legal system of Asian countries (for example, in China, South Korea, Japan). The comparative method is used to evaluate and compare the works of national and foreign legal scholars on the importance, essence and development of norms regarding the legal framework regarding AI technologies in connection with EU recommendations and basic requirements in legal acts), as well as the empirical method (reflected in the author's survey of state officials and employees working in state and local government institutions on the current legal framework, future prospects and basic requirements for AI, as well as also a vision of how AI technologies can help prevent corruption and conflict of interest and what this requires from the legal framework).

The article analyzes the data from a survey conducted by the author on the basic requirements for the implementation of AI legal regulation, which is based on the balance of interests of society, the state and officials whose activities will be controlled by AI in the future, as well as future prospects for the implementation of AI in the public sector to prevent corruption and conflict of interest: who should implement AI, according to what operating principles AI should operate, etc.

To obtain information for the article, an analysis of scientific literature and regulatory acts, analysis of international experience and modeling were conducted. In this article, the author will attempt to determine the principles upon which the development of AI technologies should be based, based on the understandings of various entities of international law, namely international organizations, global corporations, and nation states. The author will define these principles themselves, compare them, and evaluate them based on the realities and expectations of government officials themselves.

4. Results and Discussion

Since 2016, most countries with a digital economy have entered the technology race to develop AI technologies. More than 60 countries have already adopted national strategies in the field of AI.[3] According to the author, this is why it is so important today to determine the ethical framework for the development of AI, limit the possibilities of its unethical use, and direct the energy of developers and the ideas of legislators in a direction that ensures maximum safety and benefit to society.

For a long time, the development of ethical documents on the rules of use of AI has been the main trend in their regulation.[4] Currently, more than a thousand different laws, guidelines, principles and codes dedicated to the ethics of AI have been adopted around the world. Most of them contain several key principles: security, confidentiality, non-discrimination, responsibility, etc.[5]

One of the most famous examples is the Asilomar Principles of AI, adopted in January 2017 in Asilomar (California, USA) by such global corporations as Google, Google DeepMind, Facebook, Apple, OpenAI and others. These principles set out the humane goals, values, perspectives for the creation, use and research of AI technologies, their relationship with science, culture, production, and the rejection of the use of AI technologies in activities hostile to humanity, for example:

- 1) AI systems should be safe and secure throughout their entire life cycle, and the regular operation of AI should be easily verifiable;

- 2) any participation of an autonomous AI system in a court decision should be sufficiently justified and accessible to the scrutiny of competent authorities;
- 3) AI systems with a high degree of autonomy should be designed so that their goals and behavior are consistent with human values throughout their work;
- 4) people should have the right to access, process and control personal data, if AI systems have the ability to analyze and use this data;
- 5) the application of AI systems to personal data should not unreasonably reduce the real or subjectively perceived freedom of people;
- 6) AI technologies should benefit as many people as possible;
- 7) the economic benefits created by AI should be widely distributed for the benefit of all humanity;
- 8) people should be given the opportunity to take over decision-making functions after or instead of AI to achieve objectivity and quality in decision making, etc.[6]

The **Vatican**, as the seat of the highest spiritual leadership of the Roman Catholic Church, takes a rather active position regarding the regulation of AI. At the Vatican's initiative, a document was signed with the UN, EU, IBM and Microsoft that regulates ethical rules in the field of AI.[7] The influential Vatican Institute and senior officials from the EU, UN, IBM and Microsoft have joined forces to promote the development of AI ethics by signing the Rome Call for an Ethics in AI.

The document outlines six key principles for the ethical development of AI: transparency; accountability; objectivity; reliability; security and privacy.[7]

The author notes that the Vatican, in defining these principles, wants them to guide the activities of AI technologies as an assistant to humanity, which should act for the benefit of man and contribute to his spiritual development. Therefore, this document states that *"AI technologies must never be used to exploit people, especially the most vulnerable. Instead, AI must be used to help people develop their capabilities and save the planet"*.[7]

The author is interested in the experience of some leading countries in the field of implementing AI technologies. The **Singapore** Academy of Law's Law Reform Committee published three reports in 2020 and 2021 discussing the ethical principles that should underpin the legal activities of AI technologies.[8]

Law and essential interests: AI systems should be designed and deployed in a way that complies with the law and does not infringe the essential interests of individuals protected by law – two key concerns regarding liability for AI systems are (i) the lack of mental state of the subject, such as knowledge or intent; and (ii) the AI system's "decision" to act is the result of a long causal chain involving different actors at different stages of the system's design and implementation.

Accounting for the impacts of AI systems: Designers of AI systems should consider the potential consequences of reasonably foreseeable impacts of AI systems throughout their life cycle. Existing principles are likely to be sufficient and can be relied upon to fairly allocate liability. However, policymakers may need to tailor interventions by developing principles specific to specific scenarios.

Well-being and safety: AI systems should be rational, fair and free from intentional or unintentional bias. The intended and unintended impacts of AI systems on well-being and safety indicators should be assessed and harms should be minimized, taking into account factors such as human emotions, empathy and personal attitudes.

Risk management for human well-being: AI system designers should properly assess and eliminate or control the risks of using AI systems in terms of safety and well-being. Policymakers

will need to consider whether mandatory risk management standards should be introduced, and then determine the form and specifics of these standards.

Respect for values and culture: AI systems should be designed taking into account social values, cultural diversity and values in different societies. It is particularly important for effective AI systems to take into account social values and cultural norms.

Transparency: When designing AI systems, they should be as transparent as possible and it should be understood how and why an AI system made a particular decision or acted as it did. Transparency means the ability to trace, explain, verify and interpret all aspects of an AI system and its outcomes to the extent possible. The aim is not only to regulate AI appropriately, but also to make it safe. One possible response to various AI system problems related to tracking, explaining, and verifying tasks is the requirement that AI systems build in mechanisms that record data input as intelligently as possible and ensure logical decision-making.

Accountability: Holding responsible individuals accountable for the malfunctioning of AI systems, depending on their roles and context.

Ethical Data Use: Appropriate privacy and personal data management practices to protect personal data.[9]

In 2017, **China** passed the Cybersecurity Law, which regulates the use of AI by network operators and service providers. The law states that *“when collecting and using personal information, network operators must adhere to the principles of lawfulness, fairness and transparency, disclose the rules for data collection and use, and clearly define the purpose, method and scope of information collection and use, and obtain the consent of the individual for information collection”*. [10]

In May 2019, the Chinese Academy of Sciences, Beijing Academy of AI, Peking University, Tsinghua University, the Institute of Automation and the Institute of Computer Technology, market leaders Baidu, Alibaba published the “Beijing AI Principles”. [11] The principles call for “building a human community with a shared future and realizing AI that is beneficial to humanity and nature,” and include, among others: serving human values such as privacy, dignity and freedom, constant attention to AI safety, inclusiveness, openness, supporting international cooperation, and long-term planning. Meanwhile, China's National Industrial Information Security Development Research Center issued an official document in January 2021, “Developing a New AI Infrastructure in 2020,” which also stated the need to *“strengthen the ability to assess, prevent and control risks and promote the healthy development of AI technology”*, as well as *“ensure the security of AI technology, product security, data security and application security, and eliminate security problems inherent in AI itself and related to its use”*. [12] The author notes that these regulations demonstrate China's intention to strengthen oversight of AI systems. This contradicts the publicly known stereotype that the development of AI technologies is unrelated to respecting human rights. On the contrary, these laws demonstrate the Chinese government's commitment to prioritizing human rights when implementing AI technologies.

The US President issued a guide to regulating AI applications in November 2020, which “sets out policy considerations that should guide the regulatory framework for the use of AI”. [13]

In January 2021, the US Congress passed legislation that tasked the National Institute of Standards and Technology with developing guidance for companies to mitigate the risks associated with the development and use of AI. [14] The legislation also establishes the National AI Initiatives Authority to develop and implement a US approach to the technology. [12]

Speaking at the Munich Security Conference in February 2021, US President Joe Biden stated that *“the United States should formulate rules that will govern the development of technologies and*

norms of behavior in cyberspace, AI, biotechnology, so that they are used for human advancement, not human destruction”.[12]

Despite the fact that the US is one of the leaders in the practical implementation of AI technologies, the author notes the fragmented approach to regulating this area. Specifically, regulation is largely delegated to the federal level or determined by the position of the US President. The author notes that this approach does not establish the US as a leader in defining the legal principles for the implementation of AI technologies.

Therefore, the author adheres to the basic legal principles at the national level, which, although reflected in different ways in each state, are generally consistent with the general basic principles in Table 1.

Table 1. Comparison of the main principles of artificial intelligence technologies in specific countries

The Basic Principles of AI Technologies			
Asilomar AI Principles	Vatican	Singapore	China
Safety and reliability	Transparency	Law and essential interests	Legality
Validity and verifiability	Accountability	Accounting for the impacts of AI systems	Justice
Conformity to human values	Objectivity	Well-being and security	Transparency
Access to processing and control of personal data	Reliability	Managing risks to human well-being	Disclosure of data collection and use rules
Prohibition of limiting the real or subjectively perceived freedom of people	Safety	Respect for values and culture	Clear definition of the purpose, methods and scope of collection and use of information
Benefits most people	Confidentiality	Transparency	Obtaining a person's consent to collect information
Economic benefit for all mankind		Accountability	
Human decision-making capacity		Ethical Data Use	

Source: The author has created a table for comparative analysis.

A noticeable role in the field of global legal regulation of AI is the active activity of international organizations. One of them is **United Nations Educational, Scientific and Cultural Organization (UNESCO)**, which in 2019 adopted a resolution on the development of ethical standards in the field of AI at its general conference.[15] At the end of 2021, UNESCO adopted the Recommendation on the ethics of AI.[16] This is actually the first attempt to bring the positions of the world community closer together and prepare the basis for further legal regulation of AI technologies at the global level.[17]

In this document, UNESCO recommends focusing on the development of AI according to the following principles: human rights, openness, accessibility, participation of all interested groups. In addition to these principles, the UNESCO Recommendation notes the need to develop principles of AI ethics, as well as the need to respect gender and social minority equality, paying special attention to overcoming the “digital inequality”. [18]

However, the author draws attention to another important activity of the international organization. Specifically, UNESCO, in developing the international legal framework for AI, conducted a large-scale international research[15], which generated diverse opinions on a number of important issues, including:

- 1) **AI regulation** – two main approaches: (1) prioritizing technological development and regulation after identifying new problems (USA, China); and (2) regulation based on the precautionary principle, excluding unforeseen changes to the status quo (EU countries);
- 2) issues related to **gender inequality and minority rights**; several countries (USA, EU) consider these issues to be fundamentally important in the implementation of AI, while other countries (China and Russian Federation) believe that these issues are beyond the level of development of AI technologies;
- 3) **access to technologies**: “digital inequality” between different countries, social groups; uneven level of education and economic development; monopolization of AI technologies by global players; UNESCO calls on countries to move from competition to open cooperation, improve the exchange of knowledge, resources and tools in the field of AI, and strengthen multilateral and international cooperation.[19]

Based on this research and its conclusions, UNESCO adopted several key concepts related to the principles of implementation of AI technologies, namely:

- 1) **proportionality**: AI technologies should not exceed predefined limits to achieve legitimate goals or objectives and should be appropriate to the content;
- 2) **human oversight and determination**: humans are ethically and legally responsible at all stages of the life (operation) cycle of AI systems;
- 3) **protection of the environment and the world**: AI systems should promote peaceful interactions between all living beings and respect the natural environment, in particular when extracting natural resources;
- 4) **gender mainstreaming**: AI technologies should not reproduce gender inequalities that exist in real life, in particular in terms of wages, representation, access and stereotypes.[20]

The author notes that the UNESCO principles will play a key international role in the legal regulation of AI technologies if they are supported by specific policy measures designed to help governments overcome these challenges, with the participation of a wide range of stakeholders, from business to civil society. The author also notes that the UNESCO principles can significantly assist states in developing their own systems for assessing the ethical impact of AI, as well as strategies for preventing, mitigating, and monitoring the risks associated with the implementation of AI technologies. However, the author believes that this can only happen if public awareness of AI technologies, the principles by which these technologies will or already operate, are increased, and everyone is educated about digital rights.

The author notes that the **Council of Europe (CoE)** also plays a leading role in shaping modern approaches and policy standards for the development of AI in the context of protecting human rights, democracy, and the rule of law. The CoE calls for the creation of effective oversight mechanisms and democratic structures for monitoring the development, improvement, and implementation of AI [21].

The **CoE has established working bodies in the field of AI** on its website:

- 1) Inter-sectoral Committee of experts on Human Rights Dimensions of automated data processing and different forms of AI (MSI-AUT);
- 2) The CoE Ad Hoc Committee on AI (CAHAI)[22], which operated from November 2019 to the end of 2021. The committee has prepared a comprehensive research on the development of AI, including an analysis of the risks and opportunities of the technology, and has researched the usefulness and feasibility of developing a comprehensive legal framework for AI, which is based on the CoE standards in the field of human rights, democracy and the rule of law. The Committee

has developed approaches to European AI regulation, which resulted in the preparation of the report “Possible elements of a legal framework on AI, based on the Council of Europe’s standards on human rights, democracy and the rule of law”[23], which was approved in December 2021 by the plenary session of the Ad Hoc Committee on AI (CAHAI) of the Council of Europe for further use in the preparation of elements of the legal framework. Based on the report, the Committee on AI of the Council of Europe (CAI) has developed of the CoE Convention on the legal framework for AI.[24] The aforementioned convention was opened for signature on 5 September 2024 in Vilnius at the CoE Conference of Ministers of Justice.[25]

The author notes that this Council of Europe Framework Convention on AI and Human Rights, Democracy and the Rule of Law is the first binding international treaty on AI, the purpose of which is to ensure that the principles of human rights, democracy, and the rule of law that exist in each country are also applied to the current and future challenges posed by AI. The importance of this convention, according to the author, also lies in the fact that it was developed by all Council of Europe member states, as well as Argentina, Australia, Canada, Costa Rica, the Vatican, Israel, Japan, Mexico, Peru, Uruguay, and the United States, with the participation of international organizations such as the OECD, the OSCE, and UNESCO, as well as more than 70 representatives from civil society, business, technology, and academia. The author notes that such broad and diverse participation in the development of the convention may serve as an indication that the convention takes into account various issues and the rights of various groups.

The Convention sets out fundamental principles and rules that not only protect human rights, democracy and the rule of law, but also promote progress and technological innovation. It complements existing international standards on human rights, democracy and the rule of law, and aims to fill any legal gaps that may arise as a result of rapid technological developments in the field of human rights, democracy and the rule of law, and aims to fill any gaps that may arise as a result of rapid technological developments in the field of human rights, democracy and the rule of law. The Convention is technologically neutral, but its implementation should follow the logic of a gradual and differentiated approach, taking into account the potential negative impacts of AI on human rights, democracy and the rule of law. The Convention applies to both the public and private sectors. It provides for limited exceptions from the scope for national security, research and development. The Convention obliges all future contracting parties to address the risks posed by the activities of public and private actors throughout the AI lifecycle, while giving parties flexibility to implement the agreed commitments in accordance with their domestic legal and institutional frameworks and relevant international human rights obligations. The Convention’s implementation framework will also provide new opportunities for engagement with all stakeholders, including States that have not yet ratified or are yet to do so.

The documents define AI systems as being able to affect “individual cognitive autonomy” and as a micro-tasking tool that directly affects people’s social behaviour, purchasing activities, etc.

The High-Level Expert Group on AI was established within the **European Commission**, which prepared several key documents that formed the basis for the European strategic guidelines in the field of AI:

- 1) AI for Europe;[26]
- 2) Trust in human-made AI;[27]
- 3) Defining AI: opportunities and scientific disciplines;
- 4) A guide to trustworthy AI ethics;
- 5) Recommendations on trustworthy AI in policy and investment.[28]

The EC identifies three key components for a trustworthy AI system throughout its life cycle:

- 1) legality (compliance with regulatory enactments);

- 2) ethics (compliance with ethical standards, principles and values);
- 3) trustworthiness (ensuring robustness and stability, while emphasizing that even with good intentions, AI systems can cause unintended harm).[29]

The **European Parliament** adopted an ethical standard for AI in autumn 2020, which will apply to “AI, robotics and related technologies developed, deployed and/or used in the EU” (regardless of the location of the software, algorithm or data). In addition to social and environmental aspects, the document sets ethical boundaries for the use of AI, highlighting the need for human oversight of AI and the need for AI certification. At the same time, the European Parliament also published a regulation on liability for the operation of AI systems, which recommends reviewing liability for product quality.[30]

The EC also adopted the EU Digital Strategy, which includes principles for drafting legal provisions for the development and deployment of AI. Its aim is to achieve EU leadership in the creation and regulation of AI technologies.[31]

The author specifically notes that on March 13, 2024, the EU adopted the AI Act, which streamlined the regulatory framework for understanding a unified concept of AI at the EU level and established uniform core principles for the operation of AI technologies within the EU. Analyzing this act, the author, in the context of the article, highlights the following key ethical principles of AI technologies defined by it: safety, reliability, transparency, inclusiveness, protection of human rights, objectivity, fairness, controllability, confidentiality, and responsibility.

In May 2019, the **Organisation for Economic Co-operation and Development (OECD)** endorsed a Recommendation on AI[32], which focuses on the concept of trust in AI systems and identifies five complementary principles for the responsible governance of AI systems, designed to ensure the following:

- 1) AI should benefit people and the planet, contributing to sustainable development and well-being;
- 2) AI systems should be designed with due regard for the rule of law, human rights, democratic values, and diversity, and should include appropriate safeguards, such as the possibility of human intervention where necessary to ensure a fair and just public order;
- 3) There should be transparent and accountable disclosure of information about AI systems to enable people to understand and challenge AI-based decisions;
- 4) AI systems should operate reliably and securely throughout their lifecycle, with potential risks always assessed and managed;
- 5) Entities and individuals that develop, deploy, or operate AI systems should be accountable for their proper functioning in accordance with the above principles.[33]

To implement these principles, the OECD makes five key recommendations to governments:

- 1) Promote public and private investment in research and development to spur innovation in safe AI;
- 2) Foster the development of accessible AI ecosystems with digital infrastructure, modern technologies, and mechanisms for sharing data and knowledge;
- 3) Ensure a policy environment that paves the way for the development and deployment of trustworthy AI systems;
- 4) Facilitate the acquisition of skills and support AI professionals to safely transition to AI careers;

5) Collaborate across countries and sectors to achieve success in the governance of trustworthy and safe AI systems.[32]

The OECD is currently developing a practical guide to implementing the above-mentioned AI Recommendations. For example, in March 2020, the OECD launched the development of the AI Policy Observatory, a platform that collects and processes information on the legal framework for AI around the world.[34]

The author also considers it important to note that most documents on the legal implementation of AI technologies are recommendations, declarations, charters, etc., which incorporate various types of "soft law" in the development and implementation of AI technologies. Thus, various approaches to the legal regulation of artificial intelligence technologies appear on the part of international organizations in Table 2

Table 2. Comparison of the main principles of artificial intelligence technologies proposed by international organizations

The Basic Principles of AI Technologies		
UNESCO	OECD	European Parliament
Proportionality to the legitimate aims or objectives and consistency with the content	Benefits for people and the planet; development and improvement of well-being	Safety and reliability
Human control and decisions	Rule of law, human rights, democratic values and diversity	Transparency
Protecting the environment and peace	Transparent and accountable disclosure of information	Inclusiveness
Gender orientation	Reliable and safe operation	Protection of human rights
	Responsibility for proper functioning	Objectivity and fairness
		Controllability and responsibility
		Confidentiality

Source: The author has created a table for comparative analysis.

Beyond international organizations, there is a pressing need to create a regulatory framework that promotes cooperation, optimizes global networks, and narrows the digital divide. To this end, in 2023, the author conducted an extensive research on the implementation of AI technologies in the public sector of Latvian state and municipal authorities. Based on the results of this research, the author identified legal principles that should guide the development and implementation of AI technologies in the public sector.

Almost two-thirds of respondents (63.6%) also positively assessed the need to create AI regulation based on some special legal principles, while 8.5% denied such a need, are shown in Figure 1.

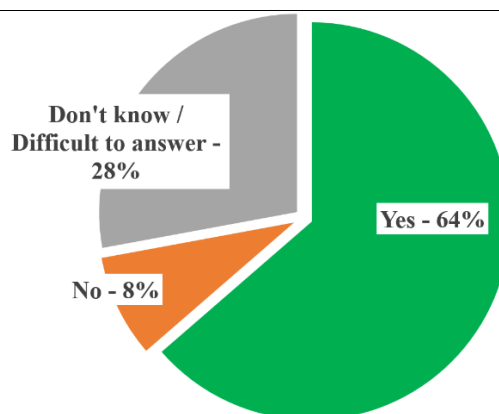


Figure 1. Respondents' opinion on whether it is necessary to create AI regulation based on any specific ethical principles

Source: based on the author's research.

In turn, of those respondents who believe that special legal principles are necessary for the legal regulation of AI technologies, security (77.8%), legality (73.9%), protection of human rights (63.1%), transparency and accountability (52.7% each) were considered the most important. Other fundamental principles are important, but did not reach the 50% mark among respondents, for example, confidentiality (46.8%), reliability (44.3%), objectivity (38.4%), fairness (35%) and manageability (29.1%), are shown in Figure 2.

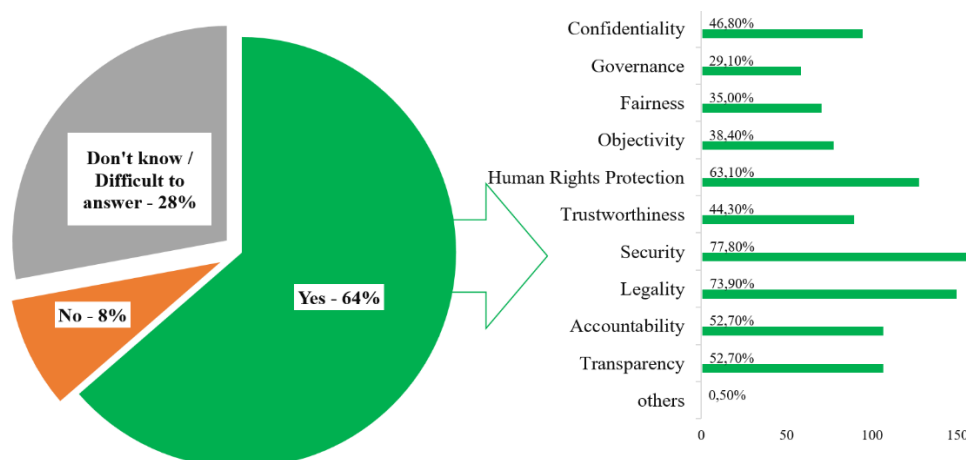


Figure 2. Respondents' opinion on the basis of which specific ethical principles it is necessary to create AI regulation

Source: based on the author's research.

One of the obstacles, but not the main one, based on the author's research, to the use of AI technologies for monitoring the activities of officials may be the lack of a sufficient degree of trust (as noted by 45.7% of respondents) on the part of the public. In particular, the increase in the degree of autonomy of AI technologies, the decrease in human control over the process of applying AI, as well as the not fully transparent decision-making logic in the operation of AI technologies may create public demand for legal restrictions on the use of AI technologies.

5. Conclusion

Analyzing the experience of international organizations in the legal regulation of AI technologies, the author emphasizes that they basically start work with the definition and implementation of ethical standards so that the AI development process is controllable and monitored, so that in the event that AI transforms into a super-powerful AI, it is always based on already established (at least three) types of ethics:

1) AI decision-making ethics – the first “entry points” are created by a person, determining certain moral principles, but in the process of self-learning, the AI system behaves autonomously, improving its parameters and making decisions that affect human life; therefore, it is necessary to guarantee the morality of AI technologies, which creates the greatest difficulties in assessing moral choices;

2) human interaction ethics and the ethics of AI software developers, the purpose of which is to analyze and prevent ethical conflicts that arise in the process of using AI, namely: violation of privacy, possible discrimination, social stratification, employment problems, etc.;

3) Professional ethics of AI system developers, which includes such principles as contribution to the development of humanity (protecting fundamental human rights, respecting cultural diversity, preventing threats to human security); compliance with regulatory enactments; respect for the privacy of other persons; fairness (objectivity, prohibition of any discrimination); security (establishing security measures for AI itself; security, manageability and necessary confidentiality, ensuring that AI users are provided with appropriate and sufficient means and information); integrity, which is based on public trust in AI technologies; accountability and social responsibility; communication and self-development.

The author believes that the principles play a very important role in the supervision and legal regulation of AI technologies and the formation of the legal regulation of AI must be based on them, thus the role of international institutions in the development and adoption of the basic legal principles of AI activities is very important. In the author's opinion, it is very important to adopt the experience of international organizations in the concept of the use of AI in preventing conflicts of interest, determining the basic principles that must be incorporated into the concept. Furthermore, it is crucial to determine that the above-mentioned basic principles are appropriate to apply in the public sector, including in matters of preventing corruption and conflicts of interest.

However, based on the author's research, AI technologies must meet the interests of society and the state. Thus, in relation to the development and application of AI technologies in the prevention of corruption in the public sector, the author believes that the following principles should be provided:

- 1) balanced stimulation of the development of AI technologies with the development of legal regulation within a certain period of time (“stimulation in parallel with regulation”);
- 2) moderate regulatory influence, based on a risk-based approach and providing for the adoption of restrictive regulations if the use of these technologies poses an objectively high risk of harming public officials subject to the control of AI technologies, as well as the interests of society and the state;
- 3) expanding the use of legal instruments, including the development of self-regulation, the creation of ethical guidelines and codes at various levels of public administration;
- 4) a humanitarian-oriented approach, which provides that the main goal of the development of AI technologies with the help of regulatory influence is to achieve a low level of corruption and avoid conflict of interest situations in the activities of public officials, thereby improving the well-being and quality of life of citizens.

In addition to the above principles, the author emphasizes that the development of AI technologies must be based on basic ethical standards and include: (1) the priority of human well-being (the

goal of ensuring human well-being must prevail over other goals of the development and use of AI systems); (2) the prohibition of causing harm at the initiative of the AI system (the development, circulation and use of an AI system that is capable of intentionally causing harm to a person on its own initiative should generally be restricted); (3) human control (to the extent possible, taking into account the required degree of autonomy of AI systems and other circumstances); (4) compliance of design with the law (the use of AI systems must not knowingly cause a violation of legal norms by the developer); (5) prevention of covert manipulation of human behavior; (6) security by design (when developing AI systems, a sufficient level of personal (i.e. state officials) and public safety must be ensured).

However, the author draws attention to the fact that the development of AI technologies poses serious challenges to the state system itself, because when solving the task of preventing corruption in the public sector, AI technologies can develop enormous autonomy of action. Therefore, achieving the goal and tasks of implementing AI in the public sector should be based on the following basic ethical norms and principles:

- 1) the principle of respect for fundamental rights: to ensure the development and application of AI-based tools and services that comply with fundamental human rights;
- 2) the principle of non-discrimination: namely, to prevent the development or intensification of discrimination between individuals or groups of persons when using AI technologies;
- 3) the principle of quality and security: to use certified sources and intangible data for decisions, judgments and data processing, using models developed in an interdisciplinary and secure technological environment;
- 4) the principle of transparency, objectivity and reliability: to make data processing methods accessible and understandable, to allow for external auditing;
- 5) the principle of user control: to abandon the approach of self-evident conclusions, but to allow the user to act as an informed participant and control his or her own choices.

Conflicts of Interest: The author declares no conflict of interest.

Data Availability Statement: This study is based on publicly available secondary sources — including national and international legal acts, conventions, institutional guidelines, and reports — all of which are cited in the references. The study also draws on primary data from a survey of public officials and employees of state and local government institutions conducted by the author; the anonymized survey data are available from the author upon reasonable request.

References

1. Doha Civil Society Declaration (2025). UNCAC CoSP11, Doha, Qatar. <https://uncaccoalition.org/uncac-cosp11-doha-civil-society-declaration/>
2. Balkin, J.M. (2015). The Path of Robotics Law. *California Law Review*. Vol. 6. pp. 45–60.
3. Demaidi, M.N. (2023). Artificial intelligence national strategy in a developing country. *AI & Soc* (2023). <https://doi.org/10.1007/s00146-023-01779-x>
4. Wagner, B. (2018). Ethics as an escape from regulation: from ethics-washing to ethics-shopping. In *being profiling: cogitas ergo sum*. Amsterdam: Amsterdam University Press. pp. 89. <https://doi.org/10.1515/9789048550180-016>
5. Galindo, L., K. Perset and F. Sheeka (2021), An overview of national AI strategies and policies, *OECD Going Digital Toolkit Notes*, No. 14, OECD Publishing, Paris, <https://doi.org/10.1787/c05140d9-en>.
6. Asilomar AI Principles. (2017). Future of Life Institute newsletter. https://ethics.cdto.center/3_8#link209

7. Pontifical Academy for Life: Yes to artificial intelligence, but with ethics (2020). Vatican News. <https://www.vaticannews.va/en/vatican-city/news/2020-05/artificial-intelligence-with-ethics-pontifical-academy-for-life.html>
8. Report Series: The Impact of Robotics and Artificial Intelligence on the Law (2021). The Singapore Academy of Law's Law Reform Committee. https://www.sal.org.sg/Resources-Tools/Law-Reform/Robotics_AI_Series
9. Nicolet-Serra, L. (2021). Singapore Academy of Law Considers the Impact of Robotics and Artificial Intelligence on the Law. The National Law Review. <https://www.natlawreview.com/article/singapore-academy-law-considers-impact-robotics-and-artificial-intelligence-law>
10. Roberts, H., Cowls, J., Morley, J., Taddeo, M., Wang, V., Floridi, L. (2021). The Chinese approach to artificial intelligence: an analysis of policy, ethics, and regulation. AI & Society, vol.36. https://www.researchgate.net/publication/342246048_The_Chinese_approach_to_artificial_intelligence_an_analysis_of_policy_ethics_and_regulation
11. Sterling, B. (2019). The Beijing Artificial Intelligence Principles. <https://www.wired.com/beyond-the-beyond/2019/06/beijing-artificial-intelligence-principles/>
12. Overly, S., Heikkilä, M. (2021). China wants to dominate AI. The U.S. and Europe need each other to tame it. <https://www.politico.com/news/2021/03/02/china-us-europe-ai-regulation-472120>
13. Combs, K., Levine, G., Blankstein, S. (2021). Podcast: Non-Binding Guidance: FDA Regulatory Developments in AI and Machine Learning. <https://www.mondaq.com/unitedstates/healthcare/1043796/podcast-non-binding-guidance-fda-regulatory-developments-in-ai-and-machine-learning>
14. William M. (Mac) Thornberry National Defense Authorization Act for Fiscal Year (2021). USA 116th Congress, H.R.6395 – 2021. <https://www.congress.gov/bill/116th-congress/house-bill/6395>
15. UNESCO Launches Worldwide Online Public Consultation on the Ethics of Artificial Intelligence (2020). UNESCO. <https://en.unesco.org/news/unesco-launches-worldwide-online-public-consultation-ethics-artificial-intelligence>
16. UNESCO Recommendation on the Ethics of Artificial Intelligence (2021). UNESCO. <https://www.ohchr.org/sites/default/files/2022-03/UNESCO.pdf>
17. World Commission on the Ethics of Scientific Knowledge and Technology (2024). UNESCO. <https://www.unesco.org/en/ethics-science-technology/comest>
18. UNESCO launches Global AI Ethics and Governance Observatory at the 2024 Global Forum on the Ethics of Artificial Intelligence (2024). UNESCO. <https://digital-skills-jobs.europa.eu/en/latest/news/unesco-launches-global-ai-ethics-and-governance-observatory-2024-global-forum-ethics>
19. Azoulay, A. (2023). Artificial Intelligence: UNESCO calls on all Governments to implement Global Ethical Framework without delay. UNESCO. <https://www.unesco.org/en/articles/artificial-intelligence-unesco-calls-all-governments-implement-global-ethical-framework-without>
20. Ethics of Artificial Intelligence (2022). UNESCO. <https://www.unesco.org/en/artificial-intelligence/recommendation-ethics>
21. Ethical frameworks (2019). Council of Europe. <https://www.coe.int/en/web/artificial-intelligence/ethical-frameworks>
22. Artificial Intelligence (2024). Council of Europe Committee on Artificial Intelligence. <https://www.coe.int/en/web/artificial-intelligence/cai>

23. Possible elements of a legal framework on artificial intelligence, based on the Council of Europe's standards on human rights, democracy and the rule of law (2021). Council of Europe, CAHAI(2021)09rev. <https://rm.coe.int/cahai-2021-09rev-elements/1680a6d90d>
24. Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (2024). Council of Europe. <https://rm.coe.int/1680afae3c>
25. CAI – Committee on Artificial Intelligence (2024). Council of Europe Committee on Artificial Intelligence. <https://rm.coe.int/terms-of-reference-of-the-committee-on-artificial-intelligence-cai-/1680ade00f>
26. Communication Artificial Intelligence for Europe (2018). European Commission, COM (2018) 237. <https://ec.europa.eu/digital-single-market/en/news/communication-artificial-intelligence-europe>
27. Building trust in human-centric AI (2025). European Commission. <https://futurium.ec.europa.eu/en>
28. Ethics guidelines for trustworthy AI (2019). European Commission. pp.10. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>
29. A definition of AI: Main capabilities and scientific disciplines (2019). European Commission, High-Level Expert Group on Artificial Intelligence. https://ec.europa.eu/newsroom/dae/document.cfm?doc_id=56341
30. Ashby, A., Syed, I., Clement-Jones, T. (2021). Regulating Artificial Intelligence: Where are we now? Where are we heading? Lexology. pp. 2. <https://www.lexology.com/library/detail.aspx?g=772fb63a-c1e2-44b9-89a4-2c042bd17841>
31. Zou, G., Zhang, H. (2021). Future of IP – China: A closer look at governmental and regulatory support of AI. ManagingIP. <https://www.managingip.com/article/b1qsl1ghs6s37w/future-of-ip-china-a-closer-look-at-governmental-and-regulatory-support-of-ai>
32. Measuring the Digital Transformation: A Roadmap for the Future (2019). ESAO. https://www.oecd-ilibrary.org/science-and-technology/measuring-the-digital-transformation_9789264311992-en
33. Council Recommendation on Artificial Intelligence (2019). ESAO, OECD/LEGAL/0449. <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>
34. AI Policy Observatory. (2025). OECD. <https://oecd.ai/en/>