
*Natalia Mushak**, *Viktor Muraviov***, *Michał Jackowski****

DOI: <https://doi.org/10.15804/rop2026105>

ARTIFICIAL INTELLIGENCE AND HUMAN RIGHTS: THE ROLE OF HUDERIA AND FRIA IN BUILDING ETHICAL FUTURES

Keywords: Council of Europe, AI governance, Fundamental Rights, FRIA, HUDERIA, Horizontal Application, Human Rights, Regulatory Compliance

ABSTRACT: In recent years, the term “artificial intelligence” (AI) has become ubiquitous, sparking wide-ranging discussions among scientists, policymakers, and practitioners across multiple fields. AI has the potential to revolutionize nearly all sectors by enhancing transparency, accountability, and decision-making. However, like any powerful technology, it can be used for both beneficial and harmful purposes. While AI creates unprecedented opportunities, it also introduces significant risks that can negatively affect individuals and society. These risks include systematic bias, digital exclusion, misuse by authoritarian actors, loss of privacy, discrimination, arbitrary detention, disinformation, election manipulation, mass surveillance, as well as concerns around human rights protection, safety, liability, security, and fairness. Given these profound implications, effective regulation of AI has become essential. Two emerging regulatory frameworks – the Fundamental Rights Impact Assessment (FRIA) under the European Union AI Act and the Human Rights Due Diligence for AI (HUDERIA) – seek to address this need by embedding fundamental rights considerations into corporate AI governance. FRIA, as a mandatory assessment mechanism, requires companies to proactively identify and mitigate AI-related human rights risks. In contrast, HUDERIA offers a voluntary framework aligned with broader ethical principles. However, the coexistence of these approaches remains fragmented, raising concerns about inconsistent corporate compliance and enforcement. The article examines how

* Department of European and Comparative Law, University of Gdansk (Poland); nataliia.mushak@ug.edu.pl.

** Institute of International Relations of Kyiv Taras Shevchenko National University (Ukraine); vikimur7@gmail.com.

*** SWPS University (Poland), visiting scholar at Cambridge University (UK); m.jackowski@dsk-kancelaria.pl

the ongoing debate over the horizontal application of human rights in AI governance is shaping corporate responsibilities. It argues that a unified AI governance framework integrating FRIA and HUDERIA could not only enhance regulatory compliance but also catalyse a strategic shift in corporate approaches to fundamental rights.

INTRODUCTION

The use of Artificial Intelligence (hereinafter – AI) platforms in the majority of European countries may be followed by many risks, which may create negative impact on everyday life of people. Those risks include systematic bias, digital exclusion, misuse of authoritarian actors, privacy loss, discrimination, arbitrary detention, disinformation, survey the population, attempts to manipulate elections to name a few. Many of these risks may have an objective ground, which reflect historic inequalities in societies, lack of internet access to or benefits from digital technologies, people speaking minority languages, lack of digital literacy of access to devices, etc.

The Council of Europe and its bodies have become the first international organisations, which has taken responsibilities for tackling the problem of AI risks management on international level. The European Commission for the Efficiency of Justice (hereinafter – CEPEJ) of the Council of Europe adopted in December 2018 the first European Ethical Charter on the use of artificial intelligence in judicial systems and their environment (hereinafter – Charter) (European ethical Charter on the use of Artificial Intelligence..., 2018). The Charter provides a framework of principles that can guide policy makers, legislators and justice professionals when they grapple with the rapid development of AI in national judicial processes (European ethical Charter on the use of Artificial Intelligence..., 2018).

The CEPEJ's view as set out in the Charter is that the application of AI in the field of justice can contribute to improve the efficiency and quality and must be implemented in a responsible manner which complies with the fundamental rights guaranteed in particular in the European Convention on Human Rights (hereinafter – ECHR) and the Council of Europe Convention on the Protection of Personal Data. For the CEPEJ, it

is essential to ensure that AI remains a tool in the service of the general interest and that its use respects individual human rights.

The CEPEJ has identified the following core principles to be respected in the field of AI and justice: a) principle of respect of fundamental rights: ensuring that the design and implementation of artificial intelligence tools and services are compatible with fundamental rights; b) principle of non-discrimination: specifically preventing the development or intensification of any discrimination between individuals or groups of individuals; c) principle of quality and security: with regard to the processing of judicial decisions and data, using certified sources and intangible data with models conceived in a multi-disciplinary manner, in a secure technological environment; d) principle of transparency, impartiality and fairness: making data processing methods accessible and understandable, authorising external audits; e) principle “under user control”: precluding a prescriptive approach and ensuring that users are informed actors and in control of their choices (European ethical Charter on the use of Artificial Intelligence..., 2018).

THE COUNCIL OF EUROPE COMMITTEE ON ARTIFICIAL INTELLIGENCE (CAI) AND FRAMEWORK CONVENTION ON AI

Another body of the Council of Europe, which deals with the AI risks management is the Committee on Artificial Intelligence (hereinafter – CAI). It was set up by the Committee of Ministers under Article 17 of the Statute of the Council of Europe (Statute of the Council of Europe, 1949). Article 17 stipulates that the Committee of Ministers may set up advisory and technical committees for specific purposes as it may deem desirable. Moreover, the establishment of CAI complies with the provisions of the Resolution CM/Res (2021)3 on intergovernmental committees and subordinate bodies, their terms of reference and working methods (Resolution CM/Res(2021)3).

The CAI’s activities on AI governance and ethical standards correlates with the concept of a Framework Convention (Framework Convention

Concept, 2011) as a possible legal and policy vehicle to establish broad, cooperative, and adaptable AI governance principles internationally. At the same time, CAI encourages the idea of international cooperation and governance frameworks for AI.

In 2024 the Council of Europe adopted the Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (hereinafter – Framework Convention) (Explanatory Report to the Council of Europe Framework Convention..., 2024), which became the first-ever international legally binding treaty in the area of AI (Explanatory Report to the Council of Europe Framework Convention..., 2024). The Framework Convention complements existing international standards on human rights, democracy and the rule of law, and aims to fill any legal gaps that may result from rapid technological advances. It aims to ensure that activities within the lifecycle of artificial intelligence systems are fully consistent with human rights, democracy and the rule of law, while being conducive to technological progress and innovation.

The Preamble sets the foundation for the Council of Europe's AI Convention, emphasizing the need to develop a globally applicable legal framework for artificial intelligence that upholds human rights, democracy, and the rule of law. It acknowledges both the positive potential of AI – such as fostering innovation, gender equality, and sustainable development – and the risks, including discrimination, surveillance, and threats to individual dignity and autonomy.

The provisions of the document encourage the participants for international cooperation taking into consideration digital literacy, trust, and a flexible framework that can evolve with technology. The legal basis for the Framework Convention is a set of international legal and policy instruments on artificial intelligence. They include: Declaration of the Committee of Ministers of the Council of Europe on the manipulative capabilities of algorithmic processes, adopted on 13 February 2019 – Declaration (13/02/2019)¹; Recommendation on Artificial Intelligence adopted by the OECD Council on 22 May 2019 (the “OECD AI Principles”); Recommendation of the Committee of Ministers of the Council of Europe to member States on the human rights impacts of algorithmic systems, adopted on 8 April 2020 – CM/Rec(2020)¹; Resolutions and

Recommendations of the Parliamentary Assembly of the Council of Europe, examining the opportunities and risks of artificial intelligence for human rights, democracy, and the rule of law and endorsing a set of core ethical principles that should be applied to AI systems; UNESCO Recommendation on the Ethics of Artificial Intelligence adopted on 23 November 2021; G7 Hiroshima Process International Guiding Principles for Organisations Developing Advanced AI Systems and Hiroshima Process International Code of Conduct for Organizations Developing Advanced AI Systems (adopted on 30 October 2023); and EU Regulation laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act) adopted on 13 March 2024.

In our opinion, the aspect of “unified governance” has a comprehensive application under the Article 7 “Human dignity and individual autonomy” of the Framework Convention. It highlights the significance of human dignity and individual autonomy as part of human-centric regulation and governance of the activities within the lifecycle of artificial intelligence systems. In other words, it means the avoidance of dehumanizing individuals, undermining their autonomy, or reducing them to mere data points, as well as to prevent the anthropomorphizing of artificial intelligence systems in ways that compromise human dignity.

At the same time within the lifecycle of artificial intelligence systems the fundamental principle of transparency, as set forth in Article 8, calls for the obligation to ensure openness and clarity in the governance of activities. It requires that the decision-making processes and general functioning of such systems should be understandable to the relevant artificial intelligence actors, and, where appropriate and necessary, to other stakeholders concerned, in a manner consistent with applicable legal standards and principles. In certain circumstances, this may also involve the provision of additional information, involving, for instance, the algorithms used, subjects relating to security, commercial confidentiality, intellectual property rights etc. The measures by which transparency is to be ensured shall depend on various factors, such as the type of artificial intelligence system, the context and purpose of its use, its function, and the background of the relevant actor or affected stakeholder.

Moreover, relevant measures may include, as appropriate, recording key considerations such as data provenance, training methodologies, validity of data sources, documentation and transparency on training, testing and validation data used, risk mitigation efforts, and processes and decisions implemented, in order to aid a comprehensive understanding of how the artificial intelligence system's outputs are derived and impact human rights, democracy and the rule of law. This will, in particular, help to ensure accountability and enable persons concerned to contest the use or outcomes of artificial intelligence systems (Art. 8 "Transparency and oversight") (Explanatory Report to the Council of Europe Framework..., 2024). Among the relevant measures could be listed different training methods, validity of data sources, risk mitigation efforts to understand how a certain artificial intelligence system generates its outputs and affects human rights, the rule of law, and democracy. In such a way the above-mentioned measures will facilitate accountability and enable individuals to challenge the use or outcomes of these systems.

According to Article 16 "Risk and impact management framework" the provisions of the document involve the relevant risks and potential impacts to human rights, democracy and the rule of law by the development of a methodology. The Framework Convention requires identifying, assessing, preventing, and mitigating risks to human rights, democracy, and the rule of law throughout the AI system's lifecycle. This should be done using a clear, objective methodology. These obligations help implement the Convention's core requirements, especially in Chapters II and III, and must reflect key principles such as transparency, oversight, accountability, and responsibility.

This provision intends to establish a uniform approach to distinguishing, analysing, and assessing the risks and impacts of such systems. Taking into consideration legal, economic, political and technological contexts the participants are being recognized as the best actors in making regulatory decisions and having flexibility in governing these processes.

Parties to the Framework Convention retain the flexibility to use or adapt the document, in whole or in part, to develop new approaches to risk assessment or to modify existing ones, in accordance with their national legal frameworks. This flexibility is permissible if Parties to the Framework

Convention fully comply with their obligations under the Framework Convention, particularly the baseline requirements for risk and impact management.

HUMAN RIGHTS IMPACT ASSESSMENT IN AI GOVERNANCE: THE ROLE OF HUDERIA

The Council of Europe's Ad Hoc Committee on AI (hereinafter – CAHAI) also shares a concern over how best to govern AI systems. The CAHAI's feasibility study (hereinafter – “the study”) concluded that measures must be put in place to radically reorient technology companies toward embedding human rights, democracy, and rule of law principles in their products and activities. These include the development of a “uniform model” for human rights impact assessments (HRIAs) of AI systems, which could form one tool in the toolbox for AI developers to meet their human rights due diligence obligations. The CAHAI has called this tool a Human Rights, Democracy, and Rule of Law Impact Assessment (hereinafter – HUDERIA).

In 2024 the CAHAI launched a Rights-Based Framework for Responsible AI Governance – a special aid, which provides guidance and a structured approach to carry out risk and impact assessments for AI systems throughout their lifecycle (Methodology for the risk..., 2024). It offers a cohesive end-to-end process for identifying contexts and applications in which the deployment of AI systems would likely pose significant levels of risk to human rights, the functioning of democracy, and the observance of the rule of law (Human Rights, Democracy, and the Rule of Law Impact Assessment for AI Systems..., 2021).

HUDERIA reflects an effort to bridge technical expertise with human rights considerations, offering both public and private entities a structured approach to AI risk and impact management (Werner, 2024).

The HUDERIA methodology is a stand-alone, non-legally binding guidance document using an algorithm-neutral and practice-based approach. It can be used by both public and private actors for identifying and addressing risks and impacts to human rights, democracy and the rule of law (Council of Europe Framework Convention..., 2024). At the same

time it is a flexible instrument for the Parties of the Framework Convention, to use or adapt it, in whole or in part, to develop new approaches to risk assessment or to refine existing ones, in accordance with their applicable laws (HUDERIA..., 2024).

The Framework Convention complements existing international standards on human rights, democracy and the rule of law, and aims to fill any legal gaps that may result from rapid technological advances. It aims to ensure that activities within the lifecycle of artificial intelligence systems are fully consistent with human rights, democracy and the rule of law, while being conducive to technological progress and innovation.

The HUDERIA methodology is to be complemented by the HUDERIA model to be adopted in 2025, which will provide supporting materials and resources, including flexible tools and scalable recommendations (HUDERIA..., 2024). The HUDERIA methodology offers the combination of general and specific guidance and flexibility by allowing room for adaptation in the practical implementation. At the general level, the HUDERIA describes high-level concepts, processes and elements guiding risk and impact assessment activities of AI systems that could have impacts on human rights, democracy and the rule of law.

The HUDERIA Model will be used at the specific level and will provide supporting materials and resources (such as flexible tools relevant for different elements of the HUDERIA process and scalable recommendations) that can aid in the implementation of the HUDERIA Methodology. These resources are referred to throughout the text and will provide a library of knowledge that can facilitate consideration of risks and impacts related to human rights, democracy, and the rule of law, including in other approaches to risk management. HUDERIA is a multi-step governance process offering an anticipatory approach to the governance of AI design.

Moreover, HUDERIA consists of four interconnected key phases that include Context-Based Risk Analysis (COBRA), Stakeholder Engagement Process (SEP), Risk and Impact Assessment (RIA) and Mitigation Plan (MP). Each of these phases is undergirded by the ambient Project Summary Report (PS Report), which provides the key source of documentation required for continuous accountability and deliberation (Human Rights, Democracy, and the Rule of Law..., 2021).

The COBRA provides a preliminary indication of the risks that AI systems could pose throughout their lifecycle through a risk arrangement mechanism. It includes three significant steps. The first step identifies key risk factors – specific characteristics within an AI system’s lifecycle context – that heighten the likelihood of adverse impacts on human rights, democracy, and the rule of law. These risk factors are grouped into three categories: the system’s application context, design and development context, and deployment context.

The second step is the analysis of factors as well as the checking of potential adverse impacts on human rights, democracy and the rule of law. The third one means the building on this information, the triage process focuses on identifying and prioritizing systems with significant risks, ensuring that the HUDERIA Methodology remains proportional and not overly burdensome for minimal or low-risk AI systems. This process also supports informed decision-making on whether the benefits of building or deploying an AI system outweigh its risks, particularly regarding potential impacts on human rights, democracy, and the rule of law.

The SEP enhances the Risk and Impact Assessment by incorporating the perspectives of potentially affected individuals identified during the COBRA stage. Tailored to the identified risk factors and potential impacts, stakeholder engagement can take various forms and levels of participation. This process not only improves the quality of risk analysis but also fosters transparency, builds trust, and enhances the usability and performance of the AI system.

The RIA provides a detailed evaluation of the potential and actual impacts of AI system activities on human rights, democracy, and the rule of law, focusing particularly on systems posing significant risks identified during the COBRA triage. It involves re-examining, contextualizing, and expanding upon potential harms, while assessing key risk variables such as scale, scope, reversibility, and likelihood to prioritize and manage risks effectively. Building on the COBRA analysis and potential SEP insights, this step ensures a comprehensive understanding of the risks to inform mitigation and governance strategies.

The MP element of the HUDERIA process outlines actions and strategies to address adverse impacts and mitigate identified harms. It involves

formulating targeted measures based on the severity and likelihood of these harms and developing a comprehensive plan to implement them. Where appropriate, it also includes establishing mechanisms for affected individuals and other stakeholders to access remedies (HUDERIA..., 2022).

HUDERIA prompts regular reassessments through the Iterative Revisitation (IR) phase, to ensure that AI systems continue to operate safely and ethically as the context and technology evolve and public is protected from emerging risks throughout the system's life cycle (Human Rights, Democracy, and the Rule of Law..., 2021). The HUDERIA Methodology and the HUDERIA Model are supposed to prevent and mitigate risks in the application of AI in corporate governance by ensuring that AI tools are used ethically, responsibly, and in compliance with governance principles.

There are several categories of risks that are proper for the application of AI in corporate governance. They include ethical and legal concerns over handling sensitive company and stakeholder data; the risks to inherit or amplify existing biases; complexity in allotment of responsibility for decisions made or suggested by AI; overreliance on technology and inability by AI to replace human judgment in governance decisions.

The governance and effective enforcement of existing law on fundamental rights and safety requirements applicable to AI systems is being one of the cornerstones of the EU Artificial Intelligence Act. The need for democratic governance of artificial intelligence was implemented both by the Parliamentary Assembly Recommendation 2181 (2020) on 22 October 2020 (Recommendation 2181, 2020) and Parliamentary Assembly Resolution 2341 (2020) with the purpose to establish a regulatory legal framework for the design, development and application of artificial intelligence (AI) at European and international levels taking into account the Council of Europe's standards on human rights, democracy and the rule of law. The necessary procedures for the assessment, designation and notification of conformity assessment bodies are being regulated under the Article 30 "Notifying authorities". According to Mauritz Kop it is crucial that certification bodies as well as notified bodies operate with full independence, free from any financial or political conflicts of interest. Furthermore, the degree to which companies are allowed to demonstrate

compliance with the new AI “product safety regime” through risk-based self-assessment and self-certification -without oversight from third-party notified bodies – will significantly influence how the Regulation impacts business practices and, consequently, the protection and promotion of our core societal values. Internal self-assessments, particularly for high-risk systems, are insufficient given the potential risks involved (Kop, 2021).

Self-regulation, even when accompanied by internal AI impact assessments – whether voluntary or mandatory – is not adequate. Companies inherently operate under different incentives than the state and are not primarily driven by the pursuit of social welfare. To truly safeguard public interest and uphold democratic values, mandatory third-party audits for all high-risk AI systems are necessary (Kop, 2021).

FROM VOLUNTARY ETHICS TO BINDING LAW: THE ROLE OF FRIA IN EU AI REGULATION

Unlike HUDERIA which is a voluntary framework aligned with broader ethical principles, FRIA, as a mandatory assessment mechanism, requires companies to proactively identify and mitigate AI-related human rights risks. Its use is foreseen by the EU AI Act (Regulation (EU) 2024/1689). For instance, the document determines the conditions under which an AI system shall be considered to be high-risk and fall under the obligations, including FRIA. It also provides that AI systems referred to in Annex III shall be deemed as high-risk (Regulation (EU) 2024/1689).

EU AI Act identifies high-risk AI systems as the AI systems listed in any of the following areas: biometrics, in so far as their use is permitted under relevant Union or national law; critical infrastructure; education and vocational training; employment, workers’ management and access to self-employment; access to and enjoyment of essential private services and essential public services and benefits; law enforcement, in so far as their use is permitted under relevant Union or national law; migration, asylum and border control management, in so far as their use is permitted under relevant Union or national law; administration of justice and democratic processes.

Furthermore, Art. 27 of the EU AI Act stipulates that high-risk AI systems shall be designed and developed in such a way as to ensure that their operation is sufficiently transparent to enable deployers to interpret a system's output and use it appropriately. It means that an appropriate type and degree of transparency shall be ensured with a view to achieving compliance with the relevant obligations of the provider and deployer. Also, the document provides that deployers shall perform an assessment of the impact on fundamental rights that the use of such system may produce. For that purpose, deployers shall perform an assessment consisting of: a description of the deployer's processes in which the high-risk AI system will be used in line with its intended purpose; a description of the period of time within which, and the frequency with which, each high-risk AI system is intended to be used; the categories of natural persons and groups likely to be affected by its use in the specific context; the specific risks of harm likely to have an impact on the categories of natural persons or groups of persons identified pursuant to point (c) of this paragraph, taking into account the information given by the provider pursuant to Article 13; a description of the implementation of human oversight measures, according to the instructions for use; the measures to be taken in the case of the materialization of those risks, including the arrangements for internal governance and complaint mechanisms (Regulation (EU) 2024/1689).

The AI Office shall develop a template for a questionnaire, including through an automated tool, to facilitate deployers in complying with their obligations under this Article in a simplified manner.¹ The impact assessment should identify the deployer's relevant processes in which the high-risk AI system will be used in line with its intended purpose, and should include a description of the period of time and frequency in which the system is intended to be used as well as of specific categories of natural persons and groups who are likely to be affected in the specific context of use.

¹ Note: Art.3.AI Office' means the Commission's function of contributing to the implementation, monitoring and supervision of AI systems and general-purpose AI models, and AI governance, provided for in Commission Decision of 24 January 2024. Act.

The assessment should also include the identification of specific risks of harm likely to have an impact on the fundamental rights of those persons or groups. Deployers should determine measures to be taken in the case of a materialization of those risks, including, for example, governance arrangements in that specific context of use, such as arrangements for human oversight according to the instructions of use. After performing that impact assessment, the deployer should notify the relevant market surveillance authority.

In general, FRIA foresees the interaction with other assessments such as the General Data Protection Regulation (GDPR) (Regulation (EU) 2016/679), that came into effect on 25 May 2018, and sets rules for how organizations can collect, use, store, and share personal data. GDPR protects individuals' personal data and privacy across the EU. It also applies to companies outside the EU if they handle EU residents' data. If personal data is processed the GDPR is to be applied. The data subjects' rights include: right to access to personal data, which is any information that can directly or indirectly identify a person (e.g., name, email, IP address, biometric data, location data); right to rectification that is right to fix errors; right to erasure which is so called "right to be forgotten"; right to restrict processing; right to data portability that is right to move data to another provider; right to object that is right stop certain processing; rights around automated decision-making and profiling.

The description for each right identified should contain potential impacts (direct or indirect); likelihood and severity of impact; affected groups (e.g., workers, students, migrants, consumers); cumulative or systemic risks (bias, surveillance, exclusion, chilling effects).

In comparison with the HUDERIA, which explicitly requires a mitigation plan document, the mitigation measures under FRIA are integrated into the risk management system and impact assessment (Regulation (EU) 2016/679). The mitigation measures, designed to reduce the likelihood, severity, or scope of harm caused by an AI system to human rights, democracy, or rule of law may include: technical measures (bias testing, explainability, human oversight); organizational measures including training, governance, accountability structures); procedural safeguards (complaint channel for applicants that is the appeal mechanisms, notice to

affected persons); residual risks (what cannot be fully mitigated) (Convention 108+..., 2018). The enforcement of the EU AI Act by member-states is supervised by national Data Protection Authorities (e.g. CNIL in France, ICO in UK pre-Brexit, AEPD in Spain, etc.).

CONCLUSION

The integration of artificial intelligence into European societies and judicial systems brings not only opportunities for innovation and efficiency but also significant risks that must be carefully managed. International governance initiatives – particularly those led by Council of Europe – play a pivotal role in ensuring that AI development and deployment align with human rights, democracy, and the rule of law. Through legal instruments such as the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, AI governance frameworks are evolving to promote transparency, accountability, and human dignity, providing both a legal and ethical foundation for responsible technological progress.

A unified and flexible governance approach is essential to mitigate AI-related risks effectively. Tools such as the Council of Europe Ad Hoc Committee on Artificial Intelligence (CAHAI)'s HUDERIA methodology provide structured, rights-based mechanisms for risk and impact assessment, enabling states and private actors to anticipate and address potential harms. By embedding principles of transparency, non-discrimination, and fundamental rights protection into the entire AI lifecycle, Europe sets an international precedent for balancing technological advancement with the preservation of democratic values and individual freedoms.

Different stakeholders shall be aware of AI application and usage. The way one draws modern technologies directly shapes the future of our community, making it imperative that democratic values, human rights and fundamental freedoms are embedded at every stage. In order to maintain this awareness, essential instruments include AI impact and conformity assessments, best practices, technology roadmaps, and codes of conduct. All these mechanisms

must be implemented by inclusive multidisciplinary teams, who use them to monitor, validate, and benchmark AI systems effectively.

BIBLIOGRAPHY:

- A risk-based approach to assessing and mitigating adverse impacts developed for the Council of Europe's Framework Convention. (2021). Retrieved from: <https://www.turing.ac.uk/research/research-projects/human-rights-democracy-and-rule-law-impact-assessment-ai-systems-huderia>.
- Convention 108+ Convention for the protection of individuals with regard to the processing of personal data. (2018). Retrieved from: https://www.europarl.europa.eu/meetdocs/2014_2019/plmrep/COMMITTEES/LIBE/DV/2018/09-10/Convention_108_EN.pdf.
- Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law Vilnius. (2024). Retrieved from: <https://rm.coe.int/1680afae67>.
- Declaration by the Committee of Ministers on the manipulative capabilities of algorithmic processes. (2019). The Committee of Ministers on 13 February 2019 at the 1337th meeting of the Ministers' Deputies. Retrieved from: <https://rm.coe.int/090000168092dd4b>.
- European Ethical Charter on the use of artificial intelligence (AI) in judicial systems and their environment. (2018). Retrieved from: <https://www.coe.int/en/web/cepej/cepej-european-ethical-charter-on-the-use-of-artificial-intelligence-ai-in-judicial-systems-and-their-environment>.
- Explanatory Report to the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law Vilnius. (2024). Retrieved from: <https://rm.coe.int/1680afae67>.
- Framework Convention Concept. (2011). Retrieved from: <https://unece.org/fileadmin/DAM/hlm/sessions/docs2011/informal.notice.5.pdf>.
- G7 Hiroshima Process International Guiding Principles for Organisations Developing Advanced AI Systems and Hiroshima Process International Code of Conduct for Organizations Developing Advanced AI Systems. (2023). Retrieved from: <https://digital-strategy.ec.europa.eu/en/library/hiroshima-process-international-code-conduct-advanced-ai-systems>.
- HUDERIA – risk and impact assessment of AI systems. (2022). Retrieved from: <https://www.coe.int/en/web/artificial-intelligence/huderia-risk-and-impact-assessment-of-ai-systems>.

- HUDERIA: New tool to assess the impact of AI systems on human rights. (2024). Retrieved from: <https://www.coe.int/en/web/portal/-/huderia-new-tool-to-assess-the-impact-of-ai-systems-on-human-rights>.
- Human Rights, Democracy, and the Rule of Law Impact Assessment for AI Systems (HUDERIA) A risk-based approach to assessing and mitigating adverse impacts developed for the Council of Europe's Framework Convention. (2021). Retrieved from: <https://rm.coe.int/cai-2024-16rev2-methodology-for-the-risk-and-impact-assessment-of-arti/1680b2a09f>.
- Kop, M. (2021). EU Artificial Intelligence Act: The European Approach to AI. Retrieved from: https://futurium.ec.europa.eu/sites/default/files/2021-10/Kop_EU%20Artificial%20Intelligence%20Act%20-%20The%20European%20Approach%20to%20AI_21092021_0.pdf.
- Methodology for the risk and impact assessment of artificial intelligence systems from the point of view of human rights, democracy and the rule of law (HUDERIA Methodology). (2024). Retrieved from: <https://rm.coe.int/cai-2024-16rev2-methodology-for-the-risk-and-impact-assessment-of-arti/1680b2a09f>.
- Parliamentary Assembly Resolution 2341. (2020). Retrieved from: <https://pace.coe.int/en/files/28803/html>.
- Recommendation 2181. (2020). Need for democratic governance of artificial intelligence, *adopted by the Standing Committee*, acting on behalf of the Assembly, on 22 October 2020. Retrieved from: <https://pace.coe.int/en/files/28804/html>.
- Recommendation of the Committee of Ministers of the Council of Europe to member States on the human rights impacts of algorithmic systems, adopted on 8 April 2020 – CM/Rec(2020)1. Retrieved from: <https://policehumanrightsresources.org/recommendation-cm-rec20201-of-the-committee-of-ministers-to-member-states-on-the-human-rights-impacts-of-algorithmic-systems>.
- Recommendation on Artificial Intelligence adopted by the OECD Council on 22 May 2019 (the “OECD AI Principles”). (2019). Retrieved from: <https://oecd.ai/en/ai-principles>.
- Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation). Retrieved from: <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>.
- 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). Retrieved from: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>.

Resolution CM/Res(2021)3 on intergovernmental committees and subordinate bodies, their terms of reference and working methods (Adopted by the Committee of Ministers on 12 May 2021 at the 1404th meeting of the Ministers' Deputies). Retrieved from: <https://rm.coe.int/0900001680a27292#:~:text=12%20May%202021%20Resolution%20CM%2FRes%282021%293%20on%20intergovernmental%20committees,at%20th%201404th%20meeting%20of%20th%20Ministers%27%20Deputies%29>.

Statute of the Council of Europe (1949). Retrieved from: <https://rm.coe.int/1680306052>.

UNESCO Recommendation on the Ethics of Artificial Intelligence (2021). Retrieved from: <https://www.unesco.org/en/articles/recommendation-ethics-artificial-intelligence>.

Werner J. (2024). Council of Europe Introduces HUDERIA for AI Risk and Impact Assessments. Retrieved from: <https://babl.ai/council-of-europe-introduces-huderia-for-ai-risk-and-impact-assessments/>.