



Artificial Intelligence and Liability for Damages

Assoc. Prof. Dr. Fikret ERKAN

fikreterkan@hotmail.com, ORCID: 0009-0004-9462-7572

Abstract

This study explores the evolving legal frameworks surrounding artificial intelligence (AI) liability at both international and national levels. As AI technologies increasingly operate with autonomy and unpredictability, traditional legal doctrines—such as fault-based liability and strict liability—face growing challenges. The article provides a comparative analysis of recent case law from China, Switzerland, Ireland, France, Germany, the Czech Republic, and the European Union, highlighting diverse judicial responses to AI-generated harms, including copyright infringement, product liability, and algorithmic decision-making. In particular, it examines landmark developments such as the EU Artificial Intelligence Act (Regulation 2024/1689) and the Council of Europe's 2024 Framework Convention on AI and Human Rights. The study emphasizes the necessity of strengthening judicial oversight and national legislation in areas like criminal accountability, administrative transparency, and compensation mechanisms. Ultimately, the article offers a normative framework for establishing a more responsible, rights-based AI governance system.

Keywords: Artificial Intelligence, Legal Liability, AI Act, Copyright, Product Liability, Algorithmic Accountability, Human Rights, EU Law, International Regulation, Judicial Oversight

I-INTRODUCTION

Artificial intelligence (AI), one of the most significant technological revolutions of the 21st century, has brought about profound transformations in a wide range of fields, from daily life and public services to commerce and justice systems. In particular, the increasing role of AI systems in decision-making processes in recent years necessitates a thorough evaluation of not only the technical but also the ethical and legal dimensions of this technology. The potential for autonomous AI systems to produce erroneous or unforeseen outcomes poses serious risks of harm to individuals and institutions, making it essential to redefine existing liability regimes.

While current legal systems are shaped around classical perpetrator-victim relationships and fault-based liability principles, the autonomous and learning capabilities of AI challenge the application

of traditional liability doctrines. Consequently, the question of “who is liable” for damages caused by AI systems has evolved into not only a technical but also a normative issue.

On the international level, the European Union’s Artificial Intelligence Act (Regulation EU 2024/1689) and the Council of Europe’s Framework Convention on Artificial Intelligence, Human Rights, Democracy, and the Rule of Law represent the first comprehensive regulatory initiatives in this domain. At the same time, the United Nations and other international organizations have developed guiding documents addressing data protection, human rights, and ethical responsibility principles related to AI applications.

In this context, the central research question of the study is as follows:
What kind of legal liability regime should be constructed in response to damages caused by AI systems, and how can the boundaries of this regime be defined at national and international levels?

To address this research question, the following methodological approach will be adopted:

- A comparative analysis of the existing legal framework for AI, with reference to international documents (particularly from the EU, the Council of Europe, and the United Nations);
- A reconsideration of the elements of liability (fault, damage, causal link) in the context of AI, and a critical examination of fault-based and strict liability regimes;
- An evaluation of liability scenarios in administrative and criminal law through case studies;
- Finally, the study will propose normative models and policy recommendations to address existing legal gaps.

Accordingly, the study aims to offer a multidimensional assessment not only from the perspective of positive law, but also in light of fundamental rights, the rule of law, and democratic accountability principles.

II- ARTIFICIAL INTELLIGENCE

1.1. The Concept, Tools, and Classification of Artificial Intelligence

1.1.1. The Concept of Artificial Intelligence

Artificial Intelligence (AI) refers to the ability of computers to demonstrate human-like intelligence. In other words, it enables machines to imitate human capabilities such as learning, problem-solving, decision-making, language understanding, and even creative thinking.

Definitions of AI in international legal instruments provide a normative framework to explain the functions and impacts of AI systems. In this regard, both the Council of Europe's Framework Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law¹ and the EU Artificial Intelligence Act (Regulation 2024/1689)² offer legal definitions of AI systems.

According to Article 2 of **the Council of Europe Convention**, an AI system is defined as: "...a machine-based system that, for explicit or implicit objectives, infers how to generate outputs such as predictions, content, recommendations, or decisions influencing physical or virtual environments, based on the input it receives."

Article 3(1) of **the European Union's Regulation 2024/1689** provides the following definition: "An AI system means a machine-based system designed to operate with varying levels of autonomy and that may adapt after deployment. It infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments, for explicit or implicit objectives."

The common ground in both definitions is the recognition of AI as a machine-based system. These definitions emphasize that AI systems produce various outputs based on the inputs they receive, and that such outputs can influence both physical and digital environments. Furthermore, both definitions indicate that AI systems may operate for explicit or implicit objectives. This suggests the possibility that decision-making processes may evolve beyond human oversight. The EU's

¹ The Convention was adopted by the Committee of Ministers of the Council of Europe on 17 May 2024. It was opened for signature on 5 September 2024 in Vilnius, Lithuania, during the Council of Europe Conference of Ministers of Justice. Council of Europe. (2024). *Framework Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law*. Strasbourg. Retrieved from <https://www.coe.int/en/web/artificial-intelligence/convention>

² On 12 July 2024, Regulation (EU) 2024/1689 on Artificial Intelligence was published in the Official Journal of the European Union, becoming the first comprehensive horizontal legal framework for the regulation of AI systems across the EU. The Regulation will enter into force on 1 August 2024 in all 27 EU Member States, and most of its provisions will become applicable starting from 2 August 2026.

For the full text of the Regulation, see:

https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=OJ:L_202401689 (Accessed 01/07/2025)

regulation additionally highlights the system's level of autonomy and its capacity for post-deployment adaptation, thus offering a more technical and functional perspective.

In both definitions, AI systems are understood not merely as technical tools but as autonomous entities capable of influencing human actions. These definitions form a shared foundation for determining the boundaries of legal liability related to AI. The EU's more detailed and technically oriented definition stands out in terms of regulatory clarity, whereas the Council of Europe's definition offers a more abstract and principled framework from a fundamental rights perspective.

The fundamental purpose of AI is to imitate human intelligence in order to produce more efficient, faster, and more accurate solutions.

1.1.2. Categories of Artificial Intelligence

AI is generally divided into three main categories:

- **Artificial Narrow Intelligence (ANI):** Designed to perform specific tasks such as pattern recognition or prediction³. These systems can match human-level performance in narrowly defined domains—e.g., voice assistants or recommendation algorithms—but lack general intelligence.
- **Artificial General Intelligence (AGI):** Refers to the potential of machines to exhibit intelligence comparable to that of humans, including the capacity to learn and apply knowledge across a broad range of tasks. AGI is still under development and remains a subject of ongoing research.
- **Generative AI (GenAI):** A set of algorithms capable of creating seemingly new and realistic outputs such as text, images, and audio from training data. These systems employ deep learning techniques and large datasets to generate, summarize, or predict content⁴.

AI also encompasses subfields such as Machine Learning (ML) and Deep Learning (DL). ML enables machines to learn from experience, while DL utilizes larger datasets and complex neural network structures for deeper learning processes.

³ University of Notre Dame. (n.d.). *AI overview and definitions*. Notre Dame Learning. Retrieved June 27, 2025, from <https://learning.nd.edu/resource-library/ai-overview-and-definitions/>

⁴ Notre Dame, *AI Overview and Definitions*.

While AI is revolutionizing sectors like healthcare, finance, automotive, and education, it continues to raise ethical, safety, and labor-related concerns.

1.2. Artificial Intelligence Tools

Artificial Intelligence (AI) tools are software and systems developed to perform a wide range of tasks. These tools are built upon technologies that enable data analysis, learning, prediction, natural language processing, and more. The most commonly used AI tools include:

1. **Machine Learning (ML) Tools:** These tools enable machines to learn from data. Examples include TensorFlow, Keras, and PyTorch, which are widely used to develop machine learning and deep learning models.
2. **Natural Language Processing (NLP) Tools:** NLP tools help machines understand and process human language. Tools such as GPT-4, BERT, and SpaCy are commonly used in text analysis, machine translation, text classification, and chatbot development.
3. **Computer Vision Tools:** These are used to analyze and interpret visual data. Tools like OpenCV, YOLO (You Only Look Once), and TensorFlow are utilized for object recognition, facial recognition, and visual analysis.
4. **Data Analysis and Processing Tools:** Data analysis is a key domain within AI, focusing on the processing and interpretation of large datasets. Tools such as Pandas, NumPy, and Scikit-learn are used to analyze, clean, and transform data into meaningful insights.
5. **Automation Tools:** These tools automate routine and time-consuming processes. RPA platforms like UiPath, Automation Anywhere, and Blue Prism are widely adopted for automating business workflows.
6. **Chatbots and Virtual Assistants:** These AI systems interact with users and answer their queries. Tools such as Dialogflow, Rasa, and Microsoft Bot Framework are employed in developing intelligent virtual assistants.
7. **Predictive and Analytical Tools:** In business and scientific research, these tools are used to forecast future trends or outcomes. Examples include IBM Watson, SAS Analytics, and Google Cloud AI, which provide predictive insights based on large-scale data analysis.

8. **AI Creativity Tools:** Some AI tools are designed for creative tasks, such as generating art, music, or text. Tools like DALL·E, DeepArt, and Runway ML are used in visual arts and content production.

These tools enhance efficiency across sectors, offer creative solutions, and allow human labor to be allocated to more strategic roles. As a constantly evolving field, AI continues to produce new tools and applications at an accelerating pace.

1.3. Classification of Artificial Intelligence

AI systems can be classified according to various criteria, including their level of competence, application areas, and operational mechanisms. The main types of classification are outlined below:

1.3.1. Classification Based on Competence Levels

Artificial intelligence can be categorized into three main types according to its capabilities:

- **Narrow AI (Weak AI):** These systems specialize in specific tasks or narrowly defined domains. They can only perform functions they are explicitly designed for and cannot operate outside those parameters (Simplilearn, 2025). Examples include voice assistants (e.g., Siri, Alexa) and recommendation systems (e.g., Netflix, YouTube).
- **General AI (Strong AI):** This refers to AI systems capable of reasoning, learning, and problem-solving across various tasks in a way comparable to human intelligence. While it remains a theoretical concept, General AI is a long-term goal in the field. Despite substantial research efforts, numerous technical and ethical challenges remain (IBM Think, 2024).
- **Superintelligent AI:** This is an advanced form of AI that would surpass human intelligence in performance. Superintelligent AI could think more quickly, accurately, and creatively than humans. It remains a theoretical possibility and raises profound philosophical and ethical concerns, particularly regarding control and alignment with human values (Lumenalta, 2024).

1.3.2. Classification by Type of AI Systems

AI can also be classified according to its subfields and technological approaches:

- **Machine Learning (ML):** A subdomain of AI where machines learn from data, often by identifying patterns and correlations. ML includes supervised, unsupervised, and reinforcement learning.
- **Deep Learning (DL):** A more advanced subfield of ML that uses artificial neural networks to learn from large and complex datasets. DL is commonly applied in image recognition and natural language processing.
- **Natural Language Processing (NLP):** AI systems that understand, interpret, and generate human language. NLP is widely used in text analysis, translation, and language modeling.
- **Computer Vision:** This branch enables machines to interpret and understand visual inputs (e.g., images and videos). Applications include facial recognition and object detection.
- **Expert Systems:** These are knowledge-based systems that simulate the decision-making ability of human experts in specific domains.

1.3.3. Classification Based on Application Areas

Entering the 2020s, AI has accelerated in development due to advances in deep learning, the abundance of data, and increased access to computing power. Today, AI is delivering transformative applications across a wide range of sectors, including:

- **Healthcare:** Used for diagnosis, treatment planning, patient monitoring, and biotechnological research. AI contributes to personalized medicine, drug discovery, and clinical decision support.
- **Automotive and Transportation:** Applied in autonomous vehicles, traffic management, and road safety. AI-powered autonomous systems such as driverless cars, drones, and robots rely heavily on deep learning, computer vision, and reinforcement learning (Olabiyi, W., Akinyele, D., & Joel, E. (2025), p.10).
- **Finance:** Utilized in risk analysis, algorithmic trading, and credit scoring. AI is reshaping operational processes and decision-making in financial services.

- **Education:** Enables personalized learning, student performance prediction, and automated assessment.
- **Retail and Customer Service:** Includes applications such as chatbots, recommendation engines, and demand forecasting.
- **Administrative and Criminal Law:** AI is increasingly used in public administration and criminal justice. AI systems have begun playing roles in administrative decision-making and judicial procedures, especially as decision-support tools in courts.

1.3.4. Classification Based on Operational Methods

AI systems may also be classified according to their mode of operation:

- **Rule-Based AI:** These systems function based on a defined set of rules. They follow explicitly written instructions to perform specific tasks.
- **Data-Based AI:** These systems work by analyzing data and learning models from it. They improve over time by recognizing patterns and relationships in datasets.

2. Artificial Intelligence Tools that Facilitate Daily Life

Several artificial intelligence tools can significantly improve and simplify daily life. A few examples are briefly explained below:

1. **Virtual Assistants (Google Assistant, Siri, Alexa):** These tools help users perform tasks on their devices via voice commands, such as setting alarms, checking the weather, sending messages, or playing music quickly and efficiently.
2. **Email and Message Management (Boomerang, SaneBox):** These applications help organize and filter emails by setting reminders, automatically deleting spam, or highlighting important messages.
3. **Time and Task Management (Trello, Todoist, Notion):** These tools allow individuals to organize daily tasks, create to-do lists, track projects, and manage time effectively.
4. **Budgeting and Expense Tracking (Mint, YNAB - You Need A Budget):** AI tools that facilitate personal financial management by tracking expenditures, creating budgets, and helping users save money.

5. **Health Tracking (Fitbit, Apple Health, Google Fit):** These applications monitor daily physical activity, sleep patterns, and heart rate, enabling individuals to develop healthy living habits.
6. **Visual Editing (Prisma, DeepArt):** AI-powered photo editing applications that can turn photographs into artistic works or assist in quick and efficient image editing.
7. **Translation Applications (Google Translate, DeepL):** These tools provide fast translation across different languages, proving especially useful when traveling or encountering foreign language content.

These tools collectively contribute to a more organized, efficient, and productive daily life.

3. Individual and Efficient Use of Artificial Intelligence

Artificial intelligence (AI) offers significant potential for personal productivity and efficiency. Today, individuals can benefit from a wide range of AI tools to enhance both their personal and professional lives. Below are several ways in which AI can be used effectively:

3.1. Productivity Tools

AI can assist in the efficient execution of daily tasks:

- **Time Management and Planning:** AI-powered calendar and task management tools can help prioritize activities and schedule appointments. Virtual assistants like Google Assistant and Siri can set reminders and manage tasks.
- **Email and Message Management:** AI-based tools can filter emails, automate responses, and organize messages to help users manage their time more effectively.

3.2. Data Analysis and Research

AI can support academic or professional research and data analysis:

- **Data Analysis:** AI-powered analytical tools can identify patterns and extract key insights from large datasets.
- **Academic Writing and Research:** AI-based writing assistants can generate text, provide content suggestions, and assist in locating academic sources.

3.3. Education and Personal Development

AI supports educational processes through:

- **Language Learning:** AI-powered applications offer adaptive lessons based on personal learning speed, facilitating language acquisition.
- **Personalized Learning:** AI can deliver tailored lessons in specific subjects (e.g., law, economics, technology) and provide interactive feedback.

3.4. Creative Endeavors

AI can contribute to creative projects:

- **Art and Design:** AI tools support graphic design, photo editing, and music composition. Tools like MidJourney can generate digital art or assist music creators (Olabiyi, W., Akinyele, D., & Joel, E., 2025, p. 11).
- **Writing and Content Creation:** AI tools can aid in drafting creative texts, such as blogs, novels, or marketing copy.

3.5. Health and Quality of Life

AI tools also support personal well-being:

- **Personal Health Monitoring:** Devices and applications track health metrics and provide actionable insights on fitness, sleep, and nutrition.
- **Meditation and Mental Wellness:** AI-based mindfulness and meditation applications help manage stress and improve mental health.

3.6. Personal Assistants

AI-supported virtual assistants simplify daily tasks:

- **Virtual Assistants:** Tools such as Google Assistant, Amazon Alexa, or Apple Siri help with playing music, controlling smart home devices, or researching information.
- **Voice Command Execution:** AI enables users to perform actions using voice input—useful when multitasking.

3.7. Security and Privacy

Artificial intelligence can also be beneficial in ensuring personal security:

- **Privacy Protection:** AI-based security tools can protect against online identity theft and other security threats. For example, security software utilizing AI can detect viruses or malware.
- **Facial Recognition and Identity Verification:** Facial recognition systems can provide additional layers of security on devices such as smartphones, computers, or other gadgets.

3.8. Financial Management

Artificial intelligence can assist in better managing personal finances:

- **Budgeting and Expense Tracking:** AI can analyze spending habits to help individuals save money. Additionally, AI-based financial advisory services can be used for investment recommendations and financial planning.
- **Cryptocurrency and Investment:** For those interested in cryptocurrency or stock markets, AI-driven trading tools can support investment decisions by analyzing market trends.

Artificial intelligence is a powerful tool for enhancing efficiency, creativity, and security in individual use. By integrating AI into daily tasks, personal development, creative projects, and even health management, users can save time and improve productivity. The effective utilization of these tools not only facilitates daily life but also creates new opportunities and experiences.

Accordingly, AI has the potential to make life more efficient, creative, and secure in personal contexts. Through the use of AI in everyday activities, self-improvement efforts, creative endeavors, and health-related matters, individuals can gain time and boost their productivity. These tools can simplify life while simultaneously generating new opportunities and experiences.

4. Are States Free to Develop Artificial Intelligence Tools, or Do International Legal Constraints Exist?

In general, states are free to develop artificial intelligence tools within the framework of their domestic legal systems. However, this freedom is shaped by certain boundaries and emerging regulatory trends. While no globally binding legal framework for AI currently exists, several important international factors should be noted:

4.1. International Legal Frameworks: Numerous international organizations are making efforts to regulate the use of artificial intelligence. For instance, entities such as the European Union, the OECD, and the United Nations have issued guidelines and recommendations concerning the ethical use of AI systems. Although these regulations are not legally binding for states, they may impose certain restrictions based on intergovernmental agreements.

4.2. Human Rights and Ethical Standards: It is internationally important that AI applications comply with ethical standards, such as not violating human rights, avoiding discrimination, and protecting privacy. For example, the European Union's Artificial Intelligence Act and the Council of Europe's "Framework Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law" classify AI tools and impose special regulations on those deemed high-risk. Such regulations may set specific boundaries in the development of AI technologies.

4.3. Intellectual Property Rights: AI tools are subject to the intellectual property laws of the states in which they are developed. However, international treaties on patents and copyrights may also regulate certain aspects of AI innovation across jurisdictions.

4.4. Military and Security Use: AI systems used for defense or security purposes may be subject to additional international restrictions. For example, the use of autonomous weapons is under global scrutiny, and several initiatives call for international regulations on such systems.

4.5. Data Protection and Privacy Laws:

Legal frameworks such as the EU's General Data Protection Regulation (GDPR) significantly impact how AI systems process personal data and can require harmonization among states that develop AI tools for transnational use.

In Conclusion: While states enjoy a degree of sovereignty in developing AI tools, their activities are increasingly shaped by international human rights norms, ethical standards, and sector-specific regulations. Emerging legal instruments and multilateral agreements can impose limits and responsibilities on states to ensure that AI development aligns with global legal and ethical standards.

III- LIABILITY FOR DAMAGES CAUSED BY ARTIFICIAL INTELLIGENCE

1. Artificial Intelligence Law

A distinct legal discipline known as "Artificial Intelligence Law" has not yet been fully established. However, the legal regulation of artificial intelligence (AI) is increasingly gaining importance as a field that intersects with various branches of law. The most prominent of these include:

- **Technology and Innovation Law:** This area encompasses the development, use, and regulation of AI and other digital technologies. It includes legal issues related to technological innovation, copyright, and intellectual property rights.
- **Consumer Law:** The provision of AI products to consumers requires regulation in terms of safety, privacy, and compensation rights. Particularly, the effects of AI systems on consumers fall within this domain.
- **Data Protection and Privacy Law:** Since AI systems operate using large datasets, the protection of personal data is of critical importance. Regulations such as the European Union's General Data Protection Regulation (GDPR) ensure that AI applications comply with data security and privacy principles.
- **Tort and Liability Law:** Determining liability for damages caused by AI systems is a central concern, particularly in compensation claims and legal responsibility.
- **Human Rights and Ethics Law:** AI is directly connected to ethical issues such as human rights, equality, and non-discrimination. Accordingly, legal regulations regarding the ethical use of AI systems are increasing.

In conclusion, "Artificial Intelligence Law" is not yet a standalone field, but rather a multidisciplinary area formed by the convergence of existing legal branches. Nonetheless, ongoing debates may eventually lead to the emergence of new and more specialized fields of law in the future.

2. Competent Jurisdiction in Lawsuits for Damages Caused by Artificial Intelligence

The jurisdiction of lawsuits concerning damages caused by AI may vary depending on the type of damage and the specific nature of the incident. In general, such cases can be classified as follows:

- **Judicial Jurisdiction (Ordinary Courts):** If the damages arise from commercial or private relationships, the case typically falls under the jurisdiction of ordinary courts. For example, if an AI software misprocesses user data or a product malfunctions due to AI and causes material damage, a compensation claim may be initiated. These lawsuits generally concern commercial disputes, contractual breaches, or tortious liability.
- **Administrative Jurisdiction:** If the damage results from an AI system used in the provision of a public service—particularly when the damage concerns public order or the state’s regulatory or service obligations—then the case may fall under administrative jurisdiction. For instance, if a public institution using AI violates a citizen’s rights, the legal remedy may lie within administrative law. Furthermore, issues such as AI-induced security breaches or privacy violations related to public order may also invoke administrative jurisdiction.

In summary, if the damage originates from a private or commercial relationship, the competent court is the judicial one; if it concerns harm to the public, administrative courts are competent. However, since the legal framework surrounding AI is still evolving, the distinction is not always clear-cut.

3. Compensation Liability for Damages

3.1. Responsible Parties

Determining liability for damages caused by AI is a complex legal matter and generally involves three principal parties:

- **AI Developer or Manufacturer:** The individual or company that designs, develops, or produces the AI system may be held liable if the system malfunctions or causes harm. This liability is addressed within the framework of product liability. Developers can be held responsible for flaws in the design or defects in the code of the AI system.
- **User or Operator:** If an individual or institution uses the AI system for a particular purpose, that user may also be liable. Especially in cases where the AI system is used incorrectly or inappropriately, this may result in compensatory liability.

- **Owner or Administrator of the AI System:** If the AI system is operated by a company or organization, that entity may bear responsibility. The owner or manager is obliged to ensure the proper functioning of the AI system.

Ultimately, compensation liability may rest with different parties depending on the nature of the damage and how the AI system is used. However, there is still no uniform legal framework on this matter in most countries, and this remains an emerging area of law.

International organizations such as the European Union, the Council of Europe, the United Nations, and the OECD have issued legal instruments concerning the liability of AI systems, thereby imposing certain responsibilities on member states.

4. Legal Liability in International Instruments

4.1. Obligations Introduced by the Council of Europe’s Artificial Intelligence Framework Convention

The Council of Europe has prepared a legal instrument addressing the issue of legal liability related to artificial intelligence. On 17 May 2024, the Council adopted the “**Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law,**” marking the first binding international treaty on AI. This Convention establishes that AI systems must respect core principles such as **human dignity, equality, transparency, and accountability** throughout their entire life cycle.

The Convention includes measures such as **risk and impact assessments, informing affected individuals**, and ensuring their **right to appeal**. Member States are obligated to adopt legal, administrative, and other measures to ensure that AI systems are developed and used in accordance with human rights, democracy, and the rule of law. Accordingly, national legislation is expected to be harmonized with the Convention through appropriate secondary regulations.

4.1.1. Obligations Imposed on States by the AI Framework Convention (Council of Europe, 2024)

This Convention imposes, for the first time, **legally binding obligations** on States to ensure a **human-centric and rights-based approach** in the development, deployment, and regulation of AI systems.

Key obligations include:

4.1.1.1. Obligation to Respect Human Rights

States must ensure that AI systems do not violate fundamental rights such as:

- the right to life,
- data protection,
- prohibition of discrimination,
- freedom of expression,
- right to a fair trial, and
- privacy.

These rights must be protected throughout the design, development, use, and decommissioning of AI systems.

4.1.1.2. Transparency and Accountability

- AI systems must be understandable, traceable, and subject to review when necessary.
- Users must be informed when a decision is made by an automated system.
- A responsible human or institution must be clearly designated for the outcomes of the AI system.

4.1.1.3. Democratic Oversight and Supervision

- AI applications must be subject to democratic scrutiny and legal remedies.
- Public authorities must prevent misuse of AI systems and avoid exploitation of legal gaps.

4.1.1.4. Rule of Law and Judicial Review

- AI systems must operate within legal limits. Arbitrary or uncontrolled decisions must be avoided.
- Decisions involving AI must be subject to review by independent judicial bodies.

4.1.1.5. Risk-Based Approach

- States must classify AI systems according to risk levels and adopt preventive measures for high-risk systems.
- Risk assessments must consider the potential impact on human rights.

4.1.1.6. Protection of Vulnerable Groups

- The Convention highlights the protection of vulnerable groups such as children, the elderly, and persons with disabilities against AI-related risks.

4.1.1.7. International Cooperation and Knowledge Sharing

- States are encouraged to share best practices, enhance technical cooperation, develop oversight methods, and promote international standards.

4.1.1.8. Independent Oversight Mechanisms

- States are required to establish independent and impartial bodies to monitor compliance with the Convention.

4.1.2. Procedures and Sanctions in Case of Violations of the Convention

While the Convention outlines principles and obligations clearly, it **does not establish binding judicial procedures or direct enforcement mechanisms** in case of violations.

Instead of sanctions through criminal procedures or court rulings, the Convention relies on **institutional mechanisms** and **international dispute resolution pathways**.

4.1.2.1. Dispute Resolution (Article 28 – Amicable Settlement)

- Article 28 stipulates that disputes between parties should be resolved primarily through amicable means.
- If this fails, the **Committee of Ministers of the Council of Europe** may intervene and mediate between parties.
- However, no binding judicial decision or automatic sanction process is envisaged.

4.1.2.2. Monitoring Mechanism

- The Convention requires periodic reporting from States and encourages the sharing of good practices.
- An advisory or monitoring body may be established under the Council of Europe, although it does not have judicial authority like the European Court of Human Rights (ECHR).

4.1.2.3. Sanctions Left to National Law

- In the event of a violation, the enforcement of sanctions is left to the **domestic legal systems** of each member state.
- Thus, individuals who suffer rights violations resulting from AI usage must **seek redress through national legal avenues**.

Violation Scenario: Procedure to Be Followed	
Stage	Explanation
1. Amicable Settlement (Article 23)	Initially, parties are encouraged to reach an agreement through amicable means.
2. Participation of the Conference and Committees (Articles 23 and 28)	In dispute resolution, the “Conference of the Parties” mechanism is activated. Then:
3. Formal Dispute Settlement Mechanisms (Article 28)	If no agreement is reached, parties may resort to mediation, arbitration, or the International Court of Justice (ICJ). However, these mechanisms are only available between States Parties: EU member states cannot invoke them against each other.
4. Monitoring and Reporting (Chapter VII)	The implementation of periodic reporting and monitoring is envisaged. Additionally, international cooperation and oversight are conducted by the “Conference of the Parties.”

Consequently, unlike the European Convention on Human Rights (ECHR), this AI Framework Convention does not establish a direct individual complaint mechanism or judicial procedure. It is primarily a **framework instrument**, intended to **encourage States to implement appropriate measures in their domestic laws**, rather than impose sanctions directly.

In the event of disputes, mechanisms such as **international mediation, arbitration**, or recourse to the **International Court of Justice (ICJ)** may be pursued. The Convention introduces **preventive principles and obligations**, while leaving the implementation of **sanctions to national legal systems**.

At the international level, the Convention envisions structured cooperation, reporting, and the establishment of a **“Conference of the Parties”** to ensure compliance and promote alignment. At

the national level, recourse to administrative or judicial mechanisms regarding the obligations of the Convention is left to **each State's internal legal system**.

4.2. Obligations and Sanctions Imposed on EU Member States by the Artificial Intelligence Regulation (EU) 2024/1689

4.2.1. Primary Obligations under Regulation (EU) 2024/1689

In response to the rapid advancement of artificial intelligence (AI) technologies and the associated legal, ethical, and societal risks, the European Union (EU) identified the need for a comprehensive legal framework. Accordingly, on 21 April 2021, the European Commission published the “Proposal for an Artificial Intelligence Act,” marking the first holistic legislative initiative globally focused exclusively on AI. This proposal aimed to promote technological innovation while simultaneously safeguarding fundamental rights.

Following nearly three years of negotiations, the European Parliament and the Council of the EU officially adopted the **Artificial Intelligence Act (Regulation (EU) 2024/1689)**⁵ on 13 June 2024. This regulation applies not only to AI systems developed within the EU but also to those interacting with the EU market from abroad, thereby having a potentially global impact.

The principal objective of the regulation is to ensure that AI systems are developed and deployed in a manner that is safe, transparent, ethically sound, and respectful of fundamental rights. In this regard, EU Member States are subject to a wide range of obligations, including the classification of AI systems based on risk levels, oversight of high-risk systems, guarantees regarding safety and human oversight, adherence to data quality and transparency principles, among others.

The following sections provide a systematic overview of the key obligations imposed by the Regulation.

⁵ The Regulation is directly applicable in the Member States without requiring separate interpretation or adaptation into national law. Member States are obliged to comply with its provisions without the need for additional legislative measures.

4.2.1.1. Designation of National Competent Authorities (Articles 70 and 77)

Regulation (EU) 2024/1689 obliges Member States to establish the institutional infrastructure necessary for the effective oversight and implementation of AI systems. By 2 August 2025, each Member State is required to undertake the following institutional arrangements:

- **Notifying Authority:** Each Member State must designate at least one notifying authority responsible for the appointment and supervision of conformity assessment bodies. This authority plays a central role, particularly in issuing conformity certificates for high-risk AI systems.
- **Market Surveillance Authority:** An independent authority must be designated to monitor compliance with the regulation after AI systems are placed on the market. This body is critical for ensuring consumer safety, data protection, and the protection of fundamental rights.
- **Single Point of Contact:** Each Member State must establish a single national contact point to coordinate the implementation of the regulation and notify the European Commission accordingly.

Additionally, Member States are expected to identify independent supervisory or ethical review bodies—such as academic or public institutions—that evaluate the development, deployment, and monitoring of high-risk AI systems with a particular focus on human rights. The regulation thus mandates not only technical compliance mechanisms but also democratic accountability structures.

For example, the German Federal Government has proposed appointing the Bundesnetzagentur as the market surveillance authority and integrating advisory bodies such as the Data Ethics Commission (Datenethikkommission) to ensure ethical oversight.

4.2.1.2. Establishment of National “Regulatory Sandboxes” (Article 57)

The regulation mandates Member States to create regulatory frameworks that facilitate the controlled testing and development of innovative AI applications. Each Member State must establish at least one “AI regulatory sandbox” by 2 August 2026.

A regulatory sandbox refers to an experimental test environment supervised by regulatory authorities. Within this framework, developers are permitted to test high-risk AI systems under real-world conditions while assessing their compliance with regulatory requirements. The regulation emphasizes that these sandboxes must safeguard not only technical standards but also fundamental rights and ethical principles.

The sandbox programmes are expected to serve the following key purposes:

- **Promotion of innovation:** Small and medium-sized enterprises (SMEs) and startups may test their AI solutions with fewer bureaucratic barriers.
- **Clarification of legal uncertainties:** Developers can evaluate their systems' compliance with EU law in advance.
- **Early identification and mitigation of risks:** Potential rights violations or algorithmic biases can be detected and addressed during the testing phase.

For instance, France has launched a pilot platform titled “*AI pour la Société*” as part of its national AI strategy. This initiative supports preliminary evaluations of high-risk AI projects in sectors such as healthcare, transport, and public services.

4.2.1.3. Supervision and Oversight of High-Risk AI Systems

The regulation adopts a risk-based approach to categorising AI systems, introducing differentiated regulatory obligations according to the level of potential harm. The classification forms a core component of legal accountability and oversight mechanisms. AI systems are divided into four categories:

1. **Prohibited Systems:** These are systems considered to pose unacceptable risks to fundamental rights, such as real-time biometric identification by public authorities or subliminal techniques used to manipulate individual behaviour.
2. **High-Risk Systems:** These include systems deployed in critical domains such as infrastructure, education, employment, healthcare, and justice, where individual rights are significantly impacted.
3. **Limited Risk Systems:** Systems subject to transparency obligations, such as chatbots.

4. **Minimal Risk Systems:** Applications involving low-level user interactions with negligible risks.

Member States are specifically obliged to:

- Ensure compliance with safety, transparency, and quality standards through the establishment of oversight mechanisms.
- Implement conformity assessment procedures, either internally or via third parties, before placing high-risk systems on the market.
- Notify such systems to the European Commission and register them in the EU AI Database.

Case Example: In 2026, a German court ruled that an AI algorithm used in recruitment processes violated transparency obligations and prohibited its further use without revision. The court referenced Articles 8, 16, and 28 of the Regulation, underlining breaches of the risk-based obligations.

4.2.1.4. Transparency and Information Obligations

The Regulation imposes obligations on Member States to ensure the comprehensibility and transparency of AI systems. Specifically:

- Individuals interacting with AI systems must be clearly informed of this fact.
- Member States must establish regulatory frameworks to monitor compliance with transparency obligations.

Example: When an AI algorithm is used in recruitment, applicants must be explicitly informed that the evaluation process involves automated decision-making (Article 52).

4.2.1.5. Data Quality and Anti-Discrimination Measures

The Regulation requires that datasets used in AI systems be impartial, accurate, and relevant. Member States must:

- Develop guidelines and inspection mechanisms to prevent the use of discriminatory datasets, particularly in high-risk systems.

- Establish independent audit mechanisms for data governance processes.

Example: Article 10 mandates that training datasets be of high quality and representative of all segments of society.

4.2.1.6. Complaint Mechanisms and Legal Remedies

Effective legal and complaint mechanisms must be available for individuals whose rights are infringed by AI systems. Member States are required to:

- Guarantee access to fair compensation procedures for individuals adversely affected by AI systems.
- Strengthen the technical capacity and expert support available to courts handling such disputes.

Example: The French Data Protection Authority (CNIL) has developed an online platform that allows individuals affected by automated decisions to submit complaints directly.

4.2.1.7. Criminal and Administrative Sanctions

The Regulation provides for severe penalties for producers, importers, or distributors who fail to comply. Member States must:

- Determine and enforce administrative fines for non-compliance at the national level.
- Ensure that sanctions are proportionate, dissuasive, and effective.

Example: Article 71 allows for fines of up to €30 million or 6% of a company's global annual turnover for the unlawful deployment of high-risk systems.

4.2.1.8. Education, Awareness, and Capacity-Building Initiatives

Beyond regulation, the AI Act incorporates a strong educational component. Member States must:

- Develop awareness-raising and ethical training programmes for public and private sector personnel.
- Launch public information campaigns to promote the responsible and informed use of AI technologies.

Example: The Netherlands has launched a national education programme titled *AI4Youth* aimed at enhancing ethical AI literacy at both secondary and university levels.

Summary Table	
Obligation Category	Details
Direct Applicability	No need for domestic legal adaptation
National Notification Authorities	Notifying authority, single point of contact
Supervisory Bodies	Market surveillance and human rights oversight authorities
Regulatory Sandbox	At least one testing platform must be established
Risk Classification	High, limited, and minimal risk system differentiation
Monitoring and Supervision	Integration with sector-specific regulations
Enforcement Mechanisms	Significant financial penalties and compensation remedies

In this context, Regulation (EU) 2024/1689 imposes a series of obligations on EU Member States, including the designation of implementing authorities, the establishment of risk management frameworks, the operationalisation of regulatory sandboxes, the strengthening of oversight mechanisms, and the provision of effective sanctions and remedies in case of violations. These obligations not only ensure legal compliance but also contribute to the European Union’s broader objective of fostering a trustworthy and human rights-respecting AI ecosystem.

4.2.2. Sanctions Envisaged in Case of Non-Compliance with the Regulation

Regulation (EU) 2024/1689 on Artificial Intelligence introduces a detailed sanction regime applicable across EU Member States, with the aim of preventing unlawful uses of AI systems and enhancing deterrence. The Regulation’s sanctions section provides for a range of measures, including administrative fines, corrective actions, and market restrictions, which vary depending on the nature and severity of the infringement.

4.2.2.1. Sanctions for Prohibited Practices (Article 5)

In cases involving explicitly prohibited AI practices—such as real-time biometric surveillance, social scoring systems, or manipulative user interfaces—legal entities may be subject to the following sanctions:

- An administrative fine of up to **7% of the total worldwide annual turnover**, or
- Up to **EUR 35 million**,
- Whichever is higher shall apply.

4.2.2.2. Non-Compliance with Obligations for High-Risk Systems

Where there is a breach of obligations concerning high-risk AI systems, as set out in Articles 16, 26, 31, 33, 34, and 50 (including conformity assessment, registration, oversight, and transparency requirements), the following penalties may be imposed:

- Up to **3% of the total worldwide annual turnover**, or
- Up to **EUR 15 million**.

4.2.2.3. Provision of Incorrect or Misleading Information

If a provider or user submits incorrect, misleading, or incomplete information to the competent authorities, the applicable sanctions are:

- Up to **1% of the total worldwide annual turnover**, or
- Up to **EUR 7.5 million**.

4.2.2.4. Proportionality Measures for SMEs and Start-ups

To support the development of small and medium-sized enterprises (SMEs) and start-ups, the Regulation adopts the principle of proportionality with respect to sanctions. Accordingly, fines shall be based on:

- The **lower amount** between the percentage of turnover and the fixed monetary cap.

4.2.2.5. Additional Administrative Measures

In addition to financial penalties, the European Commission and national supervisory authorities may impose further administrative actions, including:

- **Withdrawal** of non-compliant AI systems from the market,
- **Restrictions or suspensions** on the use of such systems,
- **Mandatory implementation** of corrective measures to remedy the infringement.

Summary Table

Type of Violation	Company (Penalty)
Use of Prohibited AI	7% or €35 million
Violation of High-Risk or Transparency Rules	3% or €15 million
Incorrect/Missing Information	1% or €7.5 million
SME/Start-up	Institution-specific lower thresholds apply

Accordingly, the AI Regulation establishes an effective compliance and enforcement framework across EU Member States, ensuring strong deterrence through substantial administrative fines and corrective measures. Member States are required to establish National Competent Authorities responsible for implementing these sanctions and to guarantee compensation rights for individuals affected by non-compliance.

4.2.3. United Nations Initiatives in the Field of Artificial Intelligence and Obligations Imposed on States

The development and deployment of artificial intelligence (AI) technologies within a framework of ethics, human rights, and the rule of law has become a key concern not only for regional actors but also for the international community. In addition to regional instruments such as the Council of Europe's Framework Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law, the United Nations (UN) has undertaken several initiatives to promote the alignment of AI with global human rights and ethical standards.

Among the most prominent UN efforts in the AI domain are the reports by special rapporteurs of the Human Rights Council, the UNESCO “Recommendation on the Ethics of Artificial Intelligence,” and resolutions adopted by the UN General Assembly. These instruments emphasize that states must ensure AI systems respect human rights, avoid discrimination, and adhere to principles of transparency and accountability throughout their development and use.

In particular, UNESCO's 2021 Recommendation on the Ethics of Artificial Intelligence provides guidance to states and stakeholders for the development and implementation of AI systems that are grounded in values such as justice, human dignity, inclusiveness, environmental sustainability, and accountability. Similarly, reports issued by UN special rapporteurs call for preventive

measures to address potential adverse impacts of AI and advocate for the establishment of shared international standards through global cooperation.

However, unlike the Council of Europe's convention, there is currently no binding and comprehensive international treaty on artificial intelligence under the UN framework. Nevertheless, the UN's non-binding instruments and recommendations form a normative foundation for national regulations focused on human rights-based AI governance. In this context, the UN's contributions are pivotal to the development of global standards for regulating AI in accordance with ethical and legal principles.

In conclusion, considering the societal impact of AI technologies, the UN's efforts, when viewed in conjunction with regional instruments such as the Council of Europe's framework convention, play a critical role in defining state responsibilities and fostering international collaboration.

4.2.4. Other Key International Instruments and Organizations

A comprehensive approach to legal responsibility in the field of AI must consider not only the initiatives of the Council of Europe, the European Union, and the United Nations but also fundamental human rights instruments, international intellectual property frameworks, and data protection regimes.

The Universal Declaration of Human Rights (UDHR) and the International Covenant on Civil and Political Rights (ICCPR) provide international standards for essential rights and freedoms such as privacy, non-discrimination, and fair trial, which must be safeguarded in the design and deployment of AI systems. These documents establish the legal foundation for addressing potential intrusions on individual rights caused by AI applications.

In the domain of intellectual property, the World Intellectual Property Organization (WIPO) is formulating regulatory approaches to address the use of AI in creative processes and the resulting responsibilities related to patents, copyrights, and innovation. WIPO's work is significant in shaping the legal framework for AI-generated innovations.

Furthermore, the International Organization for Standardization (ISO) and the International Electrotechnical Commission (IEC) focus on setting technical standards for AI systems. In

particular, ISO/IEC JTC 1/SC 42 plays a key role in establishing standards that ensure AI system reliability and legal accountability.

The European Union's General Data Protection Regulation (GDPR, 2018) also plays a foundational role in AI governance by stipulating data protection standards for the processing of personal data by AI systems, including responsibility for data management, oversight, and enforcement in the event of breaches.

Additionally, institutions such as the International Telecommunication Union (ITU) and the Organisation for Economic Co-operation and Development (OECD) contribute to global standard-setting by issuing ethical guidelines and regulatory frameworks for AI, thereby promoting harmonization across jurisdictions.

Lastly, the UNIDROIT Principles and Model Laws, developed to ensure coherence in international commercial law, provide valuable references for addressing responsibility and dispute resolution in the commercial deployment of AI technologies.

These multilateral instruments and standards serve to facilitate the secure, ethical, and legally compliant use of AI on a global scale by encouraging convergence among countries and sectors. In doing so, they offer a coordinated international response to the legal and ethical challenges posed by AI technologies.

4.3. Liability of Artificial Intelligence in Administrative and Criminal Law

The increasing use of artificial intelligence (AI) technologies in public administration and criminal justice has raised significant legal and ethical concerns, particularly regarding the influence of such systems on decision-making processes. In this context, the use of AI in administrative decisions must be assessed with regard to principles of judicial review, transparency, and accountability, especially in terms of the protection of fundamental rights and freedoms⁶.

The Council of Europe has conducted a series of studies assessing the impact of artificial intelligence on human rights, democracy, and the rule of law. In its framework documents titled *'Artificial Intelligence, Human Rights, Democracy and the Rule of Law'*, the Council emphasizes

⁶ European Union Agency for Fundamental Rights (FRA), *Getting the future right – Artificial intelligence and fundamental rights*, Publications Office of the European Union, 2020

the responsibilities associated with the use of AI in administrative practices and criminal proceedings⁷. These documents underline that decisions generated by AI must be subject to human oversight and that automated decision-making processes must adhere to the principles of transparency and the obligation to provide reasoning.

Similarly, the Council of Bars and Law Societies of Europe (CCBE) has issued a recommendation report highlighting the potential risks posed by AI to the right to a fair trial, professional confidentiality, and ethical obligations. The report particularly stresses that AI should not undermine the role of judges and defense counsel in the justice system⁸.

Accordingly, the legal liability of AI in the context of administrative and criminal law must be evaluated from the perspectives of both the Council of Europe and the European Union, considering the profound implications for fundamental rights and the rule of law.

4.3.1. The Council of Europe Perspective

Legal challenges arising at the intersection of artificial intelligence technologies and the exercise of public authority have been addressed with particular attention by the Council of Europe. In this context, multidimensional efforts in both criminal and administrative law aim to ensure that the integration of AI into public functions respects fundamental rights and is carried out in a transparent and accountable manner.

4.3.1.1. Artificial Intelligence and Criminal Law

The European Committee on Crime Problems (CDPC) of the Council of Europe held a special thematic session in 2018 entitled “Artificial Intelligence and Criminal Law.” During this session, responsibility issues stemming from AI-based technologies such as autonomous vehicles and driverless systems were discussed⁹. Key topics included the attribution of responsibility in cross-border incidents and the legal qualification of an “offender” in events lacking human intervention.

Subsequently, the CDPC published a feasibility study in 2020 and, as of 2022, outlined the possibility of introducing a dedicated legal framework for AI within the scope of criminal law¹⁰.

⁷ Council of Europe, *Feasibility Study on a Legal Framework on AI Design, Development and Application based on the Council of Europe’s Standards*, CAHAI, 2021.

⁸ Council of Bars and Law Societies of Europe (CCBE), *CCBE Recommendations on the Use of Artificial Intelligence (AI) in the Justice System and by the Legal Profession*, November 2020.

⁹ Council of Europe, CDPC. (2018). *Special Session on Artificial Intelligence and Criminal Law*. Strasbourg

¹⁰ Council of Europe, CDPC. (2020). *Feasibility Study on AI and Criminal Law*

This initiative primarily aims to clarify the legal concept of liability in incidents involving autonomous transport systems or automated decision-making processes.

4.3.1.2. Artificial Intelligence and Administrative Law

Since 2022, the European Committee on Legal Co-operation (CDCJ) has undertaken a series of activities under the theme “Artificial Intelligence and Administrative Law.” As algorithmic systems increasingly influence administrative decision-making, the need to revise legal frameworks to ensure legal certainty for individuals has come to the forefront.

In this regard, the CDCJ published a comparative analysis report in 2022 on the compatibility of AI with administrative justice systems in Council of Europe member states¹¹. Moreover, the Council’s citizen-focused guide “*Administration and You*” was updated in 2024 to include algorithmic decision-making processes¹².

These documents emphasize key principles such as the legality of administrative decisions, access to remedies against algorithmic decisions, transparency of systems, and the necessity of human oversight.

4.3.1.3. Expert Reports and Recommendations

Under the auspices of the Council of Europe, the HUDERIA initiative (Human Rights, Democracy and the Rule of Law Impact Assessment of AI Systems) provides comprehensive framework documents assessing the impact of AI on human rights, democratic values, and the rule of law¹³. Supported by the Alan Turing Institute, these studies aim to develop assessment tools and compliance standards.

In addition, the Council of Europe adopted a Recommendation in 2024 on the use of AI in criminal justice and probation services. This document underlines principles such as legality, legal certainty, explainability, human oversight, and accountability¹⁴.

In sum, the Council of Europe has guided the integration of AI into criminal and administrative law not only through normative instruments but also through practical tools and guidance

¹¹ Council of Europe, CDCJ. (2022). *Comparative Analysis on AI and Administrative Justice*.

¹² Council of Europe. (2024). *Administration and You – AI Updated Edition*.

¹³ Alan Turing Institute & Council of Europe. (2022). *HUDERIA – AI Human Rights Assessment Framework*.

¹⁴ Council of Europe. (2024). *Recommendation CM/Rec(2024)1 on the Use of AI in Criminal and Probation Systems*.

documents. In this framework, significant contributions have been made on both conceptual and technical levels in the following areas:

- Clarifying the definition of the offender,
- Ensuring the legality of algorithmic administrative decisions,
- Conducting impact assessments,
- Safeguarding fundamental rights and establishing effective remedies.

The Council of Europe's pioneering efforts in this domain serve as a critical reference for the liability and enforcement models developed by institutions such as the European Union and the United Nations.

4.3.2. The European Union Perspective

The deployment of artificial intelligence technologies within public administration and criminal justice systems is regarded by the European Union (EU) not merely as a technical development but as a profound legal and ethical issue. The protection of fundamental rights and freedoms, legal predictability, and accountability constitute the core of the EU's AI policies.

4.3.2.1. The AI Act and Public Sector Applications

Through the adoption of Regulation (EU) 2024/1689—commonly referred to as the AI Act—the European Union has established binding provisions concerning the use of AI in the public and criminal justice sectors. The Regulation sets forth specific obligations for AI systems deployed by public authorities. In particular, high-risk AI systems used in areas such as migration, border control, social welfare, and criminal investigations are subject to strict oversight, conformity assessments, mandatory user training, and human oversight¹⁵.

Articles 6 and 7 of the Regulation stipulate that systems capable of affecting criminal law (e.g., facial recognition, behavior prediction, police surveillance) may only be used under specific conditions and in accordance with the principle of proportionality. For AI systems deployed in the

¹⁵ AI Act 2024/1689, Articles 6, 9, 29, 54

context of criminal justice, compliance with the Charter of Fundamental Rights of the European Union is of paramount importance¹⁶.

4.3.2.2. Use of AI in Criminal Justice and Protection of Fundamental Rights

The Charter of Fundamental Rights of the EU imposes limitations on the use of AI in criminal contexts, particularly in light of the rights to a fair trial (Article 47), privacy and data protection (Articles 7–8), and non-discrimination (Article 21). In joint statements, the European Data Protection Supervisor (EDPS) and the European Data Protection Board (EDPB) have emphasized that algorithmic systems used in criminal investigations must be designed in full compliance with the principles of “data minimization,” “purpose limitation,” and “explainability.”¹⁷

Given the substantial human rights risks posed by applications such as facial recognition, behavioral analysis, and predictive systems, the EU legal framework has developed specific liability regimes for these areas.

4.3.2.3. Use of AI in Administrative Decisions and Legal Safeguards

The EU acknowledges that AI systems used in public administration may directly impact individuals’ rights and therefore classifies them as high-risk. In domains such as social welfare, immigration procedures, and tax administration, algorithmic decision-making is permitted only if individuals retain access to appeal mechanisms, decisions are accompanied by adequate reasoning, and human oversight is guaranteed¹⁸.

Furthermore, Article 68 of the AI Act mandates that member states maintain publicly accessible registries of high-risk AI systems used by public authorities and conduct regular audits of these systems.

4.4. Shared Principles of the Council of Europe and the European Union

Both institutions converge on several fundamental principles of legal responsibility in AI governance:

¹⁶ Charter of Fundamental Rights of the European Union, 2012/C 326/02.

¹⁷ European Data Protection Board (EDPB) & European Data Protection Supervisor (EDPS). (2021). Joint Opinion on the AI Act.

¹⁸ Fundamental Rights Agency (FRA). (2022). *Data Quality and AI: Fundamental Rights Considerations for Algorithmic Decision-Making*. Vienna.

- **Human Oversight and Responsibility:** Final decision-making authority must remain with human actors; AI is to function as a supportive tool only.
- **Explainability and Transparency:** The decision-making processes of AI systems must be comprehensible and subject to audit.
- **Legal Remedies:** Individuals adversely affected by AI-supported actions must have access to effective judicial remedies.
- **Compliance with Fundamental Rights:** All AI applications must undergo scrutiny to ensure compatibility with human rights standards.

In this respect, both the Council of Europe and the European Union do not confine AI-related legal responsibility to technical compliance alone; rather, they require a fundamental rights-oriented approach. This ensures a legal safety net that mitigates structural risks arising from the fusion of artificial intelligence and public power.

IV. LIABILITY FOR DAMAGES CAUSED BY ARTIFICIAL INTELLIGENCE: NATIONAL PRACTICES

1. United States of America

1.1. Product Liability Case Against Character.AI and Google (USA, Florida, 2025)

Among the lawsuits filed against artificial intelligence systems in the United States, the groundbreaking case *Garcia v. Character.AI & Google* has set a significant precedent regarding liability for damages caused by AI. Heard in a federal court in Florida, the case involved a mother's claim that the psychological effects generated by the chatbot Character.AI led to the suicide of her 14-year-old son, Sewell Setzer.

At the core of the litigation was the allegation that the chatbot engaged with the child as a “romantic partner” and gave the impression of being a real human, exerting psychological influence strong enough to induce suicide.

U.S. District Judge Anne C. Conway ruled that this AI-based application could be legally classified as a “product” under product liability law, and that claims related to design defects were admissible for trial. The court further recognized that Character.AI was aware of the potential for psychological harm and thus owed a duty of care to implement safety measures. Accordingly, the

court affirmed that companies bear a legal duty of care to “take necessary precautions to prevent foreseeable harm.”¹⁹

Claims against Google were also included, with the court finding that Google's indirect but legally significant responsibility arises from its contribution to the AI infrastructure and model development. This approach expands liability exposure to component providers such as infrastructure or foundational model suppliers.

While the claim exclusively alleging “intentional infliction of emotional distress” was dismissed for insufficient grounds, all other allegations—namely design defects, negligence, and failure to warn—were admitted to the litigation process.

This case represents a novel legal approach addressing personal injuries caused by AI-supported systems, design defects, and foreseeable risks. It also reinforces the evolving tendency to treat software and algorithmic systems as “products” under traditional liability frameworks. Particularly in the U.S., this decision serves as a pivotal precedent affirming that AI can be held liable under personal injury law.

1.2. Kadrey v. Meta Platforms Inc. – Copyright Infringement Case (USA, 2025)

The *Kadrey v. Meta Platforms Inc.* case, adjudicated in 2025 and attracting significant public attention, marks a critical judicial examination of the legal boundaries concerning AI systems trained on copyrighted works. Plaintiffs included Pulitzer Prize-winning authors, academics, and artists who alleged that Meta utilized their copyrighted works without permission in training its open-source large language model, LLaMA.

In the federal court in San Francisco, U.S. District Judge Vince Chhabria dismissed the case on the grounds that the plaintiffs’ legal arguments were insufficient. However, the judge explicitly clarified that the dismissal did not signify that Meta’s conduct was lawful. Instead, the ruling indicated that the evidence and legal rationale presented were not sufficiently persuasive to establish copyright infringement in this instance²⁰.

¹⁹ Anzalone & Doyle Trial Lawyers. (2024, May 21). *Federal judge allows product liability suit against Google and Character.AI to proceed in teen suicide case*. Retrieved from <https://www.anzalonelaw.com/federal-judge-allows-product-liability-suit-against-google-and-character-ai-to-proceed-in-teen-suicide-case>

²⁰ AP News. (2025, June 26). *Judge dismisses authors' copyright lawsuit against Meta over AI training*. Retrieved from <https://apnews.com/article/e77968015b94fbbf38234e3178ede578>

The court emphasized that the transformative use of copyrighted works by AI systems may, under certain conditions, be protected under the “fair use” doctrine of U.S. copyright law. Nonetheless, allegations that Meta obtained copyrighted materials from pirate databases (“shadow libraries”) were taken seriously, leaving the door open for similar lawsuits by other authors.

Judge Chhabria stated: “This decision does not confirm the legality of Meta’s use of copyrighted works in training its models. It only concludes that the plaintiffs’ claims are insufficient and lack the requisite evidentiary support.”

Moreover, the court did not underestimate the impact of Meta’s practices on the creative industry. Court documents revealed that Meta systematically scraped copyrighted works from pirate sources for the LLaMA model and that this process was discussed at senior management levels within the company.

2. Canada

2.1. Moffatt v. Air Canada (2024) – Corporate Legal Liability

This case constitutes an important precedent regarding the legal liability of companies for AI-powered chatbots. In the lawsuit, an Air Canada website chatbot provided incorrect information to a passenger, which led to the purchase of a mispriced ticket. The court held that the chatbot functioned as a representative of the company and that the company was liable for the erroneous information provided²¹.

As a result, the company was ordered to pay monetary compensation to the passenger. This ruling represents a critical milestone in clarifying the liability obligations of companies deploying AI applications.

3. People's Republic of China

3.1. Li Yunkai v. Liu Yuanchun – Copyright Protection for AI-Generated Images in China (Beijing Internet Court, 2023)

Case Overview: On November 27, 2023, the Beijing Internet Court issued a landmark ruling on whether content generated by artificial intelligence (AI) is eligible for copyright protection in

²¹ Dentons Data. (2024). *Airline ordered to compensate a B.C. man because its chatbot provided inaccurate information*. <https://www.dentonsdata.com/airline-ordered-to-compensate-a-b-c-man-because-its-chatbot-provided-inaccurate-information/>

China. Plaintiff Li Yunkai filed suit against defendant Liu Yuanchun for unauthorized online use of an image created using the generative AI tool "Stable Diffusion," which Li had published on social media.

*Court's Evaluation and Reasoning*²²: The Beijing Internet Court based its decision primarily on two copyright criteria:

1. Intellectual

Input:

The court found that the plaintiff made significant creative contributions to the image's creation process, including character design, prompt selection, use of negative prompts, parameter adjustments, and selection of appropriate outputs. These contributions demonstrated an original creative process directed by a human rather than mere automatic AI output.

2. Originality:

The court assessed that the images were produced not only mechanically but also reflected the plaintiff's aesthetic preferences and creative input. The plaintiff made corrections and improvements to the outputs, determined composition and style, thereby imparting personal expressive features.

Outcome: The court ruled that although the images were AI-generated, they embodied the plaintiff's personal intellectual labor and were thus protectable as graphic works under Chinese Copyright Law. The defendant was found to have infringed the plaintiff's rights by unauthorized use and was ordered to pay 500 RMB (approximately 70 USD) in damages and 50 RMB for legal costs.

The court notably stated: "AI models are tools. The creative process remains directed by humans, with aesthetic and intellectual input underpinning content production. Therefore, AI-assisted works that involve intellectual contribution should be protected under copyright law."

²² Beijing Internet Court. (2023, November 27). *Li Yunkai v. Liu Yuanchun*. Judgment concerning the copyrightability of AI-generated artwork created using Stable Diffusion. Case summary retrieved from [China IP Law Firm Reports].

4. Legal Liability for Artificial Intelligence in European Countries: Case Studies and Court Decisions

The rapid advancement of AI technologies has led to concrete legal cases addressing liability in various European jurisdictions. This section analyzes how legal responsibility for damages caused by AI is being shaped through notable court decisions and ongoing litigation in Switzerland, Ireland, France, Germany, the Czech Republic, and at the European Union level.

4.1. Switzerland: Strict Liability in an Autonomous Vehicle Accident

Case Summary: In Switzerland, an autonomous vehicle failed to stop at a pedestrian crossing, injuring a pedestrian. The vehicle's autonomous driving system was engaged. Following the incident, compensation claims were brought against the vehicle owner and manufacturer.

Legal Issue: The question arose as to who should be held liable for the AI system's fault in a driverless car—owner, manufacturer, or software developer?

Legal Assessment: Swiss Road Traffic Act (Art. 58 RTA) imposes strict liability on both vehicle owners and manufacturers, without requiring fault. Therefore, liability for the AI error was directly assigned to the manufacturer.

Outcome: Compensation was awarded against both the vehicle owner and the manufacturer, affirming manufacturer liability for AI-based products causing physical harm in Switzerland.

4.2. Ireland: Diagnostic Error in Medical AI Software

The 2018 CervicalCheck scandal involved false-negative results in cervical cancer screening smear tests, delaying diagnosis and treatment for affected women. Patient Vicky Phelan filed a compensation claim after the incorrect test result adversely affected her treatment timeline. The court ruled in favor of Phelan, ordering damages amounting to approximately €2.5 million²³. This case serves as a key precedent in evaluating human and technology-based errors in medical diagnostics within health law.

4.3. France: Copyright Lawsuit Against Meta (2025)

²³ Wikipedia contributors. (2024). *CervicalCheck cancer scandal*. Wikipedia. https://en.wikipedia.org/wiki/CervicalCheck_cancer_scandal

Case Summary: In March 2025, three major French cultural associations—Syndicat National de l'Édition (SNE), Société des Gens de Lettres (SGDL), and Syndicat National des Auteurs et Compositeurs (SNAC)—filed suit against Meta (owner of Facebook, Instagram, WhatsApp) in Paris.

The lawsuit concerned Meta's alleged use of copyrighted literary works on a large scale for AI model training without authorization. Plaintiffs claimed this practice violated international intellectual property laws and constituted the offense of “economic parasitism” (*parasitisme économique*)²⁴.

Court Decision:

- The Paris Judicial Court (*tribunal judiciaire*) granted a preliminary injunction at the plaintiffs' request, suspending Meta's data mining activities.
- The ruling underscored the necessity of both legal authorization and transparency before copyrighted content may be included in AI training datasets.

This lawsuit is a significant indication of cultural producers in the EU asserting legal defenses against AI-related copyright infringements, and serves as a precedent regarding the commercial aspects and legal frameworks governing AI training.

4.4. Germany

4.4.1. Google Autocomplete Decisions (BGH VI ZR 269/12)

Case Summary: A German joint-stock company and its CEO (plaintiffs) filed suit against Google Inc. for personality rights violations after Google's German search engine suggested defamatory terms such as “Scientology” and “Betrug” (fraud) in autocomplete suggestions linked to their names, allegedly harming their reputation.

*Legal Assessment*²⁵:

²⁴ **Reuters** (2025, 12 Mart). *French publishers and authors file lawsuit against Meta in AI case.* Reuters. <https://www.reuters.com/technology/artificial-intelligence/french-publishers-authors-file-lawsuit-against-meta-ai-case-2025-03-12/>

²⁵ Bundesgerichtshof (BGH). (2013). *Urteil vom 14. Mai 2013 – VI ZR 269/12 [Google Autocomplete-Entscheidung]*. Karlsruhe: Bundesgerichtshof. Karar metni için bkz: <https://www.bundesgerichtshof.de/SharedDocs/Pressemitteilungen/DE/2013/2013111.html>

- Both first-instance (LG Köln) and appellate courts (OLG Köln) dismissed the case. However, on May 14, 2013, the Federal Court of Justice (Bundesgerichtshof, BGH) overturned these decisions.
- The court held that autocomplete suggestions, though based on user behavior, constitute algorithmically generated content attributable to Google.
- If such content creates false associations that infringe personality rights, a violation occurs.
- Google has a duty to act upon notifications of such infringements, though it is not required to pre-screen all suggestions proactively.

Outcome: BGH found that the plaintiffs' personality rights were violated and remanded the case for retrial. This ruling clarified that algorithmically generated content can be subject to privacy and personality rights protection, thereby delineating responsibilities for internet companies.

4.4.2. Germany (Hamburg): Kneschke v. LAION e.V. (October 2024) – Application of Text and Data Mining (TDM) Exception at the Intersection of AI and Copyright Law

Case Summary: Photographer Robert Kneschke sued LAION e.V. at the Hamburg Regional Court for copyright infringement, alleging unauthorized inclusion of his photo—uploaded to a stock photo site—in the open-access LAION-5B training dataset used for AI research. The stock photo site's terms of use prohibited automated access ("bots").

*Legal Assessment and Decision*²⁶:

- The court assessed LAION's dataset creation under the scientific research exception for TDM (Text and Data Mining) pursuant to Section 60d of the German Copyright Act (UrhG).

²⁶ Hembt, S., Lutzhöft, N., & Bond, T. (2024, October 1). *Long-awaited German judgment by the District Court of Hamburg (Kneschke v. LAION) on the text and data mining exception(s)*. Bird & Bird. Retrieved from <https://www.twobirds.com/en/insights/2024/germany/long-awaited-german-judgment-by-the-district-court-of-hamburg-kneschke-v-laion>

Additionally see: Hamburg District Court, *Kneschke v. LAION e.V.*, 310 O 227/23 (2024), karar özeti: ChatGPTIsEatingTheWorld.com

- The court held that LAION’s activity qualified as “scientific research” due to its systematic knowledge generation objective.
- Since LAION made the dataset publicly and freely available, the use was not considered commercial.
- The work was only downloaded as a “preview” for TDM purposes, thus the use did not qualify as a “temporary copy” exception but rather fell under the Section 60d TDM exception.
- The court noted that the rights holder had neither effectively opted out nor provided machine-readable consent.

Significance: This ruling is among the first European judicial precedents explicitly recognizing AI training data usage within copyright exceptions. It provides important legal guidance on the legitimacy of large datasets employed in AI training.

4.5. Prague City Court, Czech Republic – AI-Generated Content Decision (2024)

Case Overview: A user created a simple image depicting “two individuals signing an employment contract in a law office in Prague, showing only their hands,” using OpenAI’s DALL•E artificial intelligence model. The image was published on the user’s website and was later used without authorization by a law firm on its own site. The user filed a lawsuit asserting copyright claims over the work.

Legal Evaluation and Court Decision: The Prague Court emphasized the following points in its ruling²⁷:

1. **Only natural persons can hold copyright:** Under Czech Copyright Law, copyright is granted exclusively to “natural persons who produce original works through creative activity.” AI tools fall outside this scope.

²⁷ Chloupek, V., & Taimr, M. (2024, May 29). *Czech court denies copyright protection of AI-generated work in first ever ruling*. Bird & Bird. Retrieved from <https://www.twobirds.com/en/insights/2024/czech-republic/czech-court-denies-copyright-protection-of-ai-generated-work-in-first-ever-ruling>

2. **Claimant must prove creativity:** Although the plaintiff claimed that the prompt (instruction) involved unique creative contribution, they failed to document this. Therefore, the plaintiff was not recognized as an “author.”
3. **AI output cannot be copyrighted:** The court ruled that content generated by AI does not qualify as an “original personal creation” and thus cannot be protected by copyright.
4. **Prompt remains at the idea level:** The court held that the prompt constitutes an idea or theme, which is not subject to copyright protection.

Outcome and Significance:

- Works generated solely by AI are not directly protected by copyright law.
- Persons seeking copyright protection must substantiate human creative contribution with ethical documentation.
- The decision supports the widely held principle in Europe of the “human author requirement,” legally clarifying that AI outputs are not eligible for copyright protection.

4.6. Court of Justice of the European Union (CJEU) – Preliminary Ruling (Like Company v. Google Ireland, C-250/25)

Case Overview: Hungarian news publisher Like Company filed suit with the Budapest Regional Court, alleging that Google’s LLM-based chatbot (Gemini/Bard) displayed portions of their news content without authorization in response to user queries on their own website. The court referred the matter to the CJEU for a preliminary ruling between June 13, 2023, and February 7, 2024, arguing that such activity constitutes both reproduction and making available to the public under EU Copyright Directives²⁸.

Legal Issues & Questions Posed: The Budapest court sought interpretation of several key provisions within the EU Copyright Directives²⁹:

²⁸ Court of Justice of the European Union. (2025). *Case C-250/25, Like Company v. Google Ireland Limited – request for preliminary ruling*. Submitted 3 April 2025 by Budapest Körményi Törvényszék

²⁹ William Fry. (2025, May 27). *CJEU to Rule on AI Chatbots and Copyright in Landmark Case Against Google*. William Fry

1. Does the inclusion of text in chatbot responses constitute “making available to the public” under DSM Directive Article 15(1) and InfoSoc Directive Article 3(2), protecting press publishers’ rights?
2. Does training the chatbot involve “reproduction” of digital content during large language model (LLM) training under InfoSoc Directive Article 2?
3. If so, can this reproduction benefit from the Text and Data Mining (TDM) exception in DSM Directive Article 4?
4. If a response to a user query contains text directly from the original publication, does this also amount to reproduction?

This case represents one of the first preliminary rulings requested from the CJEU concerning generative AI outputs via chatbots. The ruling will clarify the boundaries of EU copyright law in relation to both software training activities and AI-generated responses provided to users.

4.7. Evaluation

These case studies demonstrate that in response to harms caused by artificial intelligence:

- Classical liability principles such as product/service defect and negligence can rapidly be applied in new contexts;
- New case law is emerging on copyright and the use of copyrighted content;
- Legal systems across different countries are swiftly adapting to the challenges posed by AI technologies.

IV. CONCLUSION AND EVALUATION

The widespread use of artificial intelligence technologies in daily life has led to the emergence of autonomous decision-making mechanisms that can directly impact individuals’ rights. These developments have necessitated that legal systems address not only the traditional liability paradigms between individuals but also new forms of responsibility created by technological actors. The unpredictable, autonomous, and adaptive nature of AI systems requires a reconsideration of existing legal liability theories and the development of new models.

The European Union’s 2024 Artificial Intelligence Regulation (Regulation 2024/1689) and the Council of Europe’s 2024 Framework Convention on Artificial Intelligence, Human Rights, Democracy, and the Rule of Law stand out as the most significant international regulatory initiatives in this field. Both instruments establish general principles for AI liability—such as transparency, human oversight, explainability, respect for fundamental rights, and accountability—and additionally provide for specific oversight and compliance mechanisms for high-risk systems.

However, these documents are not yet supported by binding judicial decisions that have gained clarity in their implementation. In this context, strengthening legal liability regimes requires not only regulatory principles but also support through court rulings, jurisprudence, and administrative practices. Effective compensation mechanisms for victims of AI-induced harm depend fundamentally on the operationalization of such enforcement frameworks.

From a criminal law perspective, questions such as “who is the perpetrator?”, “how is fault determined?”, and “what are the limits of criminal capacity?” remain contested. Objective liability regimes based on norm violations—independent of personal fault—and enhanced supervision of legal entities could provide accountability for AI-related offenses. Similarly, in administrative law, the transfer of public authority to AI systems without ensuring the transparency, justifiability, and auditability of algorithmic decisions is unacceptable.

Country-specific examples of AI-related harms illustrate the diversity and developmental levels of legal systems’ responses to technology. The cases mentioned above demonstrate that in addressing damages caused by AI systems:

- Classical liability principles (negligence, fault, product liability) are being reinterpreted;
- National courts develop varying approaches concerning copyright, data mining, and fair use;
- The presence or absence of human contribution in AI outputs becomes a fundamental criterion.

In conclusion, the establishment of effective and accessible compensation mechanisms for damages caused by AI systems is essential not only for protecting victims’ rights but also for legitimizing the technology itself. In this regard:

- National legal systems should enact regulations providing special liability regimes for AI-related harms;
- Principles in international instruments should be incorporated into domestic law with effective judicial oversight mechanisms;
- Human-centered, accountable, and auditable AI applications should be promoted in administrative and criminal law.

This study offers a normative framework for how legal systems can respond to harms caused by AI and aims to contribute to legal roadmaps for a more just and accountable technology governance.

REFERENCES

American Bar Association. (2024). *BC tribunal confirms companies remain liable for information provided by AI chatbot.*

https://www.americanbar.org/groups/business_law/resources/business-law-today/2024-february/bc-tribunal-confirms-companies-remain-liable-information-provided-ai-chatbot/

AP News. (2025, Mart...). *French publishers and authors sue Meta over copyright works used in AI training.* AP News.

<https://www.apnews.com/article/168b32059e70d0509b0a6ac407f37e8a>

Caşın, M. H., Al, D., & Başkır, N. D. (2023). **The Problem of Criminal Liability Arising from the Actions of Artificial Intelligence and Robots.** *Journal of Law and Technology*, 12(3), 45–68.

Cerri, A. (2024, April 15). *Czech court finds that AI tool DALL-E cannot be the author of a copyright work.* IPKat. Retrieved from <https://ipkitten.blogspot.com/2024/04/czech-court-finds-that-ai-tool-dall-e.html>

Chollet, F. (2024). The Path to AGI: Neurosymbolic AI and Beyond. Retrieved from <https://time.com/7012823/francois-chollet/>

Hembt, S., Lutzhöft, N., & Bond, T. (2024, October 1). *Long-awaited German judgment by the District Court of Hamburg (Kneschke v. LAION) on the text and data mining exception(s).* Bird & Bird. Retrieved from <https://www.twobirds.com/en/insights/2024/germany/long-awaited-german-judgment-by-the-district-court-of-hamburg-kneschke-v-laion>

IBM Think. (2024). Understanding the Types of AI. Retrieved from <https://www.ibm.com/think/topics/artificialintelligence-types>

Kaur, A. (2025). A comprehensive analysis of types of artificial intelligence: Classification, applications, and future directions. **International Journal of Advanced Research in Computer and Communication Engineering**, 14(2), 167, <https://doi.org/10.17148/IJARCCE.2025.14221>

Leo, S. (2024). Data cleaning automation: Evaluating AI-based vs. rule-based approaches. **Journal of Data Engineering**, 10(3), 45-60. Retrieved from <https://www.researchgate.net/publication/389357157>

Lohmann, M. F. (2016). Liability issues concerning self-driving vehicles. *European Journal of Risk Regulation*, 7(2), 335–340. <https://doi.org/10.1017/S1867299X00005754>

Lumenalta. (2024). What are the different types of AI? Retrieved from <https://lumenalta.com/insights/what-are-the-different-types-of-ai>

Olabiyyi, W., Akinyele, D., & Joel, E. (2025). The evolution of AI: From rule-based systems to data-driven intelligence. **Journal of Artificial Intelligence Evolution**, 1(1), 1-20. Retrieved from <https://www.researchgate.net/publication/388035967>

Penrose R (1989) *The emperor's new mind. Concerning computers, minds, and the laws of physics.* Oxford University Press, Oxford

Reuters. (2025, 12 Mart). *French publishers and authors file lawsuit against Meta in AI case.* <https://www.reuters.com/technology/artificial-intelligence/french-publishers-authors-file-lawsuit-against-meta-ai-case-2025-03-12/>

Rietzler, M. (2022). *Who is responsible if an AI system gives a wrong diagnosis? Analysis of the EU liability law framework of medical AI* (Master's thesis, Lund University, Faculty of Law). Lund University Publications.

SearchEnterpriseAI. (2023). Choosing between a rule-based vs. machine learning system. **TechTarget**. Retrieved June 27, 2025, from <https://www.techtarget.com/searchenterpriseai/feature/How-to-choose-between-a-rules-based-vs-machine-learning-system>

Simplilearn. (2025). Types of Artificial Intelligence Explained. Retrieved from <https://www.simplilearn.com/tutorials/artificial-intelligence-tutorial/types-of-artificial-intelligence>

Susskind R, Susskind D (2015) *The future of the professions.* Oxford University Press, Oxford

Stephenson Harwood. (2025). *CJEU to rule on AI and copyright in a landmark case against Google.* Stephenson Harwood.

Tegmark M (2017) *Life 3.0. Being human in the age of artificial intelligence.* Alfred A. Knopf, New York

Xu, B. (2024, April 16). What is meant by AGI? On the definition of artificial general intelligence [Preprint]. arXiv. <https://doi.org/10.48550/arXiv.2404.10731>

Wachter, S., Mittelstadt, B., & Floridi, L. (2017). Why a right to explanation of automated decision-making does not exist in the General Data Protection Regulation. *International Data Privacy Law*, 7(2), 76–99.

William Fry. (2025, May 27). *CJEU to Rule on AI Chatbots and Copyright in Landmark Case Against Google*. William Fry