

LUISS Guido Carli University

Degree Program in Management and Computer Science

Course of: Business Cyberlaw

**The assessment obligations for use of personal data
and artificial intelligence in business enterprises**

Supervisor: Prof. Eugenio Prosperetti

Candidate: Antonio Cagnucci Vasquez

ID Number: 289871

Academic Year: 2024/2025

Contents

INTRODUCTION	3
CHAPTER 1: TECHNICAL FOUNDATIONS OF AI AND THE STRATEGIC ROLE OF PERSONAL DATA	5
1.1 ARTIFICIAL INTELLIGENCE IN CONTEMPORARY BUSINESS CONTEXTS	5
1.2 CORE AI TECHNOLOGIES RELEVANT FOR LEGAL AND BUSINESS ANALYSIS	8
1.3 PERSONAL DATA AS THE FOUNDATION OF CONTEMPORARY AI SYSTEMS	10
1.4 FROM TECHNICAL RISK TO SOCIETAL IMPACT: WHY IMPACT ASSESSMENTS BECOME NECESSARY	13
CHAPTER 2: AI IMPACT ASSESSMENTS AS A TOOL FOR RESPONSIBLE AI GOVERNANCE	18
2.1 FROM EX POST ENFORCEMENT TO EX ANTE GOVERNANCE: THE EVOLUTION OF EU DIGITAL REGULATION	18
2.2 IMPACT ASSESSMENTS AS INSTRUMENTS OF PREVENTIVE GOVERNANCE IN EU DIGITAL LAW	22
2.3 THE RISK-BASED REGULATORY ARCHITECTURE OF THE AI ACT	28
2.4 AI IMPACT ASSESSMENTS AS INSTRUMENTS OF EU REGULATORY GOVERNANCE	36
CHAPTER 3: COMPARATIVE PERSPECTIVES ON AI GOVERNANCE AND IMPACT ASSESSMENT MECHANISMS	41
3.1 AI GOVERNANCE IN THE UNITED STATES	42
3.2 CHINA: STATE-CENTRIC AND PREVENTIVE AI GOVERNANCE	45
3.3 OTHER ASIAN APPROACHES: JAPAN AND SINGAPORE (SOFT GOVERNANCE)	47
3.4 INTERNATIONAL AND TRANSNATIONAL AI GOVERNANCE FRAMEWORKS	49
3.5 COMPARATIVE SYNTHESIS: WHAT MAKES THE EU MODEL DISTINCT	50
3.6 CONCLUSION	52
CHAPTER 4: INNOVATION VS PROTECTION IN AI REGULATION	53
4.1 THE INNOVATION-PROTECTION TENSION	53
4.2 RISKS OF OVER-REGULATION	54
4.3 RISKS OF WEAK PROTECTION	56
4.4 FINDING A MIDDLE GROUND	57
4.5 IMPACT ASSESSMENTS AS AN ENABLER OF RESPONSIBLE INNOVATION	58
4.6 FINAL REFLECTIONS	59
BIBLIOGRAPHIC REFERENCES	60

Introduction

Artificial Intelligence systems are no longer used in experimental settings but are being already deployed in real-world every-day business activities and public sector contexts, including areas like recruitment, credit assessment, insurance, healthcare and public administration. As a result, any decisions related to AI will have a direct effect on individuals, influencing access to opportunities, services, and resources, also creating concerns that go beyond technical performance. In response to those developments the European Union has taken action by having a regulatory approach centred on risk management and the protection of fundamental rights, this led to the creation of the Artificial Intelligence Act¹. Within this framework, impact assessments play the role of a central mechanism through which legal requirements are applied in organizational practices, shaping how AI systems are designed, deployed and monitored. For organizations that are continuously operating in competitive and data-driven markets, impact assessment is not only a formal compliance obligation, but a process that directly influences innovation strategies, governance structures, and risk management choices. Understanding how impact assessments work in practice is essential to evaluate whether the AI act can effectively balance innovation, accountability, and the protection of individual rights.

While the adoption of AI has led to improvements in efficiency, it has also increased concerns related to opacity, bias, accountability, and the delegation of decision-making to automated systems. These risks create a challenge for the traditional regulatory approaches, which are designed to be applied after harm has already materialised rather than preventing it before it occurs. In this context, organizations or companies play an important role in this matter as they are responsible for designing, deploying and operating AI systems in real-world settings. As a result, decisions related to the design, deployment, and use of AI systems are not only technical, but involve strategic, organizational, and governance considerations. Consequently, a gap appears between the regulatory

¹ European Commission, *Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*, COM(2021) 206 final.

requirements and their implementation for organizational practices. This gap leads to a very important question: how can we turn and apply legal requirements into everyday organizational practices while still allowing innovation to develop?

The purpose of this thesis is to examine the role of assessment obligations in regulating the use of artificial intelligence and personal data within business enterprises. Rather than only focusing on technological aspects, the analysis will cover how these obligations have an impact in organizational decision-making, governance structures, and risk management practices. By taking into account different regulatory frameworks, with more focus in the European Union's approach, the thesis will assess how assessment mechanisms can connect the legal requirements and everyday business operations.

From a legal point of view, this thesis falls within the field of Business Cyber Law and focuses on how artificial intelligence is regulated when it is used by business enterprises. The analysis does not treat AI only as a technological development, but as a set of systems whose use triggers specific legal obligations for organisations. Based on that the thesis will focus on the duties companies have to assess and manage the risks related to the use of AI systems, especially where these systems affect individuals or rely on personal data. In this case, tools like the GDPR and the AI Act are not discussed as abstract legal concepts but more as legal frameworks that place specific duties on businesses and directly influence how AI systems are designed, used, and monitored.

This thesis is structured as follows. Chapter 1 introduces the technical foundations of artificial intelligence and examines the role of personal data in AI-driven business processes, with a view to identifying the legal and regulatory issues raised by their use in practice. Chapter 2 analyses the assessment obligations applicable to business enterprises, focusing on governance mechanisms and risk management requirements under the European regulatory framework. Chapter 3 adopts a comparative perspective, examining how different jurisdictions regulate artificial intelligence and the processing of personal data. Finally, Chapter 4 discusses the tension between regulatory compliance and innovation, reflecting on how assessment obligations may support a balanced and responsible use of AI in business contexts.

Chapter 1: Technical Foundations of AI and the Strategic Role of Personal Data

1.1 Artificial Intelligence in Contemporary Business Contexts

Artificial Intelligence (AI) systems are no longer used only in experimental or research environments; they are now integrated in our everyday business operations and public sector activities. Over the last 10 years, advances related to computational power, data availability, and algorithmic approaches created the bridge for organizations to integrate AI systems into decision-making activities that were traditionally performed by humans. As a result, AI has progressively evolved from a supporting analytical tool into a core infrastructural component determining how economic and administrative decisions are made. Once AI systems are deployed in real-world decision-making contexts, their use becomes legally relevant, as such systems may directly affect individuals and therefore fall within the scope of regulatory frameworks aimed at protecting fundamental rights and ensuring lawful decision-making.

In contemporary business contexts, AI systems are increasingly used to increase efficiency, reduce costs, and raise predictive accuracy. Companies rely on algorithm models to evaluate creditworthiness, personalize marketing strategies, automate recruitment screening, manage supply chains, and detect fraud. In that same way, public authorities use AI systems in areas such as welfare allocation, tax enforcement, risk assessment in law enforcement, and healthcare resource management. These applications show that AI is no longer a secondary tool but operates at the heart of organizational decision structures. The introduction of AI into these areas raises questions about how existing legal rules apply when decisions are automated or partially automated, and whether current legal safeguards remain adequate in contexts characterised by scale, opacity, and limited human intervention.

A key characteristic of this transformation lies in the shift from decision-support systems to automated or semi-automated decision-making systems.² At first, earlier technological tools mainly assisted human decisions by providing recommendations or forecasts, but now many modern AI systems are designed to autonomously generate decisions or considerably restrict human discretion. One case where we can see this transformation is in the recruitment processes, nowadays companies use algorithmic screening tools that can filter candidates before any human review takes place. Another example are the credit and insurance markets; automated scoring systems commonly determine access to financial services with limited or no individualized human involvement. This evolution raises fundamental questions regarding accountability, transparency, and the distribution of responsibility between human actors and technical systems. From a legal perspective, this evolution raises fundamental questions about accountability, transparency, and the allocation of responsibility, as traditional legal frameworks are mainly built on the assumption of meaningful human control over decisions that affect individuals.

From a business perspective, the attractiveness of AI-powered decision-making it's a fact. Automated systems can process huge volumes of data at speeds and scales unattainable for human decision-makers, enabling organizations to operate more competitively in markets where decisions rely on data. AI systems make it easier to achieve more consistency, scalability, and the ability to recognize complicated patterns that cannot be observed easily through traditional analytical methods. As a consequence, firms have strong economic incentives to start using AI technologies, especially in highly competitive sectors where marginal gains in efficiency or accuracy can generate numerous advantages.³ At the same time, these market-driven incentives can conflict with legal requirements related to transparency, fairness, and accountability, especially where the fast adoption of AI systems limits the ability of organizations to fully assess and manage legal risks in advance. From a legal standpoint, this scaling effect challenges regulatory approaches that are

2 F. PASQUALE, *The Black Box Society*, Cambridge (Mass.), 2015.

3 H. VARIAN, *Artificial Intelligence, Economics, and Industrial Organization*, in NBER Working Papers, 2019.

traditionally designed to address harm on an individual basis, raising questions about the effectiveness of existing enforcement and remedial mechanisms when risks and impacts become systemic rather than isolated.

However, the extensive implementation of AI systems also means that it will have greater effects and importance in society. Decisions once made on a case-by-case basis by human agents are now standardized and replicated across large populations. Errors, biases, or design choices present in algorithmic systems can therefore affect thousands or millions of individuals simultaneously. This scaling effect distinguishes AI-driven decision-making from traditional forms of organizational decision-making and increases the potential consequences of flawed or opaque systems.⁴ From a legal standpoint, this scaling effect challenges regulatory approaches that are traditionally designed to address harm on an individual basis, raising questions about the adequacy of existing enforcement and remedial mechanisms when risks and impacts become systemic rather than isolated.

Moreover, AI systems also operate in situations where individuals are often not aware that algorithmic decisions are being made about them or are not able to challenge those decisions. In many cases, people do not know how these decisions are reached, which data are used, or which elements affect the outcome. This information asymmetry reinforces existing power imbalances between organizations and individuals and complicates the exercise of fundamental rights such as non-discrimination, due process, and effective remedy.⁵ From a legal perspective, these concerns are directly addressed by European regulatory frameworks such as the GDPR and the AI Act, which aim to strengthen transparency, accountability, and procedural safeguards in automated decision-making contexts.

In this context, AI should not be understood only as a neutral technological tool but as a socio-technical system present in broader organizational, economic, and administrative frameworks. The effects of AI use depend not only on technical performance but also on design choices, governance structures, and the regulatory

⁴ S. BAROCAS – A. D. SELBST, *Big Data's Disparate Impact*, in *California Law Review*, 2016, p.

⁵ EUROPEAN COMMISSION, *White Paper on Artificial Intelligence – A European Approach to Excellence and Trust*, Brussels, 2020.

environment in which systems operate. As AI systems increasingly mediate access to opportunities, services, and resources, their integration into business practices transforms them into tools with significant effects for individuals and society at large, a reality that is directly reflected in European regulatory approaches such as the GDPR and the AI Act, which require organizations to take responsibility for the design, governance, and use of AI systems when fundamental rights may be affected.

This transformation creates the need to move beyond only efficiency-oriented evaluations of AI systems. While performance metrics such as accuracy or cost reduction remain relevant, they are not enough to capture all the consequences of AI-driven decision-making. Understanding AI in contemporary business contexts, therefore, requires attention to both its economic rationale and its capacity to generate risks and impacts that extend beyond organizational boundaries. This insight helps explain why regulatory frameworks, and in particular impact assessment mechanisms, now play a central role in the governance of artificial intelligence. For this reason, European regulatory frameworks, particularly the GDPR and the AI Act, increasingly rely on impact assessment mechanisms as tools to identify, assess, and manage risks before AI systems are deployed in practice.

1.2 Core AI Technologies Relevant for Legal and Business Analysis

Understanding how artificial intelligence systems work is important to evaluate their legal and social implications. While AI technologies include many different techniques, contemporary regulatory and governance debates are more focused on systems based on machine learning. These software systems are different from earlier ones in that they do not only operate according to pre-defined instructions but instead learn patterns and correlations from data to generate outputs, predictions, or decisions.⁶ This learning-based approach makes system behaviour less predictable and more difficult to fully specify in advance, which raises legal challenges for traditional regulatory models that rely on foreseeability, control, and clearly defined rules of operation.

⁶ S. RUSSELL – P. NORVIG, *Artificial Intelligence. A Modern Approach*, 4th ed., London, 2021.

Machine learning systems are classified into different categories depending on the learning method that they use. In supervised learning, models are trained on labelled datasets, where the desired output is known in advance, allowing the system to learn associations between input data and outcomes. On the other hand, unsupervised learning systems work differently, they identify patterns or structures within unlabelled data without predefined targets. Another category is reinforcement learning, which relies on feedback mechanisms through which systems learn by trial and error, improving their behaviour according to predefined objectives. Although these differences are mainly technical, they also have importance from a legal perspective because they influence predictability, controllability, and the degree of human control over system behaviour.⁷ In legal terms, these differences affect how responsibility can be assigned, how decisions can be reviewed, and whether organizations are able to meet obligations related to accountability and oversight when AI systems are used in decision-making processes.

Many modern AI systems rely on complex models, which makes their decision-making difficult to explain. As model complexity increases, it becomes more difficult to provide clear explanations for how specific outputs are produced. This “black box” problem creates important challenges for accountability, transparency, and legal scrutiny, especially in situations where AI systems affect individual rights or have legal consequences. This makes it difficult to understand how decisions are taken, which also complicates both internal management and external review.⁸ From a regulatory standpoint, this lack of explainability can undermine legal obligations related to transparency, the ability to provide reasons for decisions, and the possibility for affected individuals to effectively challenge automated outcomes.

Another characteristic of machine learning systems with legal relevance is their dependence on data quality and context. AI systems do not operate independently, they are the reflection of the data on which they are trained, including its biases and limitations. As a result, errors present in training data can be reproduced and amplified at scale. In addition, AI systems usually adapt over time as they are

⁷ OECD, *Artificial Intelligence in Society*, Paris, 2019.

⁸ F. PASQUALE, *The Black Box Society*, Cambridge (Mass.), 2015.

presented with new data, meaning that system behaviour may change after implementation. This challenges traditional regulatory approaches that assume stable and predictable system performance.⁹ In legal terms, this raises difficulties for compliance models based on one-time assessments, as organizations may struggle to ensure that evolving systems continue to meet legal requirements throughout their operational lifecycle.

From a business perspective, these technical characteristics are the reason for the attractiveness of AI systems but at the same time they generate governance risks. The capacity to process large volumes of data, identify complex correlations, and automate decision-making processes makes it possible to improve efficiency and scalability. At the same time, opacity, adaptability, and data dependence reduce the ability of organizations to fully anticipate system outcomes or assess later impacts in advance. This conflict between innovation potential and risk exposure is key to contemporary AI governance debates and explains why there is increasing attention to tools such as impact assessments, which aim to manage uncertainty rather than eliminate it completely.¹⁰ Within the European regulatory framework, this approach is reflected in the choice to rely on procedural obligations rather than outright prohibitions. Instruments such as the GDPR and the AI Act require organizations to identify, assess, and manage risks linked to AI systems as part of their internal governance, recognizing that uncertainty is an inherent feature of these technologies rather than an exception.

1.3 Personal Data as the Foundation of Contemporary AI Systems

The operation and effectiveness of many contemporary artificial intelligence systems depend on the availability and use of large volumes of data. In business and public sector contexts, this data usually contains information relating to identified or identifiable individuals. Therefore, personal data plays an important role in the development, training, and operation of AI systems, which makes the

⁹ S. BAROCAS – M. HARDT – A. NARAYANAN, *Fairness and Machine Learning*, Cambridge (Mass.), 2019.

¹⁰ EUROPEAN COMMISSION, *White Paper on Artificial Intelligence – A European Approach to Excellence and Trust*, Brussels, 2020.

way data is collected and used a key factor in determining both system performance and social impact.¹¹ As a result, the use of personal data in AI systems brings these practices within the scope of data protection law, making data-related design and governance choices legally relevant rather than purely technical or operational decisions.

Machine learning systems are based on historical data to identify patterns and generate predictions or decisions. When this data includes personal information, the characteristics of the dataset will have an influence in the system behaviour and outcomes. The scale, variety, and level of detail of personal data available to organizations have been increasing in recent years, this led to the creation of more sophisticated and improved versions of automation and personalization. At the same time, this reliance on personal data also generates an increase in the potential for privacy risks, discriminatory outcomes, and misuse of information, especially when the data is collected or processed without enough compliance controls.¹² In legal terms, these differences affect how responsibility can be assigned, how decisions can be reviewed, and whether organizations are able to meet obligations related to accountability and oversight when AI systems are used in decision-making processes.

In regulatory terms, the role of personal data plays in the use of AI systems creates issues that are not limited to traditional data protection concerns. Decisions made with the use of Artificial Intelligence usually include the use of multiple data sources, the continuous reuse of data over time, and the generation of new information through automated analysis. Even when individual data points appear neutral or limited in nature, their aggregation and processing can produce insights that can generate damage or affect individuals' access to employment, credit, insurance, public benefits, or other essential services. This challenges conventional assumptions about transparency, purpose limitation, and individual control over personal information¹³, which are core principles of European data protection law

¹¹ OECD, *Artificial Intelligence in Society*, Paris, 2019.

¹² EUROPEAN COMMISSION, *White Paper on Artificial Intelligence – A European Approach to Excellence and Trust*, Brussels, 2020.

¹³ S. WACHTER – B. MITTELSTADT, *A Right to Reasonable Inferences*, in *Columbia Business Law Review*, 2019.

and become increasingly difficult for organizations to fully respect when data-driven AI systems operate at scale.

An important factor of this process is the distinction between data that individuals provide directly and data that is inferred or derived through the analysis of algorithms. AI systems frequently generate predictions, profiles, or risk scores that are not explicitly provided by the individual but are inferred from observed behaviour or correlations within large datasets. These inferred data points can have a huge influence in decision-making processes, even though individuals are often unaware of them. As a result, individuals may be affected by conclusions drawn about them without being able to understand how those conclusions were reached or which data contributed to them. From a legal perspective, this raises questions about how existing data protection rules apply to inferred data. When decisions are based on profiles or predictions rather than information directly provided by individuals, it becomes harder to ensure transparency and meaningful control over personal data. This creates uncertainty as to whether current legal safeguards are sufficient to protect individuals against decisions driven by large-scale inference and profiling by AI systems.

This situation generates concerns in contexts where AI systems are used to make or inform decisions with significant consequences. In those situations, personal data is not only an input used to improve system accuracy, but a central element with huge influence in the outcomes produced by the system. The opacity of inference processes and the complexity of data flows can make it challenging for individuals to exercise control over their information or to contest decisions that negatively affect them.

In addition, the extensive use of personal data in AI systems contributes to power imbalances between organizations that use these technologies and the individuals affected by them. Organizations typically possess more technical knowledge and expertise, control over data infrastructures, and the capacity to generate value from data at scale. Individuals, in the other hand, usually lack awareness of how their data is processed, how long it is retained, or how it is reused across different systems and purposes. These imbalances make it harder to achieve the practical exercise of rights related to transparency, fairness, and accountability, especially in large-scale,

automated decision-making environments where individuals have limited ability to influence or challenge how decisions are made.¹⁴

In this context, personal data should not be viewed only as a resource that can help to create more innovation or improve system performance. Instead, it must be seen and understood as a core element influencing the risks, vulnerabilities, and societal impacts associated with AI systems. Decisions about data collection, data quality, and data governance have a direct influence in how AI systems behave in practice and how their effects are distributed across different groups within society. For this reason, decisions about how personal data is collected and managed are not only technical or business choices, but also decisions that carry legal responsibility, particularly where data practices can lead to unequal treatment or harm to individuals.

Understanding the role that personal data plays throughout the AI lifecycle is therefore essential for being able to analyse and assess the implications of AI deployment. This includes not only the initial collection and training stages, but also ongoing use, adaptation, and reuse of data over time. This perspective helps us understand why contemporary AI governance frameworks continue to connect data protection principles with risk-based regulatory approaches. It clarifies the growing importance of mechanisms such as impact assessments, which their goal is to identify and reduce potential harms linked to data-powered AI systems before they are implemented in real-world contexts.¹⁵

1.4 From Technical Risk to Societal Impact: Why Impact Assessments Become Necessary

Building on the role that personal data plays in the previous section, it becomes clear that the risks created by artificial intelligence systems cannot be understood only from a technical or data protection perspective. When AI systems are

¹⁴ EUROPEAN UNION AGENCY FOR FUNDAMENTAL RIGHTS, *Getting the Future Right – Artificial Intelligence and Fundamental Rights*, Vienna, 2020.

¹⁵ EUROPEAN COMMISSION, *Proposal for a Regulation Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)*, Brussels, 2021.

introduced into real world decision making situations, their effects are not only about questions of model accuracy, efficiency, or formal compliance with legal requirements. The risks come from the interaction between technical design choices, data practices, organizational incentives, and the specific social environments in which these systems operate. It is this combination of elements that produces forms of risk that are often difficult to fit within traditional legal categories, which were mainly developed to address isolated violations rather than complex and system effects.

One of the main features of AI systems is their capacity to operate at scale. Decisions that before were made individually by human actors are now being more automated and replicated across large populations. As a result, even minor design choices or problems in the data can lead to widespread and systematic effects. As opposed to traditional decision-making processes, where errors could remain isolated or context-specific, AI systems can reproduce the same outcome repeatedly, and this can affect a large number of individuals in similar ways. This capacity for replication transforms technical limitations into potential sources of social harm.

The impact of AI systems on society becomes especially significant when they are used in areas that directly affect people's life opportunities and access to essential services. When AI systems are used in areas such as employment, credit, insurance, healthcare, or public administration, their outputs can have an influence in the individuals' economic prospects, social inclusion, or access to fundamental rights. In such settings, the consequences of automated decisions are not only operational but also normative, since they influence social outcomes and power relations. This makes it very difficult to treat AI systems as neutral tools whose effects can be evaluated only through performance metrics.

Another important challenge is that harm can be difficult to detect and address once AI systems are already operating. Many risks caused by AI systems do not appear right away or in ways that are easy to observe. Instead, they can emerge gradually, through cumulative effects, indirect consequences, or feedback loops that reinforce existing inequalities. For example, biased data or flawed assumptions can lead to decisions that seem reasonable in individual cases but over time they can

create discriminatory patterns. These characteristics challenge regulatory approaches that rely primarily on reactive enforcement or individual complaints. As a result, legal mechanisms that depend on individuals identifying harm and seeking remedies after decisions are made often struggle to respond effectively to these types of systemic and slowly emerging risks.

From a governance point of view, these dynamics reveal the limits of regulatory approaches that react only after harm occurs. Legal frameworks that focus on addressing harm only after it has occurred can struggle to provide effective protection in settings where decisions are made automatically at scale. This is because many existing legal regimes are built around identifying individual wrongdoing or discrete instances of harm, an approach that fits poorly with automated decision-making systems that produce widespread and cumulative effects. By the time adverse effects become visible, they may already part of everyday organizational practices or have affected a large number of individuals. Moreover, the complexity of AI systems and the opacity of their decision-making processes can make it challenging to establish responsibility or to fix the harm once it has occurred.

These challenges show the need for regulatory tools that can anticipate and manage risks before AI systems are introduced or used at scale. Instead of trying to eliminate all uncertainty, these tools focus on creating clear processes to identify, assess, and address potential harms at an early stage. More broadly, technology governance is moving away from a focus on outcomes alone and more toward better emphasis on process, accountability, and continuous oversight.

For those reasons, impact assessments have gained more attention as an important tool in the governance of AI systems. Impact assessments require organizations to think carefully about how an AI system will operate, what risks it may generate, and who may be affected by its use. By requiring a prior evaluation of potential impacts, these mechanisms encourage those responsible for the AI-systems to consider not only technical aspects but also legal, ethical, and societal implications. Instead, impact assessments are not meant to block innovation, but to guide it in line with public values.

Impact assessments reduce the information and power imbalances that often occur in AI-driven environments. By requiring organizations to evaluate risks and explain how decisions are made, they make AI systems more transparent and open to review by regulators, stakeholders, and the people affected. At the same time, they provide organizations with a framework for internal evaluation and risk management, making them include legal and ethical considerations into technical and business decisions.

The growing importance for the use of impact assessments shows that responsibility in AI governance is not only compliance with isolated legal rules. Instead, responsibility must be done as an ongoing process during all the lifecycle of an AI system, from design and data collection to use and monitoring. Impact assessments support this approach by integrating risk analysis into the practices of the organizations and by doing regular reviews as systems change or are used in new settings.

This approach is especially relevant for the fact that AI systems can change very quickly. As discussed before, many AI models can change over time, use new data, or are reused for purposes that were not the initial ones. These characteristics demonstrate that having static regulatory solutions is not enough and that there is the need for flexible governance mechanisms capable of neutralizing the evolving risks. Impact assessments offer a way to adjust to the flexibility that is needed while maintaining a clear focus on accountability and protection of fundamental interests.

In the end, impact assessments matter because they contribute to the connection between technological innovation and social responsibility. By turning abstract concerns about risk, fairness, and transparency into defined analytical steps, impact assessments help ensure that AI systems are evaluated not only for what they can do but for what they should do in each social and legal context. In the European context, this logic is reflected in legal instruments such as the GDPR and the AI Act, which require organizations to assess risks in advance and to justify how AI systems are designed and used when they affect individuals. This analysis and perspective form the basis of the regulatory approaches that are going to be explained and examined in the following chapter, where impact assessments are

analysed as a key component of contemporary AI governance frameworks, particularly within the European Union.

Chapter 2: AI Impact Assessments as a Tool for Responsible AI Governance

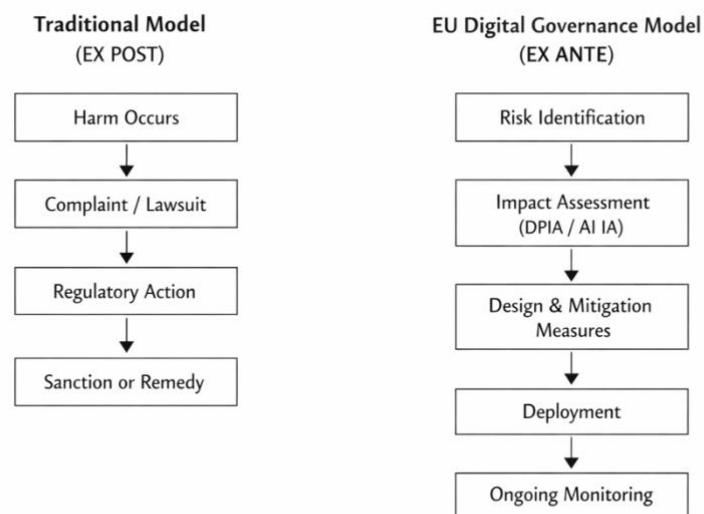


Figure 1 – From ex post enforcement to ex ante risk-based governance in EU digital regulation

2.1 From Ex Post Enforcement to Ex Ante Governance: The Evolution of EU Digital Regulation

European digital regulation has continuously moving away from relying on ex post enforcement and has instead placed more attention on ex ante governance. In simple terms, ex post enforcement means that the law intervenes only after harm or illegal conduct has already happened, for example through fines, lawsuits, or court decisions. In the other hand, ex ante governance focuses on reducing risks before damage occurs by requiring organizations to think ahead and manage potential problems before they happen. This change is based on the idea that traditional enforcement methods often do not work well when applied to digital systems that are complex and change very quickly.¹⁶

Under classical enforcement models, regulators and courts intervene only once a violation has taken place. A company may be sanctioned after a data breach, held liable following a discriminatory algorithmic decision, or fined once illegal content

¹⁶ Spark Legal & Policy Consulting, Study on Systemic Risks and their Mitigation under the DSA, 2023.

has already circulated online. These mechanisms are based on the assumption that the threat of sanctions will deter harmful behavior and that individual cases can be resolved one at a time. While this approach may work in more stable regulatory contexts, it struggles to address the realities of today's data-driven and algorithmically mediated environments.¹⁷

Digital harms often do not come from a single event but develop over time. Instead of one clear incident, people can slowly lose the control over their personal data, affected by biased automated decisions, or exposed to harmful content across many users. In these situations, a single fine or court case is rarely enough to stop the problem. Enforcement actions also tend to arrive late. By the time public authorities step in, technologies and business practices may already have changed or caused effects that are hard to reverse.¹⁸ As a result, regulators are frequently forced to react only after the harm has already taken place.

These limitations are evident in the field of data protection. Under a more reactive enforcement model, individuals harmed by misuse of personal data would need to try to get compensation only after their data has been exposed or misused. In practice, many privacy harms are intangible, difficult to quantify, or spread across large groups of people, making individual litigation ineffective. As a result, legal scholars have argued that certain digital harms are simply “too serious to be remedied after the fact,” which shows that tort-based liability is not well suited to deal with these problems in addressing systemic data-driven risks.⁵ Even where lawsuits are possible, they tend to address narrow individual cases rather than the underlying organizational practices that generate risk.¹⁹

Because of these limits, EU lawmakers have started to focus more on prevention in digital regulation. Instead of acting only after damage has already happened, recent EU laws require companies to look for risks early and deal with them before new technologies or services are used. This approach can be seen across many EU

¹⁷ Proposal for a Regulation laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act), COM(2021) 206 final, Arts. 8–15.

¹⁸ Palumbo, A., “Systemic risk management as an emerging regulatory approach in EU digital legislation,” *Technology and Regulation* (2026).

¹⁹ Bietti, E., “Can Private Law Protect Privacy in Today's Economy?”, *Balkanization*, 10 Dec. 2024.

digital laws and shows a clear shift from reacting to problems after the fact to trying to prevent them earlier toward preventing problems in advance.²⁰

This change does not mean that the EU has stopped using traditional enforcement tools. Fines, liability claims, and sanctions are still available when rules are broken. What has changed is the main focus. Instead of relying mostly on punishment after problems occur, EU law now expects organizations to build compliance and risk controls into how they operate from the start. The idea is that many digital risks can be handled better inside organizations, rather than trying to fix the damage only after it has already happened.

The GDPR is one of the first EU laws to adopt this preventive approach to digital regulation. When it entered into force in 2016, it changed how data protection rules were enforced by putting more responsibility on organizations. Companies are not only expected to follow the rules, they also to be able to show that they are doing so. One of the main tools used for this is the Data Protection Impact Assessment (DPIA). Article 35 of the GDPR requires organizations to carry out a DPIA before starting data processing activities that are likely to generate high risks to people's rights and freedoms.²¹

The DPIA obligation alters the timing and allocation of regulatory responsibility. Instead of regulators identifying violations after they occur, organizations are required to assess risks themselves in advance and to integrate safeguards into the design of their systems. These safeguards could include data minimization, bias mitigation, or improved security measures. As scholars have noted, DPIAs work as a preventive tool because they require organizations to look for risks early and take steps to avoid harm before a system is put into use. Alongside DPIAs, the GDPR introduced additional preventive mechanisms such as data protection by design and by default, as well as mandatory appointment of Data Protection Officers in certain cases.²²

The same idea can be seen in the Digital Services Act (DSA), which applies a preventive approach to online platforms and intermediaries. The DSA was

²⁰ Regulation (EU) 2022/2065 (Digital Services Act), Art. 34

²¹ Regulation (EU) 2016/679, Art. 35.

²² Regulation (EU) 2016/679, Arts. 25 and 37.

introduced because lawmakers realized that reacting only after problems appear does not work well for large platforms. For this reason, the DSA requires very large online platforms and search engines to regularly assess the risks linked to how their services operate. These assessments must cover issues such as illegal content, impacts on fundamental rights, public security, and risks related to elections.²³

Articles 34 and 35 of the DSA require designated platforms to not only identify these risks but also implement reasonable and effective mitigation measures. This includes adjusting recommendation systems, modifying content moderation practices, or redesigning platform features that contribute to harm.²⁴ The emphasis is therefore not on punishment after failure, but on continuous monitoring and proactive system design. Transparency and auditing obligations further make sure that these preventive efforts remain subject to regulatory oversight.

The proposed Artificial Intelligence Act is one of the best examples of the EU's move toward prevention in digital regulation. Instead of dealing with problems only after harm has already happened, the AI Act sets rules that apply before AI systems are used. It follows a risk-based approach, meaning that AI systems are treated differently depending on how much harm they could cause. Practices considered too dangerous are banned altogether, while high-risk AI systems must meet strict legal requirements before they can be placed on the market.²⁵

For high-risk AI systems, providers and deployers must establish risk management systems, make sure there is appropriate data governance, maintain technical documentation, implement human oversight mechanisms, and complete conformity assessments before deployment.¹³ This structure is similar to pre-market approval regimes found in other regulated sectors, like pharmaceuticals or product safety, and reflects a conscious decision to prioritize prevention over remediation. Political debates about the AI Act further confirm this orientation, as EU institutions supported preventive governance tools over expanding civil liability regimes for AI-related harm.

²³ Regulation (EU) 2022/2065, Arts. 34–35.

²⁴ Spark Legal & Policy Consulting, Study on Systemic Risks and their Mitigation under the DSA, 2023.

²⁵ Proposal for a Regulation laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act), COM(2021) 206 final, Arts. 8–15.

Taken together, these developments illustrate a broader transformation in EU digital regulation. Across data protection, online platforms, and artificial intelligence, the EU has increasingly relied on procedural obligations, internal risk assessments, and ongoing monitoring as core regulatory instruments. This approach reflects the recognition that complex, opaque, and rapidly evolving digital systems cannot be effectively governed through reactive enforcement alone. Instead, embedding accountability, foresight, and prevention into organizational decision-making has become a central feature of the EU's digital regulatory model.²⁶

2.2 Impact assessments as instruments of preventive governance in EU digital law

Impact assessments as preventive regulatory tools

Impact assessments are increasingly used in EU digital regulation as tools to deal with risk before harm occurs. Instead of reacting to violations after they have already taken place, they require organizations to think in advance about the possible negative effects of their activities. In regulatory terms, an impact assessment is a structured process through which risks are identified, evaluated, and addressed before a system or practice is put into use.²⁷ This approach shifts attention away from damage control and toward anticipation.

The growing reliance on impact assessments reflects a broader change in how the EU regulates complex technologies that are moving in a very fast pace. Traditional enforcement models are largely reactive. They rely on sanctions, liability, or court decisions once something has gone wrong. While these tools remain important, they often intervene too late to prevent harm in digital environments. Impact assessments work differently. Their main purpose is to prevent harm by forcing organizations to reflect on risks early and to adjust their

²⁶ Montagnani, M. et al., "The EU regulatory approach(es) to AI liability," *Computer Law & Security Review*, 2024.

²⁷ European Commission, Better Regulation Guidelines, SWD(2021) 305 final.

behaviour accordingly. The focus is therefore not on punishment, but on anticipation and prevention.²⁸

This preventive logic is very important in the digital field. Digital systems have the ability to operate at scale, affect large numbers of people, and improve very fast over time. Once harm has been made, its very difficult or impossible to reverse. In this context, ex post enforcement alone may be insufficient to protect individuals and society. Impact assessments respond to this problem by embedding risk analysis into the decision-making process before deployment.

A key feature of impact assessments is that they shift responsibility inside organizations. Instead of regulators identifying risks from the outside, companies themselves are required to examine how their systems work and what consequences they could produce. This internalization of responsibility is central to the EU's preventive regulatory strategy. By design, impact assessments make organizations accountable for understanding and managing the risks they create.²⁹ They require organizations to engage with the consequences of their choices rather than treating compliance as an external constraint.

Impact assessments also work as governance tools rather than simple compliance checklists. They require documentation, internal coordination, and structured decision-making. Legal, technical, and managerial actors need work together as a team to find the potential risks and decide how they will tackle them. In this way, impact assessments influence how organizations design systems, distribute the responsibilities, and justify decisions internally.⁴ Their impact of these assessments is not only about formal compliance with legal rules, but also in shaping organizational behaviour more broadly.

In practice, impact assessments also change how organizations make decisions. Before introducing a new system, someone has to stop and think about what could go wrong. This step is important because many problems appear once users complain or authorities step in. Impact assessments bring those problems forward in time. They interrupt decisions that are made too quickly or only for efficiency reasons. Instead of asking only “does this work?”, organizations are pushed to ask

²⁸ G. Majone, *Regulating Europe*, Routledge, 1996.

²⁹ L. Hood et al., *The Government of Risk*, Oxford University Press, 2001.

“is this acceptable?”. Innovation is not blocked, but it no longer moves on autopilot. It moves with more awareness of its effects.

The Data Protection Impact Assessment under the GDPR

The most developed example of an impact assessment in EU digital law is the Data Protection Impact Assessment (DPIA) under the GDPR. The DPIA works as a reference model for later regulatory tools and provides a good illustration of how impact assessments operate in practice. Its main purpose is to prevent risks to individuals before personal data processing activities begin.³⁰

Article 35 GDPR requires controllers to carry out a DPIA when a planned processing activity is expected to have significant risks to individuals’ rights and freedoms. This requirement mainly applies to processing that uses new technologies, involves large-scale profiling, or includes ongoing monitoring of individuals. The assessment must be conducted before the processing takes place, not after harm has already occurred.³¹ This timing requirement reflects the preventive logic of the GDPR.

The DPIA is not intended to be a very formal exercise. Its main function is to make sure that data protection risks are identified early and addressed through appropriate safeguards. These safeguards include technical measures, such as stronger security or data minimization, as well as organizational measures, like changes to internal procedures or limits on how data is used. Through this mechanism, the GDPR seeks to integrate data protection directly into the design of systems and business practices.³²

The DPIA also has an important role in operationalizing the accountability principle. This means that controllers are not only expected to follow the rules, but also to show that they have analysed the risks and taken the correct steps to address them. Documenting the assessment or involving supervisory authorities helps make this clear. DPIAs therefore operate both as preventive tools and as evidence of

³⁰ Regulation (EU) 2016/679 (GDPR), Recital 84.

³¹ GDPR, Article 35(1)–(3).

³² GDPR, Article 35(7).

compliance. They create a traceable record of risk-related decisions that can later be reviewed by regulators.³³

Over time, the DPIA has changed how responsible data use is understood in the EU. Today, many organizations carry out a DPIA even when the law does not strictly require it, because it is commonly viewed as good practice. Authorities, courts, and business partners often expect organizations to be able to show that risks were considered in advance. In this sense, the DPIA has had an influence beyond its formal legal obligation. It has helped normalize the idea that risk assessment is part of ordinary organizational responsibility rather than an exceptional burden imposed by regulators.

Strengths and limits of DPIAs as governance instruments

DPIAs are often regarded as one of the GDPR's most effective preventive mechanisms. By making organizations think about risks before they start processing data, DPIAs help prevent problems that would be hard to fix later. They also push organizations to take data protection seriously from the beginning, instead of dealing with it only after something goes wrong.³⁴ This often results in systems that are planned more carefully and data practices that are more responsible.

At the same time, experience with DPIAs has showed many limitations. One really discussed concern is that they may become box-ticking exercises. This risk arises when organizations focus on formally completing the assessment rather than engaging in meaningful analysis. In such cases, the DPIA exists on paper but has little influence on actual decision-making.³⁵ The effectiveness of DPIAs therefore depends mainly on how seriously organizations engage with the process.

Another problem is that the quality of DPIAs is not consistent. The GDPR gives controllers a lot of freedom in how they carry out these assessments, so results can differ a lot in practice. Smaller organizations may not have enough expertise or staff to do a thorough assessment. Larger organizations, on the other hand, may handle

³³ GDPR, Articles 5(2) and 24; Recital 85.

³⁴ A. Mantelero, "The Future of Consumer Data Protection in the EU", *Computer Law & Security Review*, 2014.

³⁵ EDPB, Guidelines on Data Protection Impact Assessment (DPIA), WP248 rev.01.

DPIAs as standard paperwork, often with little real input from technical teams.³⁶ This makes it harder to ensure that DPIAs are applied in a consistent and effective way across sectors and Member States.

These limits do not mean that DPIAs are useless. Instead, they show the difficulties of relying on this kind of regulatory approach. Impact assessments depend a lot on how organizations work internally, how decisions are made, and how much judgment is involved. Their success therefore depends not only on legal rules, but also on guidance, oversight, and enforcement by supervisory authorities. Without sufficient external supervision, there is a risk that preventive tools lose their practical impact. (me salt eel footnote 12)

Another limitation of DPIAs is that it is hard to see what effect they have in practice. Since they are carried out inside organizations, it is often unclear whether a DPIA actually helped prevent harm or whether it only recorded risks without leading to real changes. This makes it difficult for regulators to assess how effective DPIAs really are. Even so, DPIAs still serve a purpose. They force organizations to leave a written record of how risks were handled, which makes it harder to deny responsibility and helps authorities reconstruct decisions if issues arise later.

From DPIAs to AI impact assessments

While DPIAs play an important role in EU data protection law, they are not enough to address all the risks related with artificial intelligence systems. Many AI-harms go beyond the processing of personal data and cannot be fully tackled through a data protection focus alone.³⁷ This limitation becomes increasingly visible as AI systems are deployed in more complex and sensitive contexts.

AI systems can still affect people in important ways even when they do not use much personal data. Tools like automated decisions, scoring systems, or recommendations can shape outcomes and lead to unfair or unclear results, or

³⁶ F. Zuiderveen Borgesius, “Strengthening Legal Protection Against Discrimination by Algorithms”, *International Journal of Human Rights*, 2020.

³⁷ L. Floridi et al., “AI4People—An Ethical Framework for a Good AI Society”, *Minds and Machines*, 2018.

reduce human involvement. These issues are not always covered by data protection rules. Because of this the DPIAs can only deal with some of the risks linked to AI. Their main focus is on how data is used, not on how decisions are made or how they affect people in real situations.³⁸

This gap is especially evident in areas such as employment, credit, education, and public administration. In these scenarios, the main problem is often not the legality of data processing, but the effects of automated decision-making on individuals' lives. Traditional DPIAs are not designed to fully assess these impacts, especially when harm comes from decision logic, model behavior, or system outputs rather than from unlawful data use.³⁹

For these reasons, the EU has moved toward developing impact assessment obligations that focus on more things than data protection. AI impact assessments build on the DPIA model but expand its scope to cover a larger number of risks and rights. This evolution prepares the ground for the AI Act, which introduces a more comprehensive approach to risk assessment made specifically to artificial intelligence systems. The structure and role of these AI impact assessments are examined in detail in the following section.⁴⁰

The growing use of AI has therefore exposed a structural mismatch between the scope of DPIAs and the reality of algorithmic decision-making. While data protection remains important, many of the most serious concerns raised by AI relate to power, control, and social impact rather than data misuse alone. This realization has pushed EU policymakers to rethink how impact assessments should operate when technologies influence behavior, opportunities, and rights in more indirect ways. As a result, the move toward AI-specific impact assessments can be seen as an evolution rather than a break from existing EU regulatory practice.

³⁸ S. Wachter, B. Mittelstadt, "A Right to Reasonable Inferences", *Columbia Business Law Review*, 2019.

³⁹ Council of Europe, *Algorithms and Human Rights*, 2019.

⁴⁰ European Commission, Proposal for an Artificial Intelligence Act, COM(2021) 206 final.

2.3 The risk-based regulatory architecture of the AI Act

The Artificial Intelligence Act regulates AI systems according to their level of risk. Rather than relying only on sanctions after harm has occurred, it requires organizations to assess and manage risks in advance through defined procedural obligations.⁴¹ Instead of applying the same legal obligations to all AI systems, the Regulation applies different levels of regulation depending on how an AI system is used and the risks it may cause to individuals and society.⁴² The risk-based structure of the AI Act is not accidental, but determines when and how the compliance obligations apply. It shapes the circumstances in which organizations are required to carry out AI impact assessments as part of their governance framework⁴¹ This risk-based structure is legally operationalised through Articles 5 to 7 of the AI Act, which determine whether an AI system is prohibited, classified as high-risk, or subject only to limited transparency obligations.

The AI Act organises AI systems into different categories depending on the risks they pose. It distinguishes between practices that are outright prohibited, systems classified as high risk and therefore subject to strict obligations, lower-risk systems mainly governed by transparency requirements, and uses of AI that fall largely outside the scope of binding regulation.⁴² This classification has legal consequences. It determines whether an AI system may be placed on the market at all and, if so, under which conditions. More specifically, classification determines whether providers and deployers are subject to mandatory obligations under Chapter III of the AI Act, including those set out in Articles 9 to 15. Once an AI system is classified as high-risk, the provider and deployer are subject to a series of ex ante obligations, including the establishment of risk management processes, data governance safeguards, technical documentation requirements, mechanisms for human oversight, and post-market monitoring duties.⁴³

⁴¹ European Commission, Proposal for a Regulation laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act), COM(2021) 206 final, 2021, Explanatory Memorandum.

⁴² REGULATION (EU) 2024/1689, Artificial Intelligence Act, arts. 5–7.

⁴³ Artificial Intelligence Act, arts. 9–15.

This regulatory approach shows the understanding that AI systems can produce harms that are difficult to identify in advance, may accumulate over time, and can affect individuals and groups in a systemic way.⁴⁴ In areas such as recruitment, access to credit, biometric identification, or the delivery of essential public services, decisions supported or taken by AI systems can have a direct and lasting impact on individuals' prospects. At the same time, these decisions are often difficult to detect, difficult to understand, and even harder for affected persons to challenge in practice. Once harm has occurred, traditional enforcement tools such as judicial remedies or administrative sanctions are frequently too slow or too limited to offer effective protection. This limitation explains the AI Act's regulatory choice to rely on ex ante obligations that apply before an AI system is deployed, rather than depending primarily on remedies after harm has already occurred. For this reason, the AI Act places emphasis on preventive regulation seeking to address risks before they materialise through structured compliance obligations applied before deployment.⁴⁵

This regulatory approach builds on the risk-based logic already familiar from the GDPR. Under the GDPR, the extent of a controller's obligations depends on the likelihood and seriousness of risks to the rights and freedoms of individuals, most clearly reflected in the obligation to carry out Data Protection Impact Assessments.⁴⁶ The AI Act follows a comparable structure by linking stronger compliance duties to the classification of an AI system as high-risk. In doing so, the AI Act builds on the GDPR's risk-based architecture, while adapting it to address harms that are not limited to personal data processing. In this case, organisations are required to identify, document, and address risks before the system is deployed. At the same time, the range of issues examined under the AI Act is broader. Unlike GDPR impact assessments, which focus mainly on data protection and privacy, AI impact assessments are intended to capture a wider range of concerns, including

⁴⁴ M. E. KAMINSKI – G. MALGIERI, Algorithmic Impact Assessments under the GDPR, in *International Data Privacy Law*, 11(2), 2021, p. 125 ss.

⁴⁵ REGULATION (EU) 2016/679 (General Data Protection Regulation), Recitals 84 and 90; European Commission, Proposal for a Regulation laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act), COM(2021) 206 final, cit.

⁴⁶ REGULATION (EU) 2016/679, art. 35.

discrimination, transparency, and the protection of fundamental rights more generally.⁴⁷

The AI Act relies heavily on procedural safeguards because traditional legal rules are often not enough to regulate how AI systems work in practice. Many AI systems are difficult to understand, and their effects cannot always be predicted in advance. As a result, the legislator has chosen an approach that focuses on clear documentation, internal responsibility, and ongoing supervision. These procedural duties are not optional best practices, but form part of the binding compliance framework established by the AI Act. Within this framework, AI impact assessments play an important role. They provide the main tool through which legal requirements are applied in practice, guiding how AI systems are designed, put into use, and monitored over time.⁴⁸

Scope and objectives of AI impact assessments

AI impact assessments under the AI Act are not only an extension of existing DPIA obligations. Although they draw on the GDPR model, they are meant to address how AI systems operate in practice. They respond to the use of automated decision-making, the ability of AI systems to operate at scale, and the fact that important assessments are often carried out by algorithms rather than humans.⁴⁹

Unlike DPIAs, which are required when personal data processing itself creates a high risk, AI impact assessments depend on how an AI system is used and what role it plays in decision-making. This difference matters because an AI system can cause serious effects even when personal data is processed only to a limited extent, or indirectly. Some examples of this are systems used for automated ranking, scoring, or decision support in areas such as access to jobs, credit, education, or

⁴⁷ REGULATION (EU) 2024/1689, Recitals 1 and 13; European Commission, Proposal for the Artificial Intelligence Act, COM(2021) 206 final, cit.

⁴⁸ M. SPRATT, Data Protection Impact Assessments: The Intersection of Data Protection, Risk Assessments and New Governance (PhD thesis, University College Dublin, 2025).

⁴⁹ European Commission, Proposal for the Artificial Intelligence Act, cit., p. 1; REGULATION (EU) 2024/1689, Recital 70.

public services. For this reason, the AI Act does not limit impact assessments to privacy concerns. Instead, it requires attention to a wider set of risks, including discrimination, exclusion, lack of transparency, and the lack of effective human control.⁵⁰

AI impact assessments have three main objectives. These objectives are not discretionary but reflect the compliance duties imposed on providers and deployers of high-risk AI systems under the AI Act. The first one is that they require organisations to identify and assess the risks related to how an AI system is designed, trained, and used. This will make it possible to view potential problems visible at an early stage. The second objective is the requirement for organizations to take steps to address those risks, for example by changing how the system works, improving data practices, or strengthening oversight. Finally, AI impact assessments create a written record of these decisions. This documentation allows both internal governance bodies and public authorities to check if legal obligations have been met and who is responsible for managing the risks.⁵¹

AI impact assessments also influence how organisations make decisions in practice. This influence is a direct result of the procedural obligations imposed by the AI Act on providers and deployers of high-risk AI systems. By requiring risks to be analysed and addressed in advance, the impact assessments affect whether an AI system is introduced at all, how it is used, and what safeguards are put in place. In this way, impact assessments do more than ensure formal compliance. They guide internal choices about innovation and how much risk an organisation is willing to accept when using AI systems.⁵²

Just like in the GDPR, the purpose of AI impact assessments is for them to be carried out more than once. They must be reviewed again when an AI system is changed, retrained, or used in a new context. This is important because AI systems change over time and if they are used in different situations, more risks can appear.

⁵⁰ Artificial Intelligence Act, Annex III.

⁵¹ REGULATION (EU) 2016/679, art. 35(7); KAMINSKI – MALGIERI, Algorithmic Impact Assessments under the GDPR, cit.

⁵² Organisation for Economic Co-operation and Development (OECD), Recommendation of the Council on Artificial Intelligence (22 May 2019); OECD, Policy Guidance on Responsible AI, 2021.

For this reason, impact assessments are treated as an ongoing process rather than a single step.⁵³

Substantive elements of AI impact assessments

According to the AI Act, AI impact assessments must contain concrete and practical information. At least, they have to explain what the AI system does, for what is going to be used for, and where it will be used. This is important because the same AI system can create very different risks depending on factors like the sector, the people affected, and the type of decisions involved.⁵⁴

One of the main parts of the assessment is about the data used by the AI system. This includes checking if the data is accurate and suitable for the purpose of the system. If the data is incomplete or biased, the system can lead to the production of discriminatory results, especially when it is used to assess individuals or predict their behaviour. For this reason, the assessment must look not only at technical issues, but also at how data is chosen, labelled, and grouped, and at the assumptions behind those choices.⁵⁵

In addition to data issues, AI impact assessments also need to look at the decision-making process that the AI system makes. This means examining how much the system operates on its own, when humans are involved, and whether people are able to review and intervene. This is even more important for high-risk AI systems under the AI Act, especially those listed in Annex III, where automated decisions can affect people's legal position or have serious consequences for their lives. In those cases, mistakes or biased outcomes can cause real and lasting harm.⁵⁶

AI impact assessments must also deal with transparency. Even if it is not always possible to fully explain how an AI system works at a technical level, the assessment should check whether people can understand how decisions are reached and on what grounds. This includes affected individuals, public authorities, and those inside the

⁵³ European Data Protection Board, Guidelines on Data Protection Impact Assessment (WP248 rev.01,2017).

⁵⁴ REGULATION (EU) 2016/679, art. 35(7)(a); Artificial Intelligence Act, arts. 11–12.

⁵⁵ Artificial Intelligence Act, art. 10.

⁵⁶ REGULATION (EU) 2016/679, art. 22; Artificial Intelligence Act, Annex III.

organisation who rely on the system. Under the AI Act, transparency is a binding requirement for high-risk AI systems, and impact assessments provide a practical way to verify that these obligations are met in practice. For this reason, transparency requirements under the AI Act are related to the impact assessment process.⁵⁷

Finally, AI impact assessments also need to evaluate how AI systems could affect fundamental rights. Some of the effects that need to be considered are the ones on equality, non-discrimination, human dignity, and access to legal remedies. This reflects the AI Act's explicit focus on the protection of fundamental rights, particularly in the context of high-risk AI systems. Looking at these issues sets AI impact assessments apart from traditional risk management tools. They are meant to guide responsible use of AI, not just meet technical or formal legal requirements.⁵⁸

Methodological and practical challenges

Even though AI impact assessments play an important role under the AI Act, it's very hard to apply them in practice. One of the main problems is turning broad legal ideas, such as fairness, proportionality, or respect for human autonomy, into clear points that can actually be checked and written down. This is harder than in data protection, where risks can usually be linked to specific data processing activities. With AI systems, harms are often uncertain, develop over time, and depend heavily on the context in which the system is used, which makes both identifying and reducing risks more difficult.⁵⁹ These difficulties arise in relation to legal principles and fundamental rights protected under EU law, in particular those reflected in the AI Act and in the Charter of Fundamental Rights of the European Union.

Another challenge is how responsibilities are shared inside organisations. For impact assessments to be effective, they require contributions from lawyers, technical staff, and managers. In reality, these groups often work separately and do

⁵⁷ Artificial Intelligence Act, arts. 13 and 52.

⁵⁸ Artificial Intelligence Act, Recitals 1 and 28; Charter of Fundamental Rights of the European Union.

⁵⁹ KAMINSKI – MALGIERI, Algorithmic Impact Assessments under the GDPR, cit.; SPRATT, Data Protection Impact Assessments: The Intersection of Data Protection, Risk Assessments and New Governance, cit.

not communicate enough. When roles are unclear and there is no clear way to raise issues or assign responsibility, impact assessments can easily become a formal exercise that is disconnected from how AI systems are designed and used.⁶⁰ This organisational fragmentation directly affects compliance with the AI Act, which places responsibility for risk management, documentation, and oversight on clearly identified providers and deployers.

Studies on DPIAs show that when rules are unclear or weakly applied, impact assessments can turn just into simple paperwork exercise. Instead of helping to prevent risks they are completed just to meet formal requirements. The same situation can happen with AI systems, especially when they are technically complex and when regulators have less information than the organisations that build and use them. Therefore it is very important that there is the proper guidance and effective oversight for AI impact assessments are to remain useful in practice.⁶¹

At the same time, if impact assessments become too inflexible, they can lose their practical value. If organisations treat them as fixed templates, the adaptation to new technologies or changing social conditions will not be effective. The real difficulty for regulators is to give clear rules while still allowing enough room for impact assessments to be updated and adjusted over time. This tension reflects the structure of the AI Act itself, which combines detailed obligations for high-risk systems with ongoing monitoring and post-market obligations, allowing risk assessments to evolve over the lifecycle of the AI system.

AI impact assessments as a bridge between legal compliance and organizational governance

Beyond their immediate compliance function, AI impact assessments play a broader role as connective instruments between legal requirements and internal organizational governance. By requiring structured documentation, internal consultation, and formalized risk evaluation, impact assessments serve as interfaces

⁶⁰ OECD, Policy Guidance on Responsible AI, cit.

⁶¹ SPRATT, Data Protection Impact Assessments: The Intersection of Data Protection, Risk Assessments and New Governance, cit.

between law, technology, and management. This connective role is not only for organizational in nature but follows directly from the obligations imposed by the AI Act on providers and deployers of high-risk AI systems.

Within organizations, AI impact assessments make it easier to have better communication between technical teams, legal departments, compliance functions, and senior management. They provide a shared framework through which risks can be discussed, responsibilities allocated, and strategic decisions justified. In this sense, impact assessments contribute to the institutionalization of responsible AI practices, embedding legal and ethical considerations into everyday operational processes.⁶² This internal coordination is required to comply with the AI Act's obligations on risk management, documentation, and human oversight, which require structured interaction between technical, legal, and managerial functions.

From a regulatory point of view, AI impact assessments also make it easier for authorities to oversee how AI systems are used. Although these assessments are prepared inside organizations, they produce records that regulators can later review in audits, investigations or after specific incidents. This documentation function reinforces accountability without requiring continuous direct intervention by regulators, aligning with the EU's broader approach to regulation based on documented procedures and oversight.

Finally, AI impact assessments can also help build trust around the use of AI systems. When organisations are able to show that risks have been identified, discussed, and dealt with in advance, this can increase confidence among users, affected individuals, and business partners. It may also reduce reputational damage when AI systems are deployed in sensitive areas. In this sense, impact assessments are not only legal safeguards, but practical governance tools that help balance innovation with accountability and public trust.

⁶² REGULATION (EU) 2016/679, arts. 5(2) and 24; SPRATT, *Data Protection Impact Assessments: The Intersection of Data Protection, Risk Assessments and New Governance*, cit.

2.4 AI impact assessments as instruments of EU regulatory governance

Beyond compliance: AI impact assessments as governance tools

AI impact assessments under the AI Act are not supposed to work only as formal compliance checks. While they are legal obligations, their real effect goes more than only verifying whether specific rules have been respected. They are designed to influence how organizations plan, discuss, and justify the use of AI systems before those systems are put into operation.⁶³

Unlike traditional regulatory tools that focus mainly on outcomes, AI impact assessments focus more on the decision-making process itself. They require organizations to pause before deploying an AI system and to make an analysis on how it works, what risks it may create, and how those risks can be reduced. This reflection must be documented and supported by an appropriate analysis. The law regulates the way decisions are made, not only the final result of those decisions.⁶⁴

This approach changes the meaning of compliance. Organizations are no longer asked only to meet external standards once a problem has already happened. Instead, they must build internal processes that allow risks to be identified and discussed early. AI impact assessments require written reasoning, internal coordination, and responsibility for choices that are made. These steps develop behavior even when regulators are not directly involved.

Because of this, AI impact assessments operate as governance tools. They organize internal discussions between legal experts, technical teams, and management. They create moments where assumptions can be questioned and design choices reconsidered. Over time, this affects how organizations approach innovation, responsibility, and risk.

This focus on procedures is part of a more general approach in EU digital regulation. EU lawmakers accept that they cannot fully predict how complex AI

⁶³ European Commission, Proposal for an Artificial Intelligence Act, COM(2021) 206 final.

⁶⁴ European Data Protection Supervisor, *Opinion 4/2015 on Towards a New Digital Ethics*, 2015.

systems will behave in real-world settings. Instead of relying only on detailed technical rules, they require organizations to adopt structured processes that guide decisions over time. AI impact assessments are a clear example of this method, as they change legal expectations into ongoing internal practices rather than one-off checks.

This focus on procedures also changes the way risk is discussed inside organizations. When organizations are required to document and justify their choices, Decisions that would normally be made quickly or for business reasons now have to be explained in terms of risks and safeguards. Business goals are still important, but they now have to be considered together with legal and social risks. Over time, this can influence how organizations prioritize projects and evaluate new uses of AI, even beyond situations where an impact assessment is mandatory.

Internalization of responsibility and organizational self-regulation

One of the main effects of AI impact assessments is that they take responsibility inside organizations. Instead of regulators identifying risks from the outside, organizations themselves are required to analyse how their AI systems work and what effects they may have. This change is intentional and follows the EU's general approach to managing risk.

Through AI impact assessments, organizations are expected to take ownership of the risks they create. They must identify potential problems, assess how serious they are, and decide how they should be addressed. This responsibility cannot be delegated entirely to regulators or external auditors. It becomes part of internal governance and decision-making.

This also changes how accountability works. Organizations must show that they did not ignore risks but actually tried to deal with them. Documentation plays a key role here. AI impact assessments create records that explain why certain choices were made and how risks were handled. These records can later be reviewed by authorities if problems arise.⁶⁵

⁶⁵ European Data Protection Board, *Guidelines on DPIA*, WP248 rev.01.

At the same time, AI impact assessments encourage a form of organizational self-regulation. Legal rules set the framework, but most of the work happens inside the organization. This includes defining internal roles, how reviews are carried out, and ensuring that risk assessment is taken seriously in all the departments. The effectiveness of this model depends a lot on organizational culture and internal incentives.

This does not mean that regulators step aside. Public authorities can still intervene when rules are ignored or when assessments are poorly done. However, the main responsibility for managing AI risks remains with the organization using the system. Impact assessments place day-to-day responsibility inside organizations, while regulators remain in the background to supervise and enforce when needed.

This internal responsibility also affects how failures are understood when problems arise. In this way, AI impact assessments matter not only before a system is used, but also later if something goes wrong. They make it easier to see what decisions were taken inside the organization and who was responsible for them. They create a more clear link between internal decision-making and external accountability.

Flexibility, discretion, and legal uncertainty

One of the defining features of AI impact assessments is their flexible and open-ended nature. The law does not provide a single fixed method for making an assessment. Instead, it sets general requirements and leaves organizations with a degree of discretion in how they meet them. This flexibility is made on purpose because of the diversity of AI systems and use cases.

Flexibility allows impact assessments to be adapted to different contexts. A system used in healthcare creates different concerns than one used in marketing or logistics. By avoiding fixed templates, the EU allows organizations to tailor

assessments to the specific risks involved. This makes the tool more adaptable to technological change.⁶⁶

At the same time, flexibility creates uncertainty. Because obligations are not defined in purely technical terms, organizations may struggle to know what level of analysis is enough. Different actors may interpret the same requirements in different ways, leading to unequal practices across sectors and countries. This risk has already been observed in the context of DPIAs under the GDPR.⁶⁷

There is also a risk of formalism. When obligations are open-ended, some organizations will tend to focus on producing documents that look complete but have little influence on real decisions. In those cases, impact assessments are only there mainly to show compliance rather than to improve outcomes. This weakens their preventive function.

The EU accepts these risks as a trade-off. Strict rules would be easier to apply, but they would quickly become outdated. Procedural obligations are more flexible, but they depend on guidance and supervision to work properly. AI impact assessments are built around this balance between flexibility and legal certainty.⁶⁸

The way organizations respond to this uncertainty varies. Some organizations respond carefully and carry out detailed assessments to limit regulatory risk. Others do much less and focus only on what they think is needed to meet the formal requirements. Differences like this are hard to avoid when the law leaves room for judgment. This variation is almost inevitable in a system that relies on discretion. However, over time the guidance from authorities and enforcement practice are expected to reduce extremes and create more consistent standards, even without rigid rules written into the legislation.

⁶⁶ Mantelero, *Computer Law & Security Review*, 2018.

⁶⁷ EDPB, *Guidelines on DPIA*, WP248 rev.01.

⁶⁸ European Data Protection Supervisor, *Accountability on the Ground*, 2019.

A distinctive EU regulatory model

Taken together, AI impact assessments show a clear EU approach to digital regulation. Rather than relying mainly on litigation or liability after harm has occurred, the EU relies more on prevention, process, and accountability. This approach is consistent across data protection, platform regulation, and now artificial intelligence.⁶⁹

Under this approach, the law does not aim to control every technical choice. It focuses instead on how decisions are made and how risks are addressed. Organizations are allowed some flexibility, but they are expected to explain their decisions and take responsibility for the systems they put into use.

This governance-oriented approach is different from systems that depend more on courts or private lawsuits. Sanctions and liability still matter, but they are not the main tool. The core of regulation lies in changing behaviour before harm occurs, through structured internal processes.

AI impact assessments therefore should not be seen as separate legal requirements. They are part of a more general regulatory style that uses procedures, documentation, and risk management as key tools. Understanding this is very important for assessing how the EU governs AI and how this approach compares with other regulatory models. This comparison will be explored in the next chapter.

This helps to explain why EU digital law usually gives more priority to procedures instead of fixed results. The idea is to regulate how decisions are taken, instead of trying to predict every possible outcome. This allows the rules to keep working even as technology changes. AI impact assessments fit well within this approach because they can evolve without constant legislative revision. As a result, they function not only as tools for managing current risks, but also as mechanisms that allow regulation to adapt over time. This adaptive quality will be important when comparing the EU model with other global approaches in the following chapter.

⁶⁹ GDPR; Digital Services Act; Proposal for AI Act.

Chapter 3: Comparative Perspectives on AI Governance and Impact Assessment Mechanisms

EU	US	China
Ex ante, risk-based	Ex post, sectoral	Ex ante, state-centric
Mandatory impact assessments	No mandatory IA	Mandatory algorithm filing
Rights-based	Market-driven	State/security-driven

Figure 2 – Comparative approaches to AI governance and risk management

Artificial intelligence is being used exponentially in many different countries, but the rules that regulate how it is developed and used are still mostly national or regional. How these risks are regulated changes from one legal system to another, because each country approaches technology, regulation, and innovation in a different way. There is no single global model of AI governance, but rather a set of approaches that change in structure and intensity.

This chapter places the European Union’s approach within the international context in which AI is regulated. Building on the analysis of impact assessments developed in the previous chapter, the focus here is on how different jurisdictions address AI-related risks and whether they rely on assessment mechanisms to manage those risks before harm occurs. The objective is not to provide a comprehensive overview of all AI-related regulation, but to highlight the main differences in regulatory philosophy and governance techniques.

The analysis concentrates on three reference models. The United States is examined as an example of a system that relies mainly on sector-specific rules and ex-post enforcement, with limited use of mandatory preventive assessments. China represents a contrasting model, characterized by strong ex-ante controls over algorithmic systems, primarily driven by state objectives. The chapter also considers selected Asian and international frameworks, which tend to rely on soft law instruments and voluntary guidelines.

By comparing these approaches, the chapter clarifies what makes the European Union's risk-based and impact-assessment-oriented model distinctive. This comparative perspective provides the basis for the final chapter, which will examine whether such assessment obligations can effectively balance innovation with the protection of fundamental interests.

3.1 AI Governance in the United States

The United States does not have a single law specifically focused on artificial intelligence. AI is instead regulated through existing sector-based laws, enforced by federal and state authorities. This shows a regulatory tradition that prefers flexibility and limited intervention, especially in fast-developing technological fields. As a result, AI governance in the United States has developed in a fragmented and more reactive way.⁷⁰

At the federal level, there are no AI-specific rules. The laws on consumer protection, discrimination, housing, credit, and privacy are used when AI systems affect decisions in those areas. As a result, different agencies are responsible for oversight, depending on the area involved. This means that there is no central authority responsible for supervising AI but rather a set of institutions addressing AI-related problems when they happen.

One of the most active agencies in this area is the Federal Trade Commission. The FTC has explained that AI systems can break consumer protection law when they are biased or give misleading results. This has been confirmed through enforcement actions, including cases in which companies were required to delete AI models built using illegally collected data.⁷¹ These cases show that even without AI-specific legislation companies can still face legal consequences when AI systems harm consumers or violate basic legal standards.

⁷⁰ Tatevik Davtyan, *The U.S. Approach to AI Regulation: Federal Laws, Policies, and Strategies Explained* (SSRN Working Paper, 9 Sept. 2024).

⁷¹ Federal Trade Commission, *FTC Finalizes Settlement with Photo App Developer Related to Misuse of Facial Recognition Technology* (Press Release, 7 May 2021).

A similar approach is used in employment law. The Equal Employment Opportunity Commission applies anti-discrimination rules to AI tools used in hiring and the evaluation of employees. The agency shows that employers are still responsible when automated systems lead to discriminatory results. This was confirmed in a 2023 case involving an AI hiring tool that excluded older applicants.⁷² The decision showed that using AI does not reduce employers' legal obligations under civil rights law.

AI systems have also been reviewed in housing and advertising situations. The Department of Housing and Urban Development and the Department of Justice have applied fair housing law to algorithmic ad targeting practices. The most popular example is the case against Meta, which resulted in a settlement requiring changes to its advertising systems to reduce discriminatory effects.⁷³ This case marked an important step in extending traditional civil rights rules to automated and data-driven activities.

A key difference in the U.S. approach is that companies are generally not required to assess AI risks before using these systems. Federal law does not require firms to do impact assessments in advance, even for AI tools with a high risk. Attempts to introduce those requirements, including the Algorithmic Accountability Act, have failed so far. This means that legal control usually happens only after AI systems have already been used and have caused harm or broken the law.

In addition to enforcement, U.S. authorities have promoted voluntary tools to guide responsible AI development. The National Institute of Standards and Technology's AI Risk Management Framework which was released in 2023, provides a optional structure for identifying and managing AI-related risks.⁷⁴ The White House also published an AI Bill of Rights, which explains basic principles for how AI should be used, such as safety, fairness, transparency, and human

⁷² Daniel Wiessner, "Tutoring Firm Settles US Agency's First Bias Lawsuit Involving AI Software," *Reuters*, 10 August 2023.

⁷³ U.S. Department of Justice, Justice Department Secures Groundbreaking Settlement Agreement with Meta Platforms to Resolve Allegations of Discriminatory Advertising (Press Release, 21 June 2022).

⁷⁴ National Institute of Standards and Technology (NIST), *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, January 2023.

oversight.⁷⁵ These initiatives show a clear preference for guidance and best practices rather than mandatory rules.

This preference is because they have a strong focus on innovation. U.S. policymakers have often said that strict rules introduced too early can slow down technological progress or make companies less competitive. This view appears in many policy documents, which promote light regulation and warn against having different rules in different states.⁷⁶ This means that public authorities usually intervene only once concrete harm has been identified.

At the same time, the lack of clear federal regulation has encouraged action at the state and local level. Many states and cities have started to use AI-related rules, especially in areas like biometric data, automated hiring, and consumer protection. Some of these rules including a recent legislation in Colorado, introduce impact assessments or discrimination audits for certain high-risk systems.⁷⁷ While these developments show growing awareness of risks related to AI, they also increase different rules across states, as companies must comply with different rules depending on where they operate.

Overall, AI governance in the United States is characterized by decentralization, reliance on existing legal tools, and a strong emphasis on ex-post enforcement. Impact assessments are not a central element of the U.S. approach, which contrasts with the European Union's risk-based regulatory model. This difference shows a deeper gap in regulatory philosophy and has an important role in the comparative analysis in this chapter.

⁷⁵ White House Office of Science and Technology Policy, *Blueprint for an AI Bill of Rights*, October 2022.

⁷⁶ Executive Order No. 14179, *Ensuring a National Policy Framework for Artificial Intelligence*, 11 December 2025.

⁷⁷ Joe Cahill, "AI Legal Watch: January 2026 – United States AI Compliance," Baker Botts Client Update, 27 January 2026.

3.2 China: State-Centric and Preventive AI Governance

China has taken a very centralized and preventive approach to artificial intelligence governance. Unlike systems that treat AI mainly as a commercial or technical tool, Chinese authorities view AI connected to national security, social stability, and political control. As a result, AI regulation in China is done primarily by state priorities rather than individual rights or market autonomy.⁷⁸

The Chinese Communist Party plays a very important role in this approach. AI is treated as a strategic technology that must support national unity, ideological alignment, and public order. National policy documents and AI strategies connect technological development to these goals, highlighting the need for systems to follow the “correct political direction” and respect the main Socialist Values.² Under the leadership of Xi Jinping, China has developed a model of “social governance” in which the state sets strict rules in advance. National initiatives like the New Generation AI Development Plan and the AI+ strategy place security and social stability above individual freedoms, making a strong emphasis on centralized control and ideological conformity.⁷⁹

In reality, this approach encourages the use of strong controls before AI systems are started to be used. New online services especially those that use algorithms to influence content, usually need approval before they can be used. Authorities monitor online activity closely, require users to register with their real names, and place strict obligations on online platforms.⁸⁰ These rules are backed by tough penalties like fines, service suspensions and in some cases criminal penalties. The goal is not only to reduce technological risks, but also to maintain political control and social order.

this governance logic translates into extensive preventive measures. New online services, especially those using algorithms to shape content or public opinion, must receive approval before being launched. Authorities closely monitor online activity,

⁷⁸ D. Papadopoulou, *Geopolitics: AI and China; enabling ideology?* (2025).

⁷⁹ Civil Society Alliances for Digital Empowerment (CADE), *Overview of AI policy in 15 jurisdictions* (2025).

⁸⁰ M. Sheehan, *China’s AI regulations and how they get made* (2023).

require real-name registration, and impose strict compliance duties on platforms.⁴ These rules are supported by strong enforcement powers, including fines, service suspension, and, in some cases, criminal liability. The objective is not only to manage technological risks but also to maintain political control and social order.

China has also been among one of the first countries to use detailed rules specifically targeting algorithms. A key element of this system is the algorithm filing requirement. Recommendation systems and other algorithms capable of influencing public opinion must be registered in a government database before use. Companies are required to disclose how their algorithms work, what data they rely on, and what risks they may have.⁸⁰ As part of this process, providers must submit an algorithm security self-assessment, which works like a mandatory audit conducted before deployment.

Additional requirements apply to some specific technologies. The 2021 rules on algorithmic recommendations have also transparency and user controls, while the 2022 Deep Synthesis rules introduce labelling requirements and identity checks for AI-generated content.⁸⁰ New draft rules on generative AI also apply these requirements to large language models, increasing controls over training data, outputs, and content moderation. Together, these measures allow authorities to supervise how algorithms operate in society and to limit the circulation of content considered harmful to public order.

China's AI framework also includes mechanisms that are similar to impact assessments. Through the algorithm filing system, companies must assess potential risks related to their systems before deployment and explain how they intend to address them.⁸⁰ These assessments focus mainly on security concerns, such as risks to national interests or social stability. Unlike rights-based impact assessments, their main function is to make sure that there is alignment with state objectives rather than to protect individuals.

This difference becomes more evident when China's approach is compared with the one used by the European Union. Both systems focus on preventive regulation and put the obligations before AI systems are used. However, the EU's risk-based model is built around the protection of fundamental rights, transparency, and accountability. For high-risk AI systems, the EU requires assessments to protect

individuals.⁵ On the other hand, China's preventive framework is state-centered and gives the priority to national unity and social stability over personal autonomy. This comparison shows that preventive regulation alone does not guarantee rights protection. The key difference lies in how each system approaches regulation.

3.3 Other Asian Approaches: Japan and Singapore (soft governance)

Japan

Japan has taken a more pro innovation and flexible approach to AI governance. Instead of choosing strict and detailed AI-specific laws, the Japanese government approach is more on not legally mandatory frameworks and policy guidance. In 2025, Japan adopted the AI Promotion Act, which has a general strategy to encourage AI development while acknowledging potential risks but this Act does not create enforceable legal obligations. Instead, it works as a policy signal, encouraging companies to adopt responsible practices on a voluntary basis.⁸¹

This method is supported by the AI Guidelines for Business issued by METI (Ministry of Economy, Trade and Industry) and MIC (Ministry of Internal Affairs and Communications) in 2024. The goal of these guidelines is to help companies apply ethical principles in their everyday use of AI. They focus on values such as fairness, privacy protection, safety, and a human centered approach to technology. The guidelines have clear examples, checklists, and case studies to help companies apply AI ethics correctly in their internal governance structures. They also promote the inclusion of senior management and suggest treating AI governance in a way like cybersecurity governance.⁸²

Mandatory rules remain limited. While the laws on issues such as data protection and copyright are applied to AI when relevant, Japan usually avoids punitive enforcement mechanisms. The regulatory philosophy can be described as “regulation by guidance,” with a strong focus on cooperation between government

⁸¹ Inoue, K., & Kamata, C., Japan's emerging framework for responsible AI: legislation, guidelines and guidance, International Bar Association, 2025.

⁸² METI / MIC, AI Guidelines for Business, 2024.

and industry. This makes the Japanese model attractive for innovation and business development, but it also means that protections for individual rights depend a lot on voluntary compliance instead of legal enforcement.

Singapore

Singapore has chosen also flexible and practical approach. The country does not have a single law focused only on AI. The regulation is mainly based on guidance developed by agencies like the IMDA and the PDPC.⁸³ In 2019, Singapore introduced the Model AI Governance Framework, which was later updated in 2020. This framework gives guidance on responsible AI use, accountability, transparency, and fairness.

Singaporean regulators have a more active coordinating role. For example, the PDPC established an AI Ethics Advisory Council, and together with IMDA they created initiatives like the AI Verify. AI Verify works as a testing framework and sandbox that allows organizations to evaluate their AI systems for risks related to bias, transparency, and safety.⁸⁴ These tools encourage companies to use internal risk assessments and governance practices, but they are not legally required.

The Singapore model shows a more “innovation-first” regulatory culture. Companies are encouraged to experiment and improve their practices under government guidance, without facing immediate sanctions. However, like the Japanese approach, this model is not very strict. Assessment mechanisms are recommended but not mandatory, and ethical principles are not supported by mandatory legal rules. As a result, innovation is supported but legal protection for individuals remains relatively weak.

⁸³ Nayak, R., Singapore’s AI governance regulations: Requirements and governance insights, Diligent, 2025.

⁸⁴ Allen, J. G., Loo, J., & Luna, J., Governing intelligence: Singapore’s evolving AI governance framework, Cambridge Forum on AI Law and Governance, 2025.

3.4 International and Transnational AI Governance Frameworks

Internationally, AI is managed by the use of voluntary frameworks, not mandatory or strict regulations. These frameworks do not have obligations, but they help countries agree on common principles and basic standards. This has created a common understanding of how AI risks should be tackled. Organizations like the OECD, UNESCO, and the G7 have an important role in this discussion.

OECD Principles

The OECD adopted its Recommendation on Artificial Intelligence in 2019. It was one of the first international efforts about AI governance. More than 40 countries have agreed with it, and it was later accepted by the G20. The OECD principles have the idea of “trustworthy AI.” They are based on a small number of key ideas, including fairness, transparency, safety, and accountability, as keeping people at the center of AI development.⁸⁵

One important aspect of the OECD framework is its focus on risk. It encourages those who develop and use AI systems to think about possible risks and deal with them throughout the life of the system. The level of attention depends on how the AI is used and how much impact it may have. The principles also mention transparency. They suggest that people affected by AI decisions should be able to understand how those decisions are made and question them. The OECD principles are not legally mandatory. Their use is voluntary, and they do not include enforcement measures.

UNESCO Recommendation

UNESCO approved its Recommendation on the Ethics of Artificial Intelligence in 2021 supported by all its member states. This document focuses on human rights at the center of managing AI. It relates the use of AI to the protection of human dignity, fundamental rights, and democratic values.⁸⁶

⁸⁵ OECD, Recommendation of the Council on Artificial Intelligence, OECD/LEGAL/0449, 2019.

⁸⁶ UNESCO, Recommendation on the Ethics of Artificial Intelligence, 2021.

Compared to the OECD principles, the UNESCO Recommendation gives more guidance on how ethical ideas could be applied. It mentions tools like oversight bodies, impact assessments, audits, and due diligence processes to improve accountability. It also asks states to prevent harmful uses of AI, including mass surveillance and social scoring systems that affect the human rights. However, the UNESCO framework is not legally mandatory so each state decides how to apply it and the compliance depends on political commitment.

G7 and other initiatives

International coordination is also present in smaller groups like the G7. In 2023, the G7 started the Hiroshima AI Process. This process created principles and a voluntary AI code of conduct for AI systems. These documents are about managing risks, improving transparency, and monitoring how AI systems are used over time. They have a similar approach to the work done by the OECD and UNESCO.⁸⁷ However, the G7 framework is voluntary and does not impose legally binding obligations on states or companies.

Overall, international frameworks describe responsible AI, but they leave enforcement open. They support ideas like impact assessments and accountability but leave the decision on how to apply them to national governments. In this area, the EU is the one that is making mandatory legal obligations that operationalize these principles.

3.5 Comparative Synthesis: What Makes the EU Model Distinct

When these approaches are compared there are clear differences. The EU model is different from the others in many ways. The main differences are when the rules apply, how strict they are, and how enforcement works. This helps explain why the EU system is more organized and puts more legal responsibility on companies than the other approaches.

⁸⁷ Ernst & Young, G7 AI Principles and Code of Conduct – Key Points, 2023.

First, the EU uses a preventive method. Under the AI Act, AI systems that have a high risk must be checked and meet certain conditions before they can be used on the market. This is different from the United States, where action usually happens only after damage has already happened. Japan and Singapore also do not require mandatory checks before AI systems are used, they rely on guidance and monitoring after the AI systems are used.

Second, the EU relies on binding legislation. While international frameworks go more for voluntary compliance, the EU takes ethical principles into legal obligations. If companies do not follow these rules, they can face fines and other penalties. This approach leads to less flexibility, but it also creates consistent rules in all EU countries. As a result, companies have the same standards no matter where they operate within the EU.

Third, the EU's model is rights-based. Its regulatory framework gives more priority to the protection of fundamental rights like non-discrimination, privacy, and human dignity As opposed to the U.S. approach which prioritizes market efficiency and innovation, intervening usually through existing laws. Moving on to the China's approach, this one focuses on state control and social stability rather than individual autonomy. These differences show that similar regulatory tools can achieve very different political and legal objectives.

Finally, the EU focuses more on how companies manage AI inside their organizations instead of waiting for the people to bring lawsuits after problems occur. Companies are expected to set up internal processes to identify and manage risks before AI systems are used. In the United States, responsibility often comes only after a specific case, such as a lawsuit or a regulatory investigation. The EU takes a different approach by asking companies to deal with risks early. For this reason, impact assessments play a central role in the EU system and are meant to guide everyday decision-making, not just to meet formal requirements.

3.6 Conclusion

This chapter has shown that there is no single “ideal” model of AI governance. Each approach has different priorities and political choices. The U.S. and many Asian countries go for more flexibility and innovation, but this can leave gaps in protection. China’s state-centered model ensures strong control but often limits individual rights. The European Union, while less innovation-friendly in the short term, offers the most coherent and structured legal framework.

In all the models, impact assessments and accountability are becoming more common. What really differs is whether companies must use them or not. The EU approach shows that impact assessments can help reduce risks and increase trust, instead of simply slowing down innovation.

This analysis prepares the ground for Chapter 4. The next step is to assess whether the EU’s approach can balance innovation and protection in practice, and whether its ambitious regulatory framework delivers the outcomes it promises.

Chapter 4: Innovation vs Protection in AI Regulation

4.1 The Innovation–Protection Tension

Artificial Intelligence its becoming a key factor of success for many businesses. Many companies use AI for different things like creating new products, improving services, and reducing costs. In addition to that, AI systems can work with large amounts of data and find patterns that humans cannot see. Thanks to those benefits, companies can make better decisions and compete in the market. That is why AI is seen as a key factor of innovation and economic growth. But one of the main components that AI needs to work properly is data. In most cases, AI systems need large datasets to learn and improve over time.⁸⁸

Most of the time, in today's world that is driven by data, Businesses often collect data about customers, users, and markets. This data includes things like personal information, online behaviour, or consumer preferences. The more data an AI system has, the better it usually performs and this is the reasons they collect and process as much data as possible. In this sense, data is the main resource behind AI innovation.⁸⁹

At the same time, the law exists to limit harm that can result from data use. Not all data processing is allowed. The use of personal data faces risks like being misused, leaked, or used in unfair ways. This is why data protection law exists to prevent these risks. In the European Union, the GDPR has strict rules on how personal data can be collected and used.⁹⁰ One important rule is data minimisation. This means that companies should only use the data that is necessary for a specific purpose. This legal requirement can be challenging for improving the performance and development AI, which often benefits from using large and diverse datasets.

This creates a clear legal tension. On one side, companies want more data to improve AI systems and stay competitive. On the other side, the law creates limits

⁸⁸ European Commission, White Paper on Artificial Intelligence – A European approach to excellence and trust, COM(2020) 65 final.

⁸⁹ European Commission, *A European Strategy for Data*, COM(2020) 66 final.

⁹⁰ Regulation (EU) 2016/679 (General Data Protection Regulation), OJ L 119, 4.5.2016.

with the purpose of protecting individuals. This tension is not about moral ideas, but about legal rules that companies must follow. Businesses must respect privacy, avoid discrimination, and protect fundamental rights. If an AI system causes harm the law must decide who is responsible. Someone must be accountable for the decision made by the system.

Fundamental rights such as privacy, equality, and protection from unfair treatment are protected by EU law. AI systems can affect these rights if they are not properly controlled. Because of this, the European Union tries to promote innovation while also ensuring that AI systems are safe and lawful. The central legal problem is how to support AI development without breaking the rules that protect individuals and society.

This tension has an effect in how the businesses plan all of their AI projects because companies need to think not only about what AI can technically do, but also about what the law allows. For example, a company could have access to large datasets, but GDPR limits how personal data can be used. This creates legal uncertainty, especially when AI systems change over time. An AI system that is legal at first can later become illegal if it starts using more data. Because of this, businesses must review their AI systems after some time to remain compliant. This shows that the innovation–protection tension is a real legal issue for companies using AI.

4.2 Risks of Over-Regulation

The use of regulation is essential and necessary in most of the cases, but it can also lead to numerous disadvantages if it becomes too strict. If the AI laws put many heavy obligations, then the companies will have it very difficult to develop new technologies quickly. Strict legal requirements can reduce innovation and in a fast-changing field like AI, it can reduce it even more than in other fields. If every new AI system must go through long approval processes or complex checks, companies may delay or cancel projects.

Compliance costs are a big issue specially for small companies because large companies usually have legal teams and technical experts who deal with regulation.

They can afford audits, certifications, and legal advice. Small and medium-sized businesses do not have the same resources so even a small compliance requirement can be expensive. If an AI regulation requires detailed documentation or constant monitoring, the cost can be very high for smaller firms and startups are key because they contribute a lot to the economy.

Complex laws also affect small companies more than large ones. Complex laws are difficult to understand and apply. SMEs often rely on a smaller number of employees who already manage many tasks. Time spent on legal compliance is time taken away from product development or innovation. In some cases, small businesses may avoid using AI at all to reduce legal risk. This is not the intention of regulation, but it can be an unintended effect.

There is also a bigger economic risk. If regulation in Europe becomes too strict, companies may choose to move their AI activities to other regions. The United States and China are often seen as more flexible environments for AI development. Policymakers have expressed concern that Europe could fall behind if legal requirements are too heavy. Innovation may move to places where regulation is more simple, even if the EU aims to protect users.

Over-regulation does not mean that regulation is wrong. It means that rules can sometimes go further than needed. If compliance becomes too costly or too slow, innovation gets affected. The legal challenge is to design rules that protect people without discouraging legitimate business activity.

Over-regulation can also generate problems when there are too many legal rules and they are unclear. If companies do not understand what is required, they may take a very cautious approach to avoid fines. This can lead to companies stopping projects that are actually lawful. This behaviour does not increase safety, but it reduces innovation. From a legal point of view, unclear or overly detailed rules can make it more difficult for businesses to start using and benefiting from AI, even when risks are low. This shows how strict regulation can have unintended legal effects on innovation.

4.3 Risks of Weak Protection

While too much regulation can slow innovation, too little regulation can also create serious problems. If AI systems are used without legal limits, individuals can be harmed. One of the main risks is the violation of privacy. If AI systems use personal data without proper protection, private information can be leaked. Without clear legal rules, people could lose control over their own data.

Discrimination is also one of the main risks. AI systems can copy unfair patterns from the data they are trained on, so if these systems are used in hiring, credit scoring, or insurance, they could give an unfair disadvantage to certain groups of people. Without clear legal rules, companies may use AI systems that are unfair and that could break equality laws.

There can also appear safety risks related to the use of AI. AI systems used in healthcare or transport can make decisions that affect physical safety. If these systems are not properly tested or regulated, their mistakes can cause real harm. When legal protection is weak, individuals often suffer the consequences of AI failures.

Another thing is that weak protection also affects trust because people are more likely to use AI systems if they believe they are safe and fair. If users hear about data breaches, biased algorithms, or unsafe AI decisions, they will lose their trust to those systems and will stop using them. Trust is important for market adoption. European policy has clearly stated that trust is necessary for the use of AI technologies.⁹¹

From a legal point of view, weak regulation creates uncertainty. Companies may not know which standards apply or what is expected of them. This can seem attractive at first, but it increases legal risk. If an AI system has negative effects, people can take companies to court using existing laws. Judges could take decisions on these cases without clear rules, and this will make the situation uncertain for businesses.

⁹¹ European Parliament, Artificial Intelligence and Fundamental Rights, 2020.

In many cases, weak regulation leads to stricter rules later. A major incident can make lawmakers to react quickly with strong restrictions. This can be more disruptive than gradual regulation. Therefore, a lack of protection today can result in heavier regulation tomorrow.

Weak protection also creates long-term legal risks for businesses. If AI systems cause harm and there are no clear legal standards, companies could face lawsuits based on general liability or consumer protection law. Court decisions in these cases can be unpredictable, which increases legal uncertainty. In addition, serious incidents may push lawmakers to introduce stricter rules very quickly. Sudden regulation is often harder to comply with than gradual rules. From a business law perspective, weak protection today can lead to stronger and more disruptive regulation later.

4.4 Finding a Middle Ground

Because over-regulation and weak protection can create problem the best path is a balanced legal approach. The goal is not to block innovation, but to guide it. One way to do this is through risk-based regulation. This means that legal obligations depend on how risky an AI system is.

The proposed EU AI Act follows this approach. It blocks only a some AI practices that are considered unacceptable. AI systems that have a high risk like those used in healthcare or recruitment, are subject to stricter rules. Low-risk AI systems face fewer obligations.⁹² This allows regulators to focus on the areas where harm is more likely.

Flexibility is also important because AI technology changes quickly. If laws cannot change easily, they can become old and useless. Flexible rules allow regulators to adapt to new developments. Tools like regulatory sandboxes allow companies to test AI systems under supervision, without full legal pressure.

⁹² Irish Data Protection Commission, Guidance on Data Protection Impact Assessments (DPIAs).

Another important tool is impact assessments. Instead of telling companies exactly how to build AI systems, the law can require them to assess risks themselves. This shifts responsibility to businesses but also gives them more freedom to innovate within legal limits. Law and innovation do not have to be enemies. Clear legal boundaries can actually help companies innovate with more certainty.

4.5 Impact Assessments as an Enabler of Responsible Innovation

Impact assessments play a key legal role in balancing innovation and protection. They require companies to examine how an AI system can affect people before it is used. This helps identify legal risks early.

A similar tool already exists in data protection law. Under GDPR, companies must conduct a Data Protection Impact Assessment (DPIA) when processing has high risks to individuals. A DPIA forces companies to analyse risks and take steps to reduce them. It also helps demonstrate compliance to regulators.

AI impact assessments work in a similar way. Instead of reacting after harm is already done, companies think about the risks during development. This can improve design choices and reduce legal exposure. Most of the times, fixing problems early is usually cheaper than dealing with lawsuits later.

Impact assessments also support trust. When companies show that they have assessed risks, users and regulators are more likely to trust their systems. So trust does not only become a legal advantage but also an economic advantage. Experience with DPIAs shows that early compliance can reduce legal problems and increase confidence. Error! Bookmark not defined.

From a business point of view, impact assessments also improve internal decision-making. They require cooperation between legal, technical, and management teams. This will help the companies identify legal risks early. Instead of treating compliance as a final step, it becomes part of development. Experience with GDPR shows that DPIAs help companies avoid serious compliance problems later. The same logic applies to AI impact assessments, which reduce legal surprises and support responsible innovation.

4.6 Final Reflections

Balancing innovation and protection in AI regulation is not easy. This chapter showed that law must support innovation while preventing harm. Too much regulation can slow progress, but not enough regulation can lead to harm.

AI regulation is one of the fields that will continue to evolve a lot because technology is also evolving at a high pace. New AI systems will raise new legal questions that existing rules may not address. For this reason, regulators and businesses must remain flexible. Law cannot stop innovation, but it can guide it. The EU approach shows an attempt to act early and prevent harm while still allowing progress. Even if the balance is not perfect, it provides legal direction for businesses working with AI.

The European Union follows a more preventive approach. Its goal is to act before the harm happens, protect fundamental rights, and still allow businesses to innovate. This balance is difficult, but necessary. Business cyber law must continue to adapt as AI evolves. Finding the right balance will remain a key challenge, but it is essential for long-term and lawful AI development.

Bibliographic References

- [1] European Commission (2021). *Proposal for a Regulation laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act)*, COM(2021) 206 final.
- [2] Pasquale, F. (2015). *The Black Box Society*. Cambridge (Mass.).
- [3] Varian, H. (2019). Artificial Intelligence, Economics, and Industrial Organization. *NBER Working Papers*.
- [4] Barocas, S., & Selbst, A. D. (2016). Big Data's Disparate Impact. *California Law Review*.
- [5] European Commission (2020). *White Paper on Artificial Intelligence – A European approach to excellence and trust*, COM(2020) 65 final.
- [6] Russell, S., & Norvig, P. (2021). *Artificial Intelligence: A Modern Approach* (4th ed.). London.
- [7] OECD (2019). *Artificial Intelligence in Society*. Paris.
- [8] Barocas, S., Hardt, M., & Narayanan, A. (2019). *Fairness and Machine Learning*. Cambridge (Mass.).
- [9] Wachter, S., & Mittelstadt, B. (2019). A Right to Reasonable Inferences. *Columbia Business Law Review*.
- [10] European Union Agency for Fundamental Rights (2020). *Getting the Future Right – Artificial Intelligence and Fundamental Rights*. Vienna.
- [11] Spark Legal & Policy Consulting (2023). *Study on Systemic Risks and their Mitigation under the Digital Services Act*.
- [12] Palumbo, A. (2026). Systemic risk management as an emerging regulatory approach in EU digital legislation. *Technology and Regulation*.
- [13] Bietti, E. (2024). Can Private Law Protect Privacy in Today's Economy? *Balkanization*.
- [14] European Parliament and Council (2022). Regulation (EU) 2022/2065 on a Single Market for Digital Services (Digital Services Act). Official Journal of the European Union. Available at: <https://eur-lex.europa.eu/eli/reg/2022/2065/oj>
- [15] European Parliament and Council (2016). Regulation (EU) 2016/679 of 27 April 2016 on the protection of natural persons with regard to the processing of personal data (General Data Protection Regulation), OJ L 119, 1–88. Available at: <https://eur-lex.europa.eu/eli/reg/2016/679/oj>
- [16] Montagnani, M., et al. (2024). The EU regulatory approach(es) to AI liability. *Computer Law & Security Review*.
- [17] European Commission (2021). *Better Regulation Guidelines*, SWD(2021) 305 final.
- [18] Majone, G. (1996). *Regulating Europe*. Routledge.

- [19] Hood, C., Rothstein, H., & Baldwin, R. (2001). *The Government of Risk*. Oxford University Press.
- [20] Mantelero, A. (2014). The Future of Consumer Data Protection in the EU. *Computer Law & Security Review*.
- [21] European Data Protection Board (2017). *Guidelines on Data Protection Impact Assessment (DPIA)*, WP248 rev.01.
- [22] Zuiderveen Borgesius, F. (2020). Strengthening Legal Protection Against Discrimination by Algorithms. *International Journal of Human Rights*.
- [23] Floridi, L., et al. (2018). AI4People – An Ethical Framework for a Good AI Society. *Minds and Machines*.
- [24] Council of Europe (2019). *Algorithms and Human Rights*.
- [25] European Commission (2021). Proposal for an Artificial Intelligence Act, COM(2021) 206 final.
- [26] European Parliament and Council (2024). Regulation (EU) 2024/1689 (Artificial Intelligence Act). Official Journal of the European Union. Available at: <https://eur-lex.europa.eu>
- [27] Kaminski, M. E., & Malgieri, G. (2021). Algorithmic Impact Assessments under the GDPR. *International Data Privacy Law*.
- [28] Spratt, M. (2025). *Data Protection Impact Assessments: The Intersection of Data Protection, Risk Assessments and New Governance* (PhD thesis, University College Dublin).
- [29] OECD (2021). *Policy Guidance on Responsible AI*.
- [30] European Data Protection Supervisor (2015). *Opinion 4/2015 on Towards a New Digital Ethics*.
- [31] European Data Protection Supervisor (2018). *Accountability on the Ground*.
- [32] Datvyan, T. (2024). The U.S. Approach to AI Regulation: Federal laws, policies, and strategies explained. *SSRN Working Paper*, 9 September 2024.
- [33] Federal Trade Commission (2021). FTC Finalizes Settlement with Photo App Developer Related to the Use of Facial Recognition Technology (Press Release, 7 May 2021).
- [34] Wiesner, D. (2023). Tutoring Firm Settles US Agency's First Bias Lawsuit Involving AI Software. *Reuters*, 10 August 2023.
- [35] U.S. Department of Justice (2022). Justice Department Secures Groundbreaking Settlement Agreement to Resolve Claims of Discriminatory Advertising (Press Release, 21 June 2022).
- [36] National Institute of Standards and Technology (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, January 2023.
- [37] White House Office of Science and Technology Policy (2022). *Blueprint for an AI Bill of Rights*, October 2022.

- [38] Executive Order No. 14179 (2025). Ensuring a National Policy Framework for Artificial Intelligence, 11 December 2025.
- [39] Cahill, J. (2026). AI Legal Watch: January 2026 – United States AI Compliance. *Baker Botts Client Update*, 27 January 2026.
- [40] Papadopoulou, D. (2025). Geopolitics: AI and China; enabling ideology?
- [41] Civil Society Alliances for Digital Empowerment (CADE) (2025). *Overview of AI policy in 15 jurisdictions*.
- [42] Sheehan, M. (2023). China’s AI Regulations and How They Get Made.
- [43] Inoue, K., & Kamata, C. (2025). Japan’s Emerging Framework for Responsible AI Legislation, Guidelines and Guidance. *International Bar Association*.
- [44] METI / MIC (2024). *AI Guidelines for Business*.
- [45] Nayak, R. (2025). Singapore’s AI Governance Regulations: Requirements and Governance Insights. *Diligent*.
- [46] Allen, J. G., Loo, J., & Luna, J. (2025). Governing Intelligence: Singapore’s Evolving AI Governance Framework. *Cambridge Forum on AI Law and Governance*.
- [47] OECD (2019). Recommendation of the Council on Artificial Intelligence, OECD/LEGAL/0449.
- [48] UNESCO (2021). *Recommendation on the Ethics of Artificial Intelligence*.
- [49] Ernst & Young (2023). *G7 AI Principles and Code of Conduct – Key Points*.
- [50] European Commission (2020). *A European Strategy for Data*, COM(2020) 66 final.
- [51] Charter of Fundamental Rights of the European Union (2012). Official Journal of the European Union, C 326, 391–407.
- [52] European Parliament (2020). *Artificial Intelligence and Fundamental Rights*.
- [53] Irish Data Protection Commission. *Guidance on Data Protection Impact Assessments (DPIAs)*.
- [54] Regulation (EU) 2016/679 (GDPR) (2016). Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data, OJ L 119, 4.5.2016.