

The Impact of Artificial Intelligence on Legal Systems

*Normative, Ethical, and Regulatory
Perspectives*

by

Halim Bajraktari and Fatlind Mazreku

**The Impact of Artificial Intelligence on Legal Systems:
Normative, Ethical, and Regulatory Perspectives**

by Halim Bajraktari and Fatlind Mazreku

2026

Ethics International Press, UK

British Library Cataloguing in Publication Data

A catalogue record for this book is available from the British Library

Copyright © 2026 by Halim Bajraktari and Fatlind Mazreku

All rights for this book reserved. No part of this book may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, electronic, mechanical, photocopying, recording or otherwise, without the prior permission of the copyright owner.

ISBN (Hardback): 978-1-83711-754-3

ISBN (Ebook): 978-1-83711-755-0

Contents

List of Figures	vii
Preface.....	xi
Foreword	xii
Part I – Foundations	
Chapter 1: AI in Context: Technological and Legal Evolution	1
Chapter 2: Fundamentals of AI Systems and Data Ethics	28
Chapter 3: Legal Theory Meets Machine Logic.....	55
Part II – Legal Challenges in AI Innovation	
Chapter 4: Data Privacy and Consent in Machine Learning.....	81
Chapter 5: AI Accountability and Legal Liability.....	106
Chapter 6: Intellectual Property in the Age of Autonomous Creation	132
Part III – Regulation and Global Frameworks	
Chapter 7: Regulating AI: GDPR, CCPA, and the EU AI Act.....	159
Chapter 8: AI in the Public Sector: Justice, Policing, Surveillance..	187
Chapter 9: Comparative Global Models for AI Governance.....	214
Part IV – Responsible Futures	
Chapter 10: Ethics by Design: How Law Can Guide AI Innovation.....	243
Chapter 11: Toward Trustworthy AI: Compliance, Auditing, and Risk	270
Chapter 12: Co-Evolution of Law and AI Looking Ahead	296
References	323

List of Figures

Figure 1.1 Timeline of AI development: from symbolic rule systems to generative foundation models.

Figure 2.1 Machine learning taxonomy: supervised, unsupervised, and reinforcement learning.

Figure 4.1 GDPR data protection flow: lawful basis, rights, and enforcement pathways.

Figure 5.1 AI liability chain: developer, deployer, user, and affected person.

Figure 6.1 Intellectual property frameworks applicable to AI inputs and outputs.

Figure 7.1 EU AI Act risk pyramid: prohibited, high-risk, limited-risk, and GPAI tiers.

Figure 8.1 Public-sector AI: safeguard layers and risk severity across domains.

Figure 9.1 Governance maturity radar: EU, US, China, UK, and Western Balkans.

Figure 10.1 Ethics by design: lifecycle integration of legal and ethical requirements.

Figure 11.1 Accountability stack: from technical controls to regulatory enforcement.

Figure 12.1 Co-evolution model: governance maturity trajectories by jurisdiction.

Preface

This book examines the legal implications of artificial intelligence in all the areas most affected by algorithmic decision-making: privacy, accountability, intellectual property, ethics, public sector governance, and the constitutional structure of democratic institutions. It is written for legal scholars, practicing lawyers, policymakers, technology professionals, and advanced students who need a rigorous and technically informed account of how AI is changing the meaning of legal categories.

The authors bring complementary training to the project. Halim Bajraktari brings expertise in European law, international criminal justice, and human rights, with a particular focus on how public institutions acquire, enable, use, and hold accountable automated decision-making tools. Co-author Fatlind Mazreku brings expertise in machine learning systems and software engineering, with a focus on how technical design choices translate into legal risk and governance failure.

The book covers twelve chapters in four parts. Part I establishes the conceptual relationship between AI and law. Part II examines privacy, liability, and intellectual property as the doctrinal areas most directly challenged. Part III analyzes regulatory frameworks from the EU AI Act to California's CCPA and comparative international models. Part IV addresses responsible design, trustworthy artificial intelligence, and the co-evolution of law and technology over the next decade, the interplay and analysis of which engage readers as they explore topics related to the public sector, governance, ethics, innovation, trust, and the risks posed by the co-evolution of law and law enforcement, both now and in the future.

Foreword

Society is experiencing an increasing technological development, especially in Artificial Intelligence (AI) in the last decade, which has opened new horizons in... scientific, academic and technological publications, causing a deep debate on its effects on the production and use of scientific information, especially on the integrity of research and study. Today, we hear more from experts in various fields about the dilemmas and opportunities of the transformative impact of AI on modern society, especially in science and society legal and economic publications.

This book, *"The Impact of Artificial Intelligence on Legal Systems: Normative, Ethical and Regulatory Perspectives"* is dedicated to solving the dark puzzles of the present, and the bright discoveries of the future that the impact of AI will bring, about risks and the benefits that this human-enabled, human-to-human, technological approach will bring as a favor or disadvantages. A pioneering approach in justice and legal services for the use of predictive analytics in judicial decisions, by providing a model for building a system of how AI will impact the legal system, through enable able and explainable artificial intelligence, will extract best observational practices through the perspective of AI, serving the modern legal system and society as a whole in the digital age. A future focus will be to build a regional and universal system for several countries in the use of Artificial Intelligence in the justice system.

The book contains themes and opportunities for the justice system and how it will help us to facilitate these social opportunities, how to use them in the social interest in social evolution in the 21st century is bringing more and more uncertainty, doubling down on the new and

rapid opportunities that AI offers, while simultaneously addressing a question for humanity: how will we live with artificial intelligence or the opposite? One thing will be certain: man will gain some benefits, but will lose some qualities; there will be uncertainty about privacy, it will be more usable and less reliable, faster and more explorable, but less understandable for the previous ones generations. The topics of *"The Impact of Artificial Intelligence on Legal Systems"* is a reality of the current generations with a vision for the future which the authors have together reflected real facts in the book. The research problems address a dimension especially for the justice system, addressing innovations and various legislative possibilities in line with EU ethical and normative standards, and international rules that affect, endanger and help the interests of society, bringing a multitude of protector aspects that AI brings, so that we as a society can use them to our advantage in today's world, using artificial intelligence to benefit the justice system.

"The Impact of Artificial Intelligence on Legal Systems" is dedicated to a wide audience, including current and future social generations, with a high technological level, which is practiced in legal innovations that are easier and more efficient. There is a need for collaboration and research to make work easier, research and trials, at the same time to produce preventive results in preventing misuse, making people's lives easier.

The book can be used as a completely new model of light from the research that emerged in... chapters on the areas of justice, normative, ethical and regulatory perspectives, as an opportunity that brings concrete best practices from research institutions, institutes, distinguishing the research experiences of experts in the field of justice, police, criminology, criminology, economics and penology, to help and protect human rights. Book in words examines a current,

new and broad approach, a scientific, visionary, academic and technological trend that responds to the demand of the contemporary exploratory audience on justice and AI, especially in the field of law. It follows new social trends, not only in technology and economics, but with a special emphasis on the justice system. This book will serve as an example and model to encourage justice system institutions to address policies aimed at using artificial intelligence to benefit the legal system. It focuses on combating injustices and abuse, as well as confronting negative phenomena in society. It also promotes the use of artificial intelligence in other sectors to combat terrorism and other phenomena harmful to society.

"The Impact of Artificial Intelligence on Legal Systems" also examines the role of the legal system. education and training to prepare young professionals who can use advanced artificial intelligence technologies. The book also discusses the need to standardize practices for using artificial intelligence in legal systems, to guarantee the protection of individuals' rights and avoid possible discrimination that can occur from unquestionable algorithms. This includes the development of ethical and legal regulations that will protect the integrity of legal processes, as well as promoting transparency in how AI operates makes decisions that affect people's lives.

In this context, the book positions itself as an important resource for researchers and professionals who want to understand not only the transformative potential of AI, but also the challenges, opportunities, and Accountability that come with its use in legal systems, as well as in other spheres.

Dafina Abdullahu

University Ukshin Hoti Prizren, Kosovo

Part I

Foundations

Chapter 1

AI in Context: Technological and Legal Evolution

AI in context: technological and legal evolution

In 2013, a Wisconsin circuit court sentenced Eric Loomis to six years in prison. The judge cited, among other considerations, a recidivism risk score produced by a proprietary software system called COMPAS. Loomis challenged the sentence on the ground that he had not been given meaningful access to the algorithm that scored him. The Wisconsin Supreme Court rejected his challenge in 2016, holding that the score was one factor among several and that the judge had not relied on it exclusively.¹ What the decision did not resolve- and what courts across several jurisdictions are still working through- is whether a person can meaningfully contest a decision influenced by a system whose internal logic is not disclosed. That unresolved question runs through this book. The Loomis case did not begin with a conversation about artificial intelligence. It began with a sentencing hearing in a criminal court. Technology arrived as infrastructure present, consequential, and largely unexamined. That is how AI typically enters legal settings. It does not announce itself as a disruption. It arrives embedded in a workflow, a procurement contract, a vendor system licensed by a public institution. Lawyers, judges, and affected parties often encounter it only when something goes wrong. This chapter sets out the context in which that arrival makes sense legally. Its purpose is not a complete history of computer science. It is to identify the moments where AI changed the practical meaning of legal concepts such as responsibility, evidence, consent,

and authorship- and to explain why those changes require sustained legal analysis rather than occasional judicial improvisation.

From imitation to deployment

Alan Turing's 1950 paper "Computing Machinery and Intelligence" is the standard starting point for AI history, and it earns that position.² Turing proposed an operational test: if a machine could sustain a conversation indistinguishable from a human's, the question of whether it could "think" became less pressing. The test mattered philosophically, but it also encoded a legal assumption that would prove enduring- that the measure of an AI system should be what it can do in practice, not what it is metaphysically. The decades that followed did not match the early optimism. The first AI winter- a period of reduced funding and deflated expectations that followed criticism of overpromised expert systems in the 1970s and 1980s- demonstrated that technical ambition without adequate computational foundations produces limited results.³ Systems such as MYCIN, a medical expert system developed at Stanford in 1972, showed genuine capability within narrow domains but could not generalise. Their legal advantage, ironically, was their traceability. MYCIN's recommendations followed explicit rules that could be examined, disputed, and updated. If it recommended a drug treatment, a reviewer could trace the chain of reasoning back to specific facts and specific rules. That traceability disappeared as machine learning replaced rule-based systems. The transition began in earnest in the 1990s and accelerated sharply in the 2010s. Geoffrey Hinton, Yann LeCun, and Yoshua Bengio- who would share the Turing Award in 2018 for their contributions to deep learning- demonstrated that neural networks trained on large datasets could recognise patterns that no human had explicitly programmed.⁴ AlexNet's performance at the 2012 ImageNet competition is typically

cited as the inflection point: a convolutional neural network achieved error rates that outperformed prior state-of-the-art methods by a margin that surprised even specialists. The legal implications of this transition were not immediately obvious. Courts, regulators, and legislators were still developing frameworks for the earlier generation of rule-based systems when statistical learning had already displaced it. By the time GDPR entered into force in May 2018, the most consequential AI systems were already operating through learned statistical patterns that their operators could not fully explain. The third transition- to large-scale generative and foundation models- created a different order of problems. Models such as GPT-3, released by OpenAI in 2020, and its successors demonstrated that a single model trained on broad datasets could generate text, code, images, and reasoning chains across many domains. ChatGPT, released in November 2022, brought these capabilities into everyday professional and institutional use.⁵ By early 2023, lawyers were submitting AI-generated court filings, students were submitting AI-generated essays, and medical professionals were using AI to interpret images and draft clinical notes. The infrastructure had arrived faster than the governance.

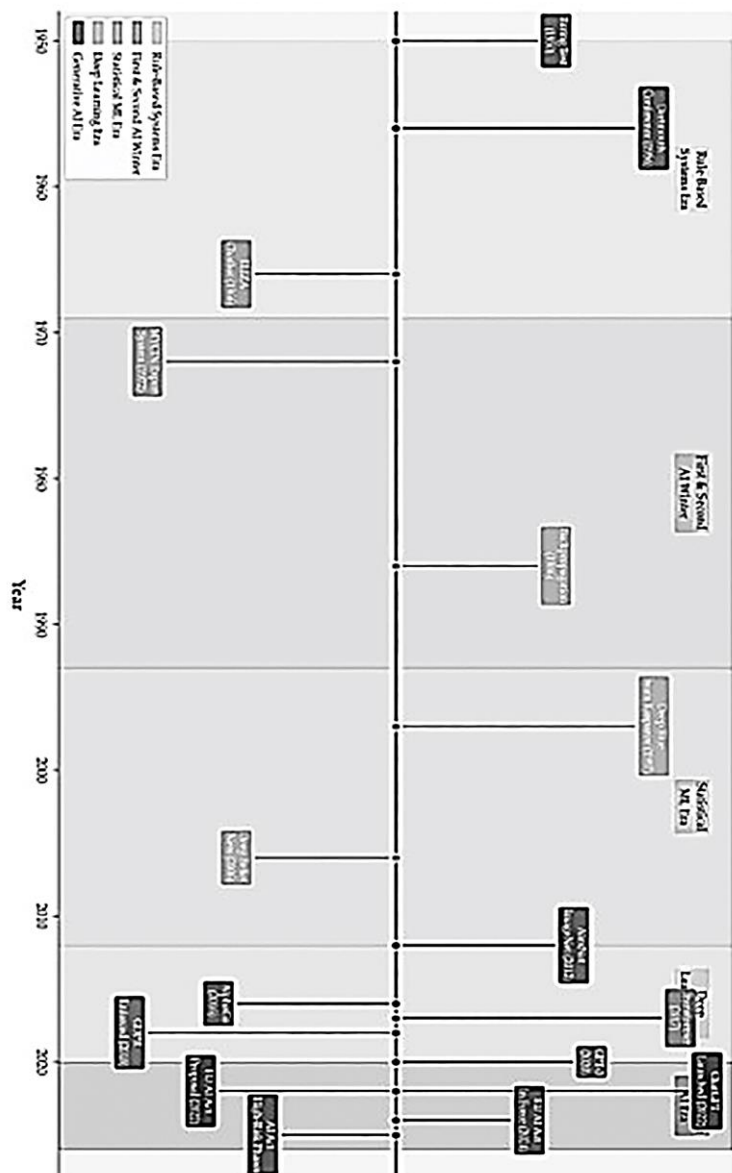


Figure 1.1 – Artificial Intelligence: Key Milestones and Legal Governance Events (1950–2026)

Figure 1.1. Key milestones in AI development and legal governance,

Three generations of AI and three legal problems

The history above is more than chronology. It describes three distinct technical paradigms, each carrying a different legal profile. **Rule-based and expert systems** (roughly 1960– 1990) represented knowledge as explicit conditional rules. MYCIN had approximately 600 rules of the form *"IF the organism is gram-positive AND the morphology is coccus, THEN consider streptococcus."* The legal advantage was transparency. A decision could be reconstructed and disputed because the reasoning was explicit. The legal disadvantage was brittleness. Expert systems failed outside their programmed domain and required constant manual updating. They were also built on assumptions provided by their developers rather than patterns learned from data, which meant that their biases were the biases of their creators, visible in the rules themselves. **Statistical machine learning** (roughly 1990–2015) trained models on historical data to find patterns. A loan model might learn from thousands of past loan records, identified only as outcomes- repaid or defaulted- without anyone specifying which features mattered. Credit scoring, fraud detection, hiring filters, and medical diagnosis tools all adopted this approach. The legal problems shifted accordingly. No longer could a reviewer examine a rule set. The model had learned associations between variables in ways that could not be fully narrated. When those associations produced discriminatory outcomes as happened with Amazon's recruiting tool, which the company quietly shelved in 2018 after discovering it had learned to penalise CVs that included words associated with women- the harm was real but the explanation was elusive.⁶ **Generative and foundation model AI** (2020– present) operates differently again. These systems are not trained for a single task but for general language or image generation across many contexts. They can be fine-tuned, prompted, or augmented for specific uses. A single model can draft a legal brief, generate a contract,

translate a document, summarise a case, analyse a medical image, and write code. Their outputs may be fluent, plausible, and wrong. The legal category most directly challenged is evidence: these systems can produce documents, citations, images, and transcripts that appear authentic but are not. *Mata v. Avianca*, the 2023 case in which lawyers submitted a brief citing six non-existent cases generated by ChatGPT, was not an anomaly. It was an early indication of a structural problem in how legal systems will need to authenticate documents in an era of generative AI.⁷ These three generations do not replace each other cleanly. Rule-based systems are still used in compliance monitoring, clinical protocols, and structured decision trees. Statistical models still power most commercial AI deployments. Generative AI is the newest addition, not the only one. Legal analysis must therefore deal with all three, often within the same institution.

Why regulatory delay is built into the problem

Law moves more slowly than technology, but this is not always a policy failure. Careful regulation requires evidence, deliberation, institutional competence, and political agreement. A rushed statutory response to a half-understood technology can create worse problems than the problem it addresses. The Sarbanes Oxley Act's accounting requirements, passed quickly after corporate scandals in 2002, were effective partly because accountancy had stable concepts. AI governance does not have that advantage. The EU AI Act illustrates the timing challenge. The European Commission published its initial proposal in April 2021.⁸ By the time the Act entered into force on 1 August 2024, AI systems had changed substantially. The proposal had been drafted largely around narrow AI systems used in specific, identifiable applications. The final text had to accommodate general-purpose AI, foundation models, and large language models—categories that barely existed when drafting began. The EU AI Act

should be read as a governance achievement given that constraint, not a perfect instrument. Regulatory delay has at least three distinct causes. The first is information asymmetry. Companies that develop AI systems know substantially more about their training data, architecture, failure modes, and internal testing than any regulator or affected user. Some of that knowledge is appropriately protected as commercial confidential information. But the imbalance means that public authorities often write rules based on what they can observe- the output- rather than what matters most- the design choices made upstream.⁹ The second cause is institutional fragmentation. AI systems raise questions simultaneously for data protection authorities, financial regulators, health regulators, employment law bodies, consumer protection agencies, competition regulators, and courts. No single institution has jurisdiction over all the questions that a single AI system might raise. An automated hiring tool might simultaneously implicate GDPR, employment discrimination law, consumer protection obligations, and CCPA if the employer serves California residents. Coordinating across those regulatory domains is not impossible, but it requires deliberate effort that often does not happen. The third cause is the difficulty of writing technology neutral rules that are not empty. A rule prohibiting "unfair automated decisions" is technology-neutral but tells organisations very little about what to do. A rule requiring specific explainability methods for neural networks is technically precise but may become obsolete when architectures change. The EU AI Act attempts a middle path by specifying objectives- risk management, technical documentation, human oversight, transparency- rather than technical methods. Whether that approach proves robust as the technology evolves is a question this book cannot definitively answer, but the chapters that follow examine the evidence.

When AI tests legal categories

Law organises social life through categories: person, property, contract, harm, fault, consent, authorship, evidence. These categories are not permanent facts about the world. They are legal tools developed to solve recurring human problems, and they have been revised when circumstances required it. The invention of corporations as legal persons, the extension of property rights to intellectual works, the development of products liability- each represented a significant rethinking of how a legal category applied to changed conditions. AI requires comparable work. Authorship is the clearest current example. Copyright law in most jurisdictions assumes human creativity as the source of protectable expression. The United States Copyright Office has stated that works created entirely by a machine, without human creative control, are not eligible for copyright.¹⁰ The UK and EU frameworks take similar positions. But the line between "created by AI" and "created with AI assistance" is contested, fact-dependent, and technologically unstable. As generative systems become more capable and more integrated into creative workflows, the authorship question will press harder against traditional copyright doctrine. Liability is equally strained. When a pharmaceutical company produces a drug that causes harm, product liability law identifies a responsible actor, a defective product, and a traceable causal chain. When an AI system produces a recommendation that leads to harm- a medical misdiagnosis, a wrongful credit denial, a police force stop based on facial recognition error- the responsibility may be distributed across developers, data providers, platform operators, deployers, users, and the institution that authorised the deployment. Traditional product liability, designed for defective physical goods, struggles with this distribution. Evidence faces a new problem of authenticity. Courts and legal systems depend on the ability to distinguish genuine

documents from fabricated ones. Generative AI makes that distinction more difficult at scale. A generated image, contract, transcript, or email can be made to appear authentic. The legal system's existing tools document authentication, expert witnesses on forgery, chain of custody rules were not designed for this volume or quality of synthetic content. Consent in data protection law is built on the assumption that a person can meaningfully authorise specific uses of their information. AI systems challenge that model because data useful for training in one context becomes a generalised model capability that may be applied in many other contexts. A person who consented to their medical records being used for a specific diagnostic study may not have contemplated that the resulting model would be deployed in an insurance context, or that re-identification from supposedly anonymised data would become possible as the model improved. **Human judgment** in legal adjudication assumes a decision-maker who can listen, reason about evidence, apply principles, and be held accountable. AI-assisted decision-making distributes that judgment across a technical system, the people who designed it, the institution that deployed it, and the human reviewer who may or may not engage meaningfully with the recommendation. None of these actors is quite the decision-maker that traditional accountability concepts assume.

International dimensions and the problem of jurisdiction

AI governance cannot be understood through one national system. A model may be developed in the United States, trained using data scraped from multiple countries, hosted on servers in Ireland, and deployed in a welfare agency in Kosovo. At each of those steps, different legal regimes may apply, different rights may be at stake, and different enforcement powers may or may not be available. The

European Union has become the most influential jurisdiction in global AI governance for reasons that parallel its influence in data protection. GDPR's extraterritorial scope- it applies to any organisation processing data of EU residents, regardless of where the organisation is established- has set a precedent that the EU AI Act follows.¹¹ The Act applies to AI systems placed on the EU market or affecting persons in the EU, whether developed inside or outside the EU. This extends EU governance standards to global developers who wish to access the European market, producing what Anu Bradford has called the Brussels effect: the tendency of EU regulation to become a de facto global standard when companies find it more efficient to design one product for the most demanding market.¹² The United States has not adopted a comparable comprehensive AI statute. Its approach is sectoral, agency-driven, and litigation-dependent. The Federal Trade Commission can use its authority over deceptive and unfair practices to address misleading AI claims. Financial regulators can apply fair lending law to algorithmic credit scoring. Employment agencies can apply non-discrimination requirements to algorithmic hiring tools. Courts can develop tort law. But there is no unified framework, and the coverage is uneven. China's model is more centralised. The Cyberspace Administration of China has issued measures on algorithmic recommendation systems, on deep synthesis, and on generative AI, all within a governance framework in which state priorities- social stability, content control, and national security- are explicitly integrated into regulatory objectives.¹³ The resulting framework is coherent but not rights- protective in the sense that liberal constitutional systems understand rights. For Kosovo and other Western Balkan jurisdictions, the situation is distinctive. These jurisdictions are aligned with European frameworks in data protection and digital governance through their European integration processes, but they may lack the administrative capacity, regulatory expertise, or judicial specialisation to enforce those frameworks

effectively. They will also be affected by AI systems designed and deployed elsewhere, without having had significant influence over their design. This position- of a jurisdiction that imports both the technology and the governance obligations but lacks the leverage to shape either is one this book treats as genuinely important rather than peripheral.

Case study: COMPAS and the limits of numerical authority

The COMPAS case merits extended treatment because it shows all three of the core problems identified in this chapter: governance without transparency, liability without clear responsibility, and a due process claim that existing doctrine could not fully resolve. COMPAS assigned recidivism risk scores on a scale of 1 to 10, derived from around 137 input variables. The score was used in sentencing, parole, and pretrial release decisions in multiple US states. ProPublica's 2016 analysis found that Black defendants were almost twice as likely as white defendants to be assigned a higher risk score than they actually deserved, while white defendants were more likely to be assigned a lower score than warranted.¹⁴ Northpointe (the company behind COMPAS, later renamed Equivant) disputed the analysis, arguing that the tool was well-calibrated- that is, that a score of 7, for example, meant roughly the same probability of reoffending regardless of race. The dispute was genuine: the two claims were not contradictory, because it is mathematically impossible to simultaneously satisfy calibration and equal false positive/false negative rates across groups when base rates differ. The disagreement was about which fairness criterion should be prioritised, not about the data.¹⁵ The Loomis decision is often cited as validating COMPAS use in criminal justice. A closer reading reveals something more limited. The Wisconsin Supreme Court accepted the use of the tool as one factor, with the

proviso that the judge had not relied on it as the determinative factor. It did not hold that COMPAS was accurate, fair, or constitutionally above challenge. It held only that the particular sentencing decision in this case survived scrutiny. The decision left open the question of what happens when a system like COMPAS is used more aggressively, or when an affected person can establish that the score did play a determinative role. What the case demonstrates for the purposes of this book is the institutional problem of numerical authority. A score presented as data appears objective. A judge who relies on a number produced by a validated statistical tool may believe they are being more rigorous, not less. But rigour requires understanding what the number means, what assumptions built it, and what fairness criteria it reflects. When that understanding is absent and a closed proprietary system ensures it is absent numerical authority fills the explanatory gap without providing genuine justification.

Case study: facial recognition and wrongful identification

The wrongful arrest of Robert Williams in Detroit in January 2020 is the best-documented example of how facial recognition error translates into harm.¹⁶ Michigan State Police used facial recognition software that returned Williams as a possible match in an investigation. Police officers arrested Williams at his home, in front of his family, and held him for thirty hours before releasing him. No further investigation was conducted before the arrest. The system's output was treated not as a lead requiring verification but as a basis for detention. Williams is Black. Research published by the National Institute of Standards and Technology (NIST) had established by 2019 that several commercially deployed facial recognition algorithms had materially higher false positive rates for darker-skinned faces,

particularly darker-skinned women.¹⁷ The vendors of these algorithms had not always disclosed this performance differential to law enforcement customers. Agencies that deployed the tools had not always tested for it. The decision to rely on a match without independent verification created the conditions for the arrest. The legal analysis requires distinguishing three separate failures. The first was technical: the algorithm produced an incorrect match. The second was governance: police used the output as a basis for arrest rather than as an investigative lead. The third was systemic: neither the procurement process nor the deployment policy adequately addressed the known demographic performance differential. Fixing only the technical failure improving the algorithm- would not have prevented the arrest if the governance failure remained. Fixing only the governance failure- requiring corroboration- would not have changed the biometric inequity embedded in the system. All three failures interact. This pattern of technical deficiency, governance failure, and systemic inequality operating together will appear repeatedly in the chapters that follow, in different domains and with different legal consequences.

Case study: generative AI and fabricated legal authority

The *Mata v. Avianca* case of 2023 raised a narrower but equally significant problem.¹⁸ Lawyers representing Roberto Mata in a personal injury claim against Avianca airline submitted a brief citing multiple cases in support of legal propositions. When Avianca's counsel checked the citations, several of the cited cases did not exist. Others existed but said nothing remotely resembling what the brief attributed to them. The lawyers had used ChatGPT to conduct legal research. ChatGPT had generated plausible sounding but fabricated case names, citations, and quotations, a phenomenon now called "hallucination" in the AI research literature, though the term

understates the legal seriousness of the problem. Judge Kevin Castel of the Southern District of New York imposed sanctions on the lawyers involved.¹⁹ The order noted that a legal citation is not ordinary text. It is a claim about institutional authority representing the court that a prior court decided a relevant question in the way described. A fabricated citation is therefore not merely an error. It is a misrepresentation to the court. The case became widely discussed as a cautionary tale about using generative AI without verification. That framing is accurate but incomplete. The deeper issue is structural. Legal research has always required verification checking that a cited case says what you claim says that it has not been overruled, and that it is from a court whose authority applies. Verification is not an optional extra for AI- assisted research. It is the minimum professional duty. The problem is that generative AI makes non-verification easier to rationalise: the output looks authoritative, the case names sound real, and the volume of work may seem to justify shortcuts. The institutional response will need to include updated professional standards, law school curricula that teach AI-specific verification practices, and court rules that address AI- generated submissions. Several bar associations and courts had issued preliminary guidance by early 2024. That guidance is welcome but does not resolve the underlying difficulty that generative AI makes fabrication cheap and detection expensive.

AI as institutional infrastructure

The three case studies above share a structural feature that deserves explicit attention. In each, AI arrived as institutional infrastructure-present in the workflow before governance frameworks had been developed to manage it. COMPAS was licensed by courts and corrections agencies. Facial recognition tools were purchased by police departments through standard procurement processes.

ChatGPT was used by practising lawyers as a research tool. None of these deployments was clandestine. All were, in a sense, routine. Infrastructure becomes invisible when it works and becomes dangerous when invisibility persists even when it fails. A database of legal records is infrastructure. A search engine is infrastructure. A risk scoring system is infrastructure. The governance failure in each case study was not primarily that the technology was misunderstood by one person. It was that the institution had not built the governance structures - policy, training, oversight, accountability appropriate to technology that affected legal rights. This observation shapes the method of the book. Legal analysis of AI cannot begin with the question "what should be prohibited?" and stop there. Most AI use is not prohibited. The governance challenge is to ensure that the AI systems that institutions routinely use are appropriate for the decisions being made, that the people relying on them understand their limitations, that affected persons can contest their outcomes, and that someone can be held accountable when they fail. Those requirements must be built into institutional design, procurement contracts, professional standards, regulatory oversight, and ultimately legal doctrine.

The scale of AI deployment in legal and public institutions

The governance urgency is intensified by the speed and breadth of AI adoption. By 2024, the US Department of Justice had identified over 1,700 algorithmic or AI- assisted decision-support systems in use across federal agencies, a number that reflected voluntary self-reporting and almost certainly understated the actual count.²⁰ The European Union Agency for Fundamental Rights found in a 2022 survey that a majority of member state law enforcement agencies had deployed or were piloting AI tools in operational policing, with

predictive analytics and surveillance applications predominating. In the healthcare sector, a 2023 analysis of Medicare claims data estimated that AI-assisted clinical decision tools influenced treatment decisions for over 200 million patient encounters annually in the United States alone. Courts have not been immune. In England and Wales, HM Courts and Tribunals Service began investing substantially in AI-assisted court scheduling, document processing, and case management from 2019 onwards as part of its court reform programme. Several US states have court-appointed risk assessment tools embedded in bail and sentencing processes as default workflow components - not optional supplements. In Kosovo and neighboring Balkan jurisdictions, AI-powered translation tools, case management software developed by international donor programmes, and procurement analytics tools sourced from major European technology vendors are all in active governmental use, often without jurisdiction-specific governance frameworks having been developed. These deployment figures matter for legal analysis because they mean that AI governance is not a speculative problem about future technology. It is a present institutional condition. The question is not whether to govern AI in legal and public-sector contexts but whether the governance that already nominally applies - through procurement law, administrative procedure, data protection requirements, and professional standards - is functional. The evidence from the case studies discussed in this chapter suggests that, in many institutions, it is not. The mismatch between deployment scale and governance capacity creates a specific regulatory risk. When governance catches up to deployment, it often does so in the form of prohibition or mandatory remediation. The Dutch toeslagenaffaire - the childcare benefit fraud investigation scandal in which an AI system wrongly flagged tens of thousands of families for investigation, leading to enormous financial and personal harm - resulted eventually in the resignation of the Dutch cabinet, financial compensation running to

hundreds of millions of euros, and a parliamentary inquiry that produced the "unprecedented injustice" report in December 2020.²¹ The cost of post-deployment remediation of governance failures vastly exceeds the cost of adequate governance at deployment.

Lessons from analogous technological transitions

AI is not the first technology to disrupt legal categories and governance frameworks. Understanding how law adapted to earlier technological transitions provides both historical context and methodological guidance. The introduction of electronic databases in the 1970s and 1980s created the first generation of computer-related legal questions: who owned data stored in electronic form, how did tort liability apply to software defects, how could evidence in electronic form be authenticated. The US Electronic Signatures in Global and National Commerce Act (2000) and the EU Electronic Commerce Directive (2000/31/EC) were legislative responses that took approximately two decades to mature after the technology became operationally significant. The proliferation of the internet created a second generation: liability for online platforms, intermediary liability, defamation in digital contexts, jurisdiction over cross-border digital transactions, and copyright in a world of easy reproduction. The Digital Millennium Copyright Act (1998), the EU Copyright Directive (2001/29/EC), and eventually the EU Digital Services Act (2022) represent successive legislative responses, each imperfect, each requiring revision as the technology evolved. The distance between the emergence of the problem and the legislative response was, in each case, measured in years or decades. AI governance has compressed this timeline under external political pressure, but the fundamental challenge has not changed. Law must generalise across cases that have not yet occurred, using concepts developed for conditions that AI changes. The EU AI Act's risk-based tiering

prohibited, high-risk, limited-risk, general-purpose represents current best legislative practice. It is also, as noted above, a snapshot of AI capabilities and deployment patterns as they appeared during the drafting process. The Act contains explicit review mechanisms, including a requirement that the Commission evaluate and update Annex III (the list of high-risk use cases) periodically. Whether those mechanisms prove adequate depends on whether the governance institutions that implement the Act develop genuine technical capacity. Pharmaceutical regulation offers a more successful historical analogy. Pre-market safety assessment for drugs was not always standard. The thalidomide disaster of the late 1950s and early 1960s—in which a drug approved in multiple countries caused severe birth defects produced what became modern pharmaceutical regulation: controlled trials, pre-market approval requirements, post-market surveillance, and mandatory reporting of adverse events.²² The analogy is not perfect: pharmaceutical regulation can assess a single drug entity whose effects can be tested in controlled conditions, while AI systems are context dependent, continuously evolving, and interact with human behaviour in ways clinical trials cannot replicate. But the underlying regulatory logic require demonstrated safety before authorization, maintains monitoring after authorization, require withdrawal when safety evidence deteriorates maps onto AI governance as a model worth examining.

The Western Balkans and the governance capacity problem

Throughout this book, the Western Balkans appears not as a footnote but as a specific governance context that illustrates a problem broader than the region itself. Many jurisdictions in Europe, Asia, Africa, and the Americas are adopting AI in legal and public institutions while operating under three simultaneous constraints: they are bound by

governance frameworks developed primarily in larger, wealthier jurisdictions; they lack the administrative capacity to implement those frameworks fully; and they are largely dependent on AI systems designed for other markets whose performance in local conditions has not been validated. Kosovo's legal system illustrates the pattern. As a jurisdiction in the process of European integration, Kosovo is progressively aligning its legislation with EU law, including the GDPR transposition through the Law on Protection of Personal Data. AI governance will similarly require alignment with the EU AI Act as integration progresses. But alignment on paper does not create the supervisory authority capacity, the judicial expertise, or the institutional culture that makes governance effective in practice. The EU AI Act itself recognises capacity constraints. Articles 64– 70 establish the AI Office and the European Artificial Intelligence Board as coordination and support mechanisms, and the AI Office established within the European Commission is intended to provide technical expertise that smaller member states and candidate countries may lack. Whether these mechanisms prove adequate for candidate jurisdictions depends on how the AI Office functions in practice a question that was not fully answerable at the time this book was completed. What can be said with confidence is that the governance problem for smaller jurisdictions is not simply a smaller version of the governance problem for large ones. It has specific features: AI systems are typically designed and validated for markets very different from those in which they are deployed; algorithmic systems imported through international aid or vendor contracts may not come with adequate documentation in local languages; supervisory authorities may have technical expertise in data protection but not in the machine learning methods that constitute most modern AI; and judicial systems may encounter AI- related legal questions before relevant precedent has been established. These features require governance responses tailored to capacity constraints

rather than simply transpositions of governance frameworks designed for the most sophisticated regulatory environments.

What governance architecture requires

The three case studies and the governance scale problem all point toward the same analytical conclusion: AI governance requires architecture, not improvisation. Architecture means designed structures - procedural, institutional, and legal - that produce governance outcomes systematically across the range of AI systems in use, rather than remediation of specific failures after they become visible. Architectural governance has several components. **Prospective coverage** ensures that governance applies before AI systems are deployed in consequential contexts, rather than only after harm has been identified. The EU AI Act's conformity assessment requirements for high-risk systems are the clearest current example of prospective governance at scale. Institutional capacity provides the technical and legal expertise needed to assess AI systems across the range of their deployment contexts. No regulatory body can be expert in all domains in which AI systems operate, so capacity-building must include sector-by-sector expertise in health, justice, financial services, employment, and public administration. Accountability chains ensure that when AI-assisted decisions harm individuals, responsibility can be traced to identifiable actors who can be required to make good and held accountable for future conduct. This requires supply-chain transparency, documentation obligations, and contractual allocation of responsibility that currently exists only partially and unevenly. Redress mechanisms ensure that people harmed by AI systems have accessible and effective means of seeking remedy - not merely theoretical legal rights that are practically inaccessible because of evidence problems, costs, or jurisdiction complexity.²³ These components are not new legal concepts. They are applications of

governance principles - legality, accountability, proportionality, effectiveness - to a new institutional setting. What is new is the specific implementation challenge: governance structures must be technically literate, because AI systems have properties that legal concepts alone cannot capture, and technically accurate, because a governance framework that describes AI systems incorrectly will be gamed or misapplied.

The method of this book

This book operates from a methodological premise: legal analysis of AI must be grounded in technically accurate description of the system being discussed. A rule-based expert system, a logistic regression model, a deep neural network, and a large language model create different risks, require different governance responses, and raise different legal questions. Treating all of them as "AI" and writing rules about "AI systems" may be inevitable at a level of abstraction, but it is insufficient for detailed analysis. At the same time, technical accuracy cannot substitute for legal reasoning. A technically accurate description of how COMPAS works does not tell you whether its use in sentencing is constitutionally permissible. A technically accurate description of how facial recognition matches are generated does not tell you what standard of corroboration police should require before acting on a match. Those are normative questions about rights, fairness, accountability, and institutional design that technical description cannot resolve. The book therefore proceeds by combining technical explanation with doctrinal analysis. Each chapter explains the relevant technical mechanisms before examining the legal frameworks that apply, the cases where courts or regulators have reached conclusions, and the areas where the law remains underdeveloped or contested. The method is comparative where the legal frameworks differ significantly across jurisdictions and on AI