

ARTICLE



REVISTA
COLETA CIENTÍFICA

Tort Liability for Damages Caused by Artificial Intelligence Systems: Current Challenges and Future Prospects

Chaabna Imen¹

<https://orcid.org/0009-0002-5481-7051>
University of 20 August 1955 Skikda, Algeria
E-mail: i.chaabna@univ-skikda.dz

Publication information

ARK: [24285/RCC.v10i20.253](https://doi.org/10.24285/RCC.v10i20.253)

ISSN: 2763-6496

Keywords:

Artificial intelligence.
Tort liability.
Causation.
Strict liability.
Compulsory insurance.

Abstract

The contemporary era is marked by an unprecedented technological revolution in the field of artificial intelligence (AI), with autonomous systems progressively integrated into critical sectors including public administration, healthcare, transportation, and essential infrastructure in Algeria and globally. This massive integration raises profound legal challenges, particularly concerning the application of classical tort liability frameworks to autonomous and opaque algorithmic entities, notably deep learning systems that evolve beyond their initial programming. This article examines, through a comparative perspective, the adequacy of the Algerian legal framework (Articles 124, 125, 138, and 140 bis of the Algerian Civil Code), French law (Article 1242, paragraph 1, of the French Civil Code), and European Union regulations (the proposed EU AI Liability Directive) in addressing damages caused by AI systems. Employing a qualitative and comparative analytical methodology, the study demonstrates the structural insufficiency of the traditional elements of tort liability—fault (*faute*), causation (*lien de causalité*), and subjective accountability—when confronted with the decision-making autonomy and algorithmic opacity of the "black box." The article concludes that an imperative legislative reform is necessary to establish a strict liability regime grounded in risk theory, coupled with the implementation of a mandatory insurance system to guarantee full victim compensation and legal certainty in the context of digital transformation.

¹ Lecturer Professor (Class A), Faculty of Law and Political Science, University of 20 August 1955 – Skikda, Algeria (2012); a Master's degree in Law Faculty of Law, University of Frères Mentouri Constantine 1, Algeria (2012); a PhD in Law University of Frères Mentouri Constantine 1, Algeria (2017).

Introduction

The contemporary global landscape is undergoing a radical and profound transformation in the nature and functionality of the devices, machines, and systems that human beings deploy in their daily, professional, and institutional activities. Whereas mechanical and electronic tools were historically conceived as inert, passive instruments wholly subordinate to human command and control—executing predefined, static instructions without deviation or autonomy—the advent and exponential proliferation of artificial intelligence (AI) systems have fundamentally disrupted this paradigm. Modern AI-driven entities, particularly those employing machine learning, deep learning, and neural network architectures, are no longer merely passive executors of human will. Instead, they constitute advanced, autonomous systems endowed with unprecedented capacities for self-directed learning, real-time environmental adaptation, autonomous decision-making, and continuous self-modification and updating, often beyond the anticipation, comprehension, or control of their original programmers, manufacturers, or end-users (Russell & Norvig, 2021).

This technological revolution has not remained confined to experimental laboratories or niche applications. On the contrary, AI systems have permeated critical, high-stakes domains of public and private life, including but not limited to: advanced medical diagnostics and robotic surgery, autonomous and semi-autonomous vehicular transportation, algorithmic financial trading and credit risk assessment, predictive policing and criminal justice decision support, administrative automation in public governance, and complex industrial manufacturing processes. In Algeria, as in many states across the Middle East and North Africa (MENA) region, governmental institutions, healthcare facilities, transportation networks, and foundational infrastructures are progressively integrating AI-driven solutions to enhance efficiency, reduce costs, and improve service delivery (Khaloui, 2025; Alioui, 2023).

Yet this accelerating integration engenders a corollary and inevitable reality: the systematic emergence of errors, malfunctions, biases, unpredictable behaviors, and resultant harms attributable to AI systems. When an autonomous surgical robot commits an intraoperative error, when an algorithmic credit-scoring system unjustly denies a loan based on opaque criteria, when a self-driving vehicle causes a fatal collision, or when a predictive policing algorithm perpetuates systemic discrimination, a fundamental legal question arises with acute urgency: **Who bears civil liability for the resulting damages?**

This question exposes a profound and structural disjunction between the foundational premises of classical civil liability law and the operational realities of autonomous AI systems. Traditional tort law frameworks—whether rooted in the Romano-Germanic civil law traditions of Algeria and France, or in common law jurisdictions—were architecturally designed and jurisprudentially refined to address human conduct. The doctrinal edifice of tort liability rests upon the normative and ontological assumption that the agent causing harm is a natural or legal person endowed with legal capacity (*capacité juridique*), intent (*intention*), discernment (*discernement*), and the ability to commit fault (*faute*) through a volitional act or omission that deviates from a legally prescribed standard of conduct (Viney & Jourdain, 2017).

However, AI systems—particularly those employing deep learning and reinforcement learning algorithms—operate fundamentally differently. They function

as algorithmic "black boxes": complex, multi-layered computational architectures whose internal decision-making processes are frequently opaque, non-linear, and epistemically inaccessible even to their designers (Burrell, 2016). These systems autonomously generate outputs, predictions, classifications, and actions based on probabilistic pattern recognition across vast datasets, continuously modifying their internal parameters through iterative learning cycles that occur post-deployment, often in environments and under conditions unforeseen by the original programmers. Consequently, when harm occurs, the traditional evidentiary pathways for establishing fault, causation, and liability become obstructed or entirely severed (Calo, 2015).

In this context, the Algerian Civil Code—promulgated by Order No. 75-58 of September 26, 1975, and subsequently amended—provides the normative foundation for civil liability. Articles 124 and 125 establish the regime of personal fault-based liability (*responsabilité pour faute personnelle*), requiring the demonstration of: (1) a material element (*élément matériel*), consisting of conduct that deviates from the legally required standard, and (2) a moral or intellectual element (*élément moral*), comprising discernment and intent, such that individuals lacking such capacity (minors, the insane, the mentally incapacitated) are exempted from liability (Algerian Civil Code, 1975; Debbache, 2019). Article 138 enshrines the principle of liability for things under one's custody (*responsabilité du fait des choses que l'on a sous sa garde*), a doctrinal descendant of French Article 1384 (now Article 1242, paragraph 1, of the reformed French Civil Code of 2016), which imposes liability upon the "custodian" (*gardien*) of a thing that causes damage, provided the custodian exercises use, control, and direction (*usage, contrôle, et direction*) over the object (Malouhi, 2009).

Yet the application of these venerable principles to autonomous AI systems reveals fundamental conceptual and operational failures. First, AI systems do not possess legal personality (*personnalité juridique*): they are neither natural persons nor juridical entities recognized under civil law. Consequently, they cannot be attributed fault, intent, or discernment. Second, the notion of custody (*garde*) presumes the custodian's continuous and effective control over the object's behavior—a condition systematically violated by autonomous AI systems that learn, adapt, and modify their conduct independently and unpredictably post-deployment. Third, the establishment of a causal nexus (*lien de causalité*)—a sine qua non of civil liability—becomes a legally insurmountable challenge when the decision-making process is shrouded in algorithmic opacity, rendering it practically and epistemologically impossible for a victim to trace the precise chain of causation from the programmer's initial code, through iterative machine learning cycles, to the specific harmful output (Scherer, 2016).

Furthermore, Algeria, like France, has adopted a product liability regime codified in Article 140 bis of the Algerian Civil Code, which transposes principles analogous to the French Articles 1245 et seq. and the European Product Liability Directive (85/374/EEC, amended by Directive 1999/34/EC). This regime permits victims to hold manufacturers liable for damages caused by defective products without proving fault, provided they demonstrate the defect, the damage, and the causal link. However, this framework encounters serious limitations when applied to AI: the development risk defense (*risque de développement*), which exonerates manufacturers for defects unknowable at the time of circulation, becomes a nearly universal shield for AI developers, given the inherent

unpredictability and emergent properties of machine learning systems (European Commission, 2020).

Internationally, the European Union has embarked on a comprehensive regulatory response. The EU Artificial Intelligence Act (Regulation (EU) 2024/1689, entered into force in August 2024) establishes a risk-based classification system, prohibiting certain high-risk AI applications and imposing stringent transparency, accountability, and conformity assessment obligations on developers and deployers of high-risk AI systems (European Parliament and Council, 2024). Additionally, the proposed EU AI Liability Directive (COM/2022/496 final) seeks to introduce rebuttable presumptions of causality and fault, reducing the evidentiary burden on victims and facilitating compensation claims (European Commission, 2022). These instruments reflect a growing doctrinal consensus that the classical fault-based liability paradigm must be supplemented—or even supplanted—by strict liability, risk-based liability, or objective liability regimes (*responsabilité objective*), which prioritize victim compensation and systemic risk distribution over the penalization of individual fault (Lohsse et al., 2019; Wagner, 2020).

Against this backdrop, the present article undertakes a systematic comparative analysis of the adequacy of the Algerian, French, and emerging EU legal frameworks in addressing tort liability for AI-generated damages. The central problematic guiding this inquiry is: To what extent are the general principles and specific provisions of civil liability law suitable and sufficient to address the challenges posed by the risks and harms generated by autonomous AI systems? Answering this question is of paramount normative and practical importance: it concerns the fundamental right to legal protection and just compensation for victims, irrespective of whether the harmful agent is a human actor or an algorithmically governed intelligent system. Without a clear, suitable, and comprehensive liability regime, victims of AI-related harms may find themselves unable to vindicate their rights due to insurmountable technical evidentiary barriers, thereby undermining the foundational principles of justice, accountability, and the rule of law.

To address this problematic, the article is structured into two principal substantive sections. First, it critically examines the challenges inherent in applying traditional civil liability doctrines to AI systems, dissecting the structural incompatibilities between classical liability elements (fault, causation, capacity, custody) and the operational realities of autonomous, opaque, and self-modifying intelligent systems. Second, it explores emerging legal frameworks and future prospects for protecting victims of AI-generated harms, focusing on the doctrinal and policy shift toward objective liability, strict product liability, and mandatory insurance mechanisms as instruments for ensuring victim compensation, distributing risk, and fostering responsible technological innovation.

This inquiry is animated by a methodological commitment to rigorous comparative legal analysis, drawing upon legislative texts, judicial precedents, doctrinal scholarship, and policy instruments from Algeria, France, and the European Union, as well as insights from the broader international discourse on AI governance, ethics, and law. By illuminating the deficiencies of existing legal frameworks and charting pathways toward reform, this article seeks to contribute to the urgent and ongoing global effort to reconcile the imperatives of technological progress with the enduring demands of justice, accountability, and human dignity.

Theoretical Foundation and Literature Review

The advent of artificial intelligence as an autonomous, decision-making, and self-learning technological phenomenon has ignited a vigorous and multifaceted scholarly debate regarding its legal status, moral agency, and the appropriate liability regimes to govern its deployment and the harms it may generate. This debate traverses doctrinal, philosophical, economic, and policy dimensions, and has produced a rich and rapidly expanding body of literature that can be organized chronologically and thematically into several key conceptual clusters.

1. The Contested Question of Electronic Legal Personality

A foundational and fiercely contested issue concerns whether AI systems—particularly highly autonomous robots and algorithmic agents—should be granted legal personality (*personnalité juridique*). Legal personality is the normative mechanism through which an entity is recognized as a subject of rights and obligations, capable of holding assets, entering contracts, suing and being sued, and bearing civil or criminal liability (Teubner, 2006).

Proponents of granting legal personality to AI systems argue that the increasing autonomy, decision-making capacity, and societal integration of intelligent machines warrant their recognition as a new category of legal subject, akin to corporations or other juridical persons. This position was notably articulated in the European Parliament's 2017 Resolution on Civil Law Rules on Robotics (2015/2103(INL)), which recommended that the European Commission consider creating a specific legal status for robots, particularly the category of "electronic persons" (*personnes électroniques*), to ensure clarity in liability allocation and facilitate compensation for victims (European Parliament, 2017).

However, this proposal has encountered substantial doctrinal resistance. Leading civil law scholars, jurists, and policymakers have raised profound objections on normative, practical, and ethical grounds. First, legal personality is inextricably linked to moral agency, intentionality, and capacity for rights and duties—attributes that AI systems, however sophisticated, fundamentally lack. AI systems do not possess consciousness, intentionality, subjective interests, or the capacity for moral reasoning; they are sophisticated instruments designed, programmed, and deployed by human actors to achieve human objectives (Solaiman, 2017). Second, granting legal personality to AI could function as a strategic liability shield for manufacturers, programmers, and deployers, enabling them to externalize risks and evade responsibility by attributing harm to the "autonomous decisions" of the AI entity itself. This would create a moral hazard, incentivizing negligent design, inadequate testing, and insufficient safeguards, while systematically denying victims full and fair compensation (Al-Qousi, 2018, pp. 72, 77). Third, legal personality traditionally entails patrimonial capacity—the ability to hold assets and satisfy judgments. An AI system, lacking patrimony, could not satisfy a compensation claim, rendering legal personality a hollow formalism devoid of practical utility for victims (Khaloui, 2025, p. 493).

Consequently, the dominant view within contemporary civil law scholarship and regulatory policy is that AI systems should remain classified as "things" (*choses*) or "products" rather than persons, with liability allocated among human and corporate actors in the design, manufacturing, deployment, and operational chain (Bertolini, 2020).

This position is reflected in the EU AI Act, which eschews the concept of electronic personality and instead focuses on risk-based regulatory obligations imposed on identifiable human and corporate actors (European Parliament and Council, 2024).

2. The Limitations of the Custody Framework (*Responsabilité du Fait des Choses*)

Given the prevailing doctrinal consensus that AI systems are "things" rather than persons, a natural recourse is to apply the well-established regime of liability for things under one's custody (*responsabilité du fait des choses que l'on a sous sa garde*), codified in Article 1242, paragraph 1, of the French Civil Code (formerly Article 1384, paragraph 1) and Article 138 of the Algerian Civil Code.

This regime, developed and refined through over a century of French jurisprudence—most notably the landmark *Cour de Cassation* decisions in *Jand'heur v. Galeries Belfortaises* (February 13, 1930, *Dalloz Périodique* 1930.1.57) and subsequent rulings—establishes that the custodian (*gardien*) of a thing is strictly liable for damages caused by that thing, regardless of fault, provided the victim establishes the intervention of the thing (*fait de la chose*) in causing the harm (Viney & Jourdain, 2017). The custodian is defined as the person who exercises use, control, and direction (*usage, contrôle, et direction*) over the thing, whether by virtue of ownership, lease, or other legal or de facto relationship (Malouhi, 2009, pp. 122–123).

However, the application of the custody framework to autonomous AI systems encounters fundamental conceptual and operational obstacles:

a) The Disintegration of Effective Control

The classic custody framework presupposes that the custodian exercises continuous, effective, and predictable control over the thing's behavior. This presumption is systematically violated by autonomous AI systems, which are designed precisely to operate independently of direct human oversight, to learn from environmental data, and to modify their decision-making algorithms autonomously and unpredictably (Calo, 2015). Once deployed, a machine learning system may encounter novel scenarios, retrain itself, adjust its internal parameters, and generate outputs that diverge fundamentally from its initial programming—outputs that may be epistemically inaccessible even to the original developer. Consequently, the notion that a programmer, manufacturer, or end-user exercises "control" over the AI's behavior becomes a legal fiction divorced from technical reality (Scherer, 2016).

b) The Fragmentation of Custody Among Multiple Actors

The lifecycle of an AI system typically involves multiple actors: data providers, algorithm designers, software developers, hardware manufacturers, system integrators, deployers, operators, and end-users. Each actor contributes to different stages of the system's creation and operation, and each may exercise partial but incomplete control over different aspects of the system's functionality. This fragmentation of custody raises acute questions: Who is the "custodian" when a self-driving car, designed by one company, using sensors manufactured by another, operating software developed by a third, and processing data provided by a fourth, causes a collision? The classical custody framework, predicated on a unitary custodian, offers no clear guidance for allocating liability among multiple, distributed contributors to a complex, autonomous system (Hubbard, 2014).

The Problem of "Causal Indeterminacy" and the Black Box

Even if one could identify a custodian, establishing the causal link between the custodian's conduct (or failure to act) and the specific harmful output remains a formidable challenge due to the algorithmic opacity of AI systems. Deep learning neural networks, for example, consist of millions or billions of weighted parameters adjusted iteratively through training on massive datasets. The resulting decision-making process is non-transparent, non-linear, and non-interpretable: neither the developer nor external experts can reliably explain why the system generated a particular output in a specific instance (Burrell, 2016). This opacity renders it practically impossible for a victim to satisfy the evidentiary burden of proving that a specific defect, omission, or negligence by the custodian caused the harm—a burden that, under traditional civil procedure, rests upon the plaintiff (*probatio diabolica*) (European Commission, 2020).

French jurisprudence has occasionally sought to mitigate this burden by recognizing rebuttable presumptions of causation when a thing "intervenes" in causing harm, but this mechanism presupposes a relatively transparent causal mechanism and a custodian exercising meaningful control. In the context of autonomous AI, these conditions are systematically absent.

3. The Black Box Dilemma and the Fracture of Causation

The metaphor of the "black box" has become central to the legal and ethical discourse on AI liability. It captures the fundamental epistemological challenge posed by AI systems whose internal workings are opaque, complex, and inaccessible to external scrutiny (Pasquale, 2015).

a) Technical Opacity

Modern AI systems, particularly those employing deep learning, function as multi-layered neural networks trained on vast datasets. The training process adjusts billions of internal parameters to optimize predictive accuracy, but the resulting model does not produce human-readable rules or explanations. Consequently, when an AI system generates a harmful output—such as an erroneous medical diagnosis, a discriminatory credit decision, or a collision by an autonomous vehicle—it is frequently impossible to reconstruct the precise causal chain that led from input data, through the network's internal transformations, to the harmful output (Selbst & Barocas, 2018).

b) Legal Implications: The Rupture of the Causal Link (*Lien de Causalité*)

Civil liability, whether fault-based or strict, requires the establishment of a causal nexus between the defendant's conduct (or the defect in a product) and the plaintiff's harm. In classical tort law, this is typically established through expert testimony, documentary evidence, or judicial inference. However, the black box nature of AI systems renders such proof extraordinarily difficult or impossible:

- **For fault-based liability** (Articles 124–125 Algerian Civil Code), the victim must prove that the programmer or operator committed a specific fault (negligence, imprudence, violation of a standard of care) that caused the harm. But if the harmful output resulted from autonomous learning and adaptation post-deployment, how can the victim trace the harm back to a specific human fault?

- **For product liability** (Article 140 bis Algerian Civil Code; French Articles 1245 et seq.), the victim must prove a defect in the product. But what constitutes a "defect" in a machine learning algorithm that functions as designed but produces an unanticipated harmful output due to emergent properties of its training data or learning process?

This rupture of causation systematically undermines the victim's ability to satisfy the evidentiary requirements of classical liability regimes, creating a legal vacuum in which harm occurs but no party can be held accountable (Abbott, 2020).

4. From Fault-Based to Risk-Based Liability: Philosophical and Policy Foundations

Faced with these profound challenges, legal scholars and policymakers have increasingly advocated for a paradigm shift from fault-based liability (*responsabilité pour faute*) to risk-based or objective liability (*responsabilité objective*) (Wagner, 2020).

a) Fault-Based Liability: Foundations and Limitations

Fault-based liability is grounded in the moral and corrective justice principle that individuals should bear the costs of harms they wrongfully inflict upon others. It presupposes moral agency, intentionality, and the capacity for fault. However, AI systems lack these attributes, and the complexity and opacity of their operation render human fault difficult or impossible to prove (Honoré, 2010).

b) Risk-Based Liability: Foundations and Advantages

Risk-based or strict liability is grounded in the distributive justice principle that those who create, control, or benefit from hazardous activities should bear the costs of the risks they generate, regardless of fault. This principle underpins product liability, liability for abnormally dangerous activities, and workers' compensation schemes (Calabresi, 1970). Applying strict liability to AI offers several advantages:

- **Victim Compensation:** By eliminating the need to prove fault or navigate the black box, strict liability facilitates compensation for victims.
- **Risk Internalization:** It incentivizes developers and deployers to invest in safety, testing, and risk mitigation by ensuring they bear the costs of harm.
- **Allocative Efficiency:** It aligns liability with the parties best positioned to prevent harm and spread risk (manufacturers, insurers) rather than individual victims.

The EU AI Liability Directive (proposed) embodies this shift by introducing rebuttable presumptions of causation and fault for high-risk AI systems, effectively moving toward a strict liability regime (European Commission, 2022).

c) Compulsory Insurance as a Complementary Mechanism

Strict liability alone may be insufficient if liable parties are judgment-proof or if identifying the liable party remains difficult. Consequently, scholars advocate for mandatory insurance schemes modeled on motor vehicle insurance, workers' compensation, or nuclear accident liability (Mohamed, 2025, p. 772). Such schemes ensure victim compensation, distribute risk broadly across society, and de-couple compensation from the often-futile search for fault (Abraham & Rabin, 2019).

Methodology

This research employs a qualitative, comparative, and analytical methodology to systematically examine the adequacy of existing civil liability frameworks in addressing damages caused by autonomous AI systems and to identify pathways for legal reform.

1. Comparative Legal Analysis

The study juxtaposes the legislative texts, doctrinal interpretations, and jurisprudential trajectories of three legal systems:

- **Algerian Civil Law:** Focusing on Articles 124, 125, 138, and 140 bis of the Algerian Civil Code (Order No. 75-58 of 1975, as amended), which govern personal fault-based liability, liability for things under custody, and product liability, respectively.
- **French Civil Law:** Investigative the evolution of Article 1242, paragraph 1, of the French Civil Code (formerly Article 1384) concerning *responsabilité du fait des choses*, the landmark jurisprudence of the *Cour de Cassation* (e.g., *Jand'heur*, 1930), and the product liability regime under Articles 1245 et seq. (transposing EU Directive 85/374/EEC).
- **European Union Regulatory Framework:** Analyzing the EU Artificial Intelligence Act (Regulation (EU) 2024/1689) and the proposed EU AI Liability Directive (COM/2022/496 final), which establish risk-based classifications, transparency obligations, and liability presumptions tailored to AI systems.

Through doctrinal analysis (examining scholarly commentary, treatises, and law review articles), legislative analysis (interpreting statutory texts and regulatory instruments), and jurisprudential analysis (reviewing leading judicial decisions), the study identifies convergences, divergences, gaps, and best practices across these systems.

2. Conceptual and Theoretical Framework

The research is grounded in tort law theory, philosophy of law, and risk regulation theory. It critically engages with foundational concepts such as fault (*faute*), causation (*lien de causalité*), custody (*garde*), strict liability (*responsabilité sans faute*), distributive justice, and risk allocation. By situating the AI liability problematic within these theoretical traditions, the study illuminates the normative and structural incompatibilities between classical tort doctrines and the operational realities of autonomous AI systems.

3. Analytical Approach

The study adopts a problem-driven analytical approach, structured around the following research questions:

- **Q1:** To what extent do the traditional elements of civil liability (fault, causation, damage, capacity) adequately capture and address harms caused by autonomous AI systems?
- **Q2:** What are the structural, evidentiary, and conceptual obstacles that prevent victims of AI-generated harms from obtaining compensation under existing liability regimes?
- **Q3:** What alternative liability frameworks (strict liability, objective liability, presumptions of causation, mandatory insurance) have been proposed or implemented in comparative jurisdictions, and how effective are they likely to be in the AI context?
- **Q4:** What legislative, regulatory, and institutional reforms are necessary to ensure adequate victim protection, incentivize responsible AI development, and promote legal certainty in Algeria and similar civil law jurisdictions?

4. Sources and Data

The research draws upon a comprehensive corpus of primary and secondary legal sources, including:

- **Legislative Texts:** Algerian Civil Code, French Civil Code, EU Regulations and Directives.
- **Judicial Decisions:** Leading French tort law cases (*Cour de Cassation*), European Court of Justice rulings on product liability.
- **Scholarly Literature:** Monographs, peer-reviewed journal articles, and law review essays by leading tort law scholars (e.g., Viney, Jourdain, Calabresi, Wagner) and AI law experts (e.g., Scherer, Calo, Bertolini, Abbott).
- **Policy Documents:** European Commission White Papers, Expert Group Reports, Parliamentary Resolutions.
- **Arab Legal Scholarship:** Including the works cited in the original manuscript (Mohamed, 2025; Khaloui, 2025; Al-Qousi, 2018; Alioui, 2023; Debbache, 2019; Malouhi, 2009).

5. Limitations

This study is primarily doctrinal and comparative in nature. It does not empirically assess the frequency, severity, or distribution of AI-generated harms, nor does it conduct quantitative analysis of litigation outcomes. Future research could usefully complement this doctrinal analysis with empirical studies of AI liability cases, surveys of judicial attitudes, and economic modeling of alternative liability and insurance regimes.

Results and Discussion

This section presents the core findings of the comparative legal analysis, organized into four thematic subsections that progressively demonstrate the systematic failure of classical civil liability frameworks to address AI-generated harms and the necessity of transitioning to risk-based strict liability and mandatory insurance regimes.

Section 1: The Disintegration of the Fault Element in the Context of Autonomous AI Systems

Fault (*faute*) is the conceptual cornerstone of traditional civil liability in both Algerian and French law. It is a normative and psychological construct comprising two essential elements:

1.1. The Material Element (*Élément Matériel*)

The material element consists of conduct—an act or omission—that deviates from the standard of care required by law, regulation, professional norms, or the general standard of a reasonable, prudent person (*bon père de famille*). This may take the form of:

- **Negligence** (*négligence*): failure to exercise due care.
- **Imprudence** (*imprudence*): reckless disregard for foreseeable risks.
- **Violation of a statutory or regulatory duty** (*violation d'une obligation légale ou réglementaire*).
-

1.2. The Moral or Intellectual Element (*Élément Moral ou Intellectuel*)

The moral element requires that the actor possess discernment (*discernement*) and legal capacity (*capacité*). Algerian Civil Code Article 125 explicitly provides that persons lacking discernment—such as minors, the insane, or the mentally

incapacitated—cannot be held liable for their harmful acts, even if those acts objectively deviate from required conduct (Algerian Civil Code, 1975; Debbache, 2019, pp. 26–27). This principle reflects the foundational premise that liability presupposes moral agency: the capacity to understand the nature and consequences of one's actions and to conform one's conduct to legal norms.

1.3. The Impossibility of Attributing Fault to AI Systems

When we attempt to apply the concept of fault to autonomous AI systems, we encounter an insurmountable conceptual and normative barrier: AI systems lack legal personality, consciousness, intentionality, and discernment. They are not moral agents; they are sophisticated computational tools that execute probabilistic inferences based on statistical patterns in data.

Consider the following scenario: A robotic surgical system, employing machine learning algorithms to optimize incision trajectories, makes an erroneous cut during a surgical procedure, resulting in severe injury to the patient. Can we attribute fault to the robot?

- **Material Element:** The robot's action (the erroneous cut) objectively deviates from the required standard (precise, safe surgical incisions). Thus, the material element appears satisfied.
- **Moral Element:** However, the robot lacks discernment, consciousness, and legal capacity. It does not "know" what it is doing in any moral or psychological sense; it executes algorithmic operations devoid of intentionality or understanding. Under Article 125 of the Algerian Civil Code, and analogous French jurisprudence, such an entity cannot be held liable for fault, as it is conceptually equivalent to a minor or insane person lacking discernment.

Therefore, the robot's "error" is not a legal fault. The legal question then shifts: Can fault be attributed to the programmer, manufacturer, or operator of the robot?

1.4. The Evidentiary Impossibility of Proving Human Fault in the Black Box Context

Even if we accept that liability should rest with human actors, the algorithmic opacity of AI systems renders proof of human fault extraordinarily difficult or impossible:

- **Programmer Fault:** The victim must prove that the programmer committed a specific error—such as a coding mistake, failure to test adequately, or violation of industry standards—that caused the harm. However, in a machine learning context, the harmful output may result from **emergent properties** of the training process, interactions among millions of parameters, or autonomous learning post-deployment—none of which can be traced to a specific programming fault (Selbst & Barocas, 2018).
- **Manufacturer Fault:** The victim must prove that the manufacturer failed to implement adequate safety measures, quality controls, or risk assessments. Yet the manufacturer may have complied with all existing standards and regulations; the harm may result from the **inherent unpredictability** of autonomous learning systems, which is not attributable to negligence (Abbott, 2020).
- **Operator Fault:** The victim must prove that the operator failed to supervise, maintain, or use the system properly. But if the system is designed to operate autonomously, without

continuous human oversight, imposing a duty of constant supervision would negate the very purpose of automation (Calo, 2015).

In practice, the victim confronts a *probatio diabolica*—a "devilish proof"—requiring expert reconstruction of the AI system's internal decision-making process, which is often technically infeasible, prohibitively expensive, or obstructed by corporate claims of trade secrecy and proprietary information (Pasquale, 2015).

Conclusion of Section 1: The fault element of traditional civil liability is structurally incompatible with autonomous AI systems. AI systems cannot commit fault due to lack of discernment, and proving human fault is rendered practically impossible by algorithmic opacity. Consequently, fault-based liability systematically fails to provide remedies for victims of AI-generated harms.

Section 2: The Fracture of Causation (*Lien de Causalité*)

The requirement of a causal nexus between the defendant's conduct (or the defect in a thing) and the plaintiff's harm is a universal and indispensable element of civil liability in all legal systems. Without causation, liability cannot be imposed, regardless of the severity of the harm or the culpability of the defendant.

2.1. Classical Causation Theories

Civil law systems, including Algeria and France, employ various theories of causation:

- **Equivalence Theory (*Théorie de l'équivalence des conditions*):** All conditions that are *sine qua non* (necessary) for the harm are considered causes.
- **Adequate Causation Theory (*Théorie de la causalité adéquate*):** Only those conditions that are foreseeable, typical, and adequate to produce the harm are considered legal causes.
- **Direct Causation Theory (*Théorie de la causalité directe*):** Only the immediate, direct cause of the harm is recognized.

Regardless of the theory adopted, the victim bears the burden of **proving** the causal link, typically through direct evidence, expert testimony, or judicial inference.

2.2. The Problem of Autonomous Learning and Self-Modification

Autonomous AI systems—particularly those employing reinforcement learning or online learning—continuously modify their internal parameters and decision-making processes based on new data and environmental feedback. This dynamic evolution creates a temporal and causal gap between the initial programming and the ultimate harmful output:

1. **Initial Programming (T_0):** The developer writes the code, selects the training data, and defines the learning algorithm.
2. **Deployment and Autonomous Learning (T_1 to T_n):** The AI system operates in the real world, encounters novel scenarios, retrains itself, adjusts its parameters, and modifies its behavior.
3. **Harmful Output (T_{n+1}):** The AI generates a harmful output—a misdiagnosis, an erroneous credit decision, a collision.

Causal Question: What *caused* the harmful output?

- Was it the initial programming at T_0 ? Possibly, but the system has since evolved far beyond its initial state.
- Was it the training data? Possibly, but the data may have been lawfully obtained, and the developer may not have foreseen the specific harmful interaction.
- Was it the autonomous learning process? Possibly, but this process is emergent, non-linear, and not directly controlled by any human actor.
- Was it the specific environmental input at T_{n+1} ? Possibly, but this input may have been unforeseeable and beyond the control of any party.

The result is causal indeterminacy: the harmful output is the product of a complex, multi-stage, distributed, and opaque process involving contributions from multiple actors and stochastic learning dynamics. Identifying a single, direct, legally cognizable cause becomes epistemologically and evidentially impossible (Scherer, 2016).

2.3. Jurisprudential Challenges

French jurisprudence has, in some contexts, relaxed the plaintiff's burden of proving causation by recognizing presumptions of causation (*présomptions de causalité*) when a thing "intervenes" in causing harm (e.g., *Cour de Cassation, Jand'heur*, 1930). However, this mechanism presupposes:

1. A relatively transparent causal mechanism (e.g., a car collides with a pedestrian; the collision caused the injury).
2. A custodian exercising meaningful control over the thing.

In the AI context, both conditions are absent: the causal mechanism is opaque, and the "custodian" does not exercise meaningful control over the autonomous system's evolving behavior.

Conclusion of Section 2: The requirement of proving a causal link is systematically obstructed by the autonomous learning, self-modification, and algorithmic opacity of AI systems. This fracture of causation compounds the failure of the fault element, rendering classical liability frameworks doubly inadequate for addressing AI-generated harms.

Section 3: The Transition to Strict Product Liability and Objective Liability

Faced with the failures of fault-based liability and the custody framework, legal systems have explored strict liability regimes, particularly product liability and objective liability (*responsabilité objective*).

3.1. Product Liability in Algerian and French Law

Algerian Law: Article 140 bis of the Algerian Civil Code, introduced to transpose principles analogous to the EU Product Liability Directive, provides:

"The producer is liable for damage caused by a defect in his product, whether he knew of the defect or not, unless he proves that the defect did not exist when the product was put into circulation, or that the defect resulted from compliance with mandatory regulatory standards."

French Law: Articles 1245 et seq. of the French Civil Code transpose Directive 85/374/EEC, establishing that the producer is liable for defective products without the need to prove fault, provided the victim establishes:

1. **The defect** (*le défaut*): a failure to provide the safety a person is entitled to expect.
2. **The damage** (*le dommage*).
3. **The causal link** (*le lien de causalité*).

3.2. Advantages of Product Liability for AI Systems

Applying product liability to AI offers significant advantages:

- **Eliminates the need to prove fault:** The victim need not prove negligence or wrongdoing by the developer.
- **Identifies a solvent defendant:** Manufacturers and developers are typically corporate entities with assets and insurance.
- **Incentivizes safety:** Liability incentivizes developers to invest in testing, validation, and risk mitigation.

3.3. Fundamental Limitations of Product Liability for AI

Despite these advantages, product liability encounters serious limitations when applied to AI:

a) Defining "Defect" in AI Systems

A "defect" traditionally means a deviation from design specifications, a manufacturing flaw, or inadequate warnings. However:

- **Design Defect:** If an AI system functions exactly as designed but produces harmful outputs due to emergent learning, is there a "defect"? The system may be operating within specification.
- **Manufacturing Defect:** Software is replicated perfectly; there is no "manufacturing" variation analogous to physical products.
- **Failure to Warn:** Warnings presuppose knowledge of risks. If the risks are emergent and unknowable at deployment, how can the developer warn against them?

b) The Development Risk Defense (*Risque de Développement*)

Article 1245-10 of the French Civil Code (and analogous provisions in Algerian law and the EU Directive) provides a defense if the producer proves:

"That the state of scientific and technical knowledge at the time the product was put into circulation was not such as to enable the existence of the defect to be discovered."

This development risk defense is a nearly universal shield for AI developers: given the inherent unpredictability and emergent properties of machine learning, developers can plausibly argue that the specific defect (the harmful output) was unknowable at the time of deployment (European Commission, 2020).

c) The Burden of Proving Causation Remains

Even under strict product liability, the victim must still prove the causal link between the defect and the damage. As discussed in Section 2, this remains extraordinarily difficult due to algorithmic opacity.

Conclusion of Section 3: While strict product liability represents a significant improvement over fault-based liability, it remains insufficient for AI systems due to difficulties in defining defects, the availability of the development risk defense, and persistent challenges in proving causation. A more robust, AI-specific strict liability regime is necessary.

Section 4: The Mandatory Insurance Paradigm and the Decoupling of Compensation from Fault

Given the evidentiary and conceptual failures of traditional liability regimes, scholars and policymakers have increasingly advocated for mandatory insurance schemes as the optimal mechanism for ensuring victim compensation, distributing risk, and fostering responsible AI development.

4.1. Theoretical Foundations: From Corrective to Distributive Justice

- **Corrective Justice:** Fault-based liability is grounded in the principle that wrongdoers should compensate their victims. This is appropriate when harms result from culpable human conduct.
- **Distributive Justice:** Strict liability and mandatory insurance are grounded in the principle that the costs of technologically generated risks should be distributed broadly across society, rather than concentrated on individual victims. This reflects the recognition that society as a whole benefits from technological innovation, and thus should bear its costs (Calabresi, 1970).

4.2. Models of Mandatory Insurance

Several models exist, drawn from analogous domains:

a) Motor Vehicle Insurance

Most jurisdictions require motor vehicle owners to carry compulsory third-party liability insurance. This ensures that victims of traffic accidents can obtain compensation regardless of the driver's solvency or fault. The system operates efficiently, compensates millions of victims annually, and is financed through risk-based premiums (Abraham & Rabin, 2019).

b) Workers' Compensation

Workers injured on the job receive compensation from mandatory employer-funded insurance schemes, without needing to prove employer fault. This "no-fault" system ensures swift, predictable compensation and avoids costly litigation (Fishback & Kantor, 2000).

c) Nuclear Accident Liability

Due to the catastrophic potential of nuclear accidents and the difficulty of proving causation, international conventions (e.g., the Paris Convention on Third Party Liability in the Field of Nuclear Energy, 1960) impose **strict liability** on nuclear operators and require them to carry mandatory insurance up to specified limits (OECD, 2018).

4.3. Applying Mandatory Insurance to AI Systems

A mandatory insurance regime for AI could operate as follows:

1. **Obligation:** Developers, manufacturers, or deployers of high-risk AI systems (as classified under the EU AI Act or analogous national frameworks) are required to obtain liability insurance covering harms caused by their systems.
2. **Risk-Based Premiums:** Insurance premiums are calibrated based on the risk profile of the AI system, incentivizing developers to invest in safety and risk mitigation.
3. **Compensation Fund:** In cases where the liable party cannot be identified or is insolvent, a public or industry-funded compensation fund provides compensation to victims, analogous to motor vehicle accident compensation funds.

4. **No-Fault Mechanism:** Victims need only prove the damage and the causal link to the AI system (facilitated by evidentiary presumptions), without proving fault or defect.

4.4. Advantages of Mandatory Insurance

- **Ensures Victim Compensation:** By decoupling compensation from the often-futile search for fault, insurance ensures swift and predictable redress.
- **Distributes Risk:** Costs are spread across all users and beneficiaries of AI technology, rather than concentrated on individual victims.
- **Incentivizes Safety:** Risk-based premiums create financial incentives for developers to invest in robust safety measures, testing, and risk mitigation.
- **Reduces Litigation Costs:** No-fault mechanisms reduce the need for costly, protracted litigation over fault and causation.
- **Promotes Innovation:** By providing legal certainty and predictable liability costs, insurance facilitates investment and innovation in AI technologies (Mohamed, 2025, p. 772).

4.5. Challenges and Counterarguments

- **Moral Hazard:** If developers know they are insured, they may take excessive risks. This can be mitigated through risk-based premiums, policy exclusions for gross negligence, and regulatory oversight.
- **Actuarial Uncertainty:** Insurers require reliable data to calculate premiums. Given the novelty of AI risks, initial premium-setting may be difficult. However, experience from motor vehicle insurance demonstrates that actuarial knowledge develops over time.
- **Defining "High-Risk" AI:** Determining which AI systems require mandatory insurance necessitates clear regulatory classification, as provided by the EU AI Act's risk-based taxonomy.

Conclusion of Section 4: Mandatory insurance, coupled with strict or objective liability, represents the most promising pathway for ensuring comprehensive victim protection, efficient risk distribution, and responsible AI development. It shifts the focus from penalizing fault to guaranteeing compensation and managing systemic risk.

Final Considerations and Conclusion

The rapid proliferation and integration of artificial intelligence systems into critical domains of public and private life—from healthcare and transportation to finance, administration, and beyond—has precipitated a profound and urgent legal crisis. The foundational premises of classical civil liability law, meticulously constructed over centuries to regulate human conduct, have encountered a new and fundamentally incompatible reality: autonomous, self-learning, opaque algorithmic agents that operate beyond the predictable bounds of human control, intent, and foreseeability.

This article has systematically demonstrated that the traditional doctrinal architecture of tort liability—grounded in fault (*faute*), causation (*lien de causalité*), legal capacity (*capacité juridique*), and custody (*garde*)—suffers from structural, conceptual, and evidentiary failures when applied to AI-generated harms. Specifically:

1. **The fault element is inapplicable:** AI systems lack discernment, intentionality, and legal personality, rendering them incapable of committing fault. Simultaneously, the algorithmic opacity of the "black box" renders proof of human fault by programmers,

manufacturers, or operators a *probatio diabolica*, systematically obstructing victim claims.

2. **The causation requirement is fractured:** Autonomous learning, self-modification, and emergent properties create a temporal and causal gap between initial programming and harmful outputs, rendering the establishment of a direct causal link epistemologically and evidentially insurmountable.
3. **The custody framework is obsolete:** The classical notion of custody presupposes continuous, effective control over a thing's behavior—a condition systematically violated by autonomous AI systems that learn, adapt, and decide independently post-deployment.
4. **Product liability is insufficient:** While strict product liability eliminates the need to prove fault, it remains hampered by difficulties in defining defects in adaptive learning systems, the availability of the development risk defense, and persistent challenges in proving causation.

The cumulative effect of these failures is the creation of a legal vacuum: a zone in which serious, unjust harms occur, but victims are systematically denied access to compensation and redress due to insurmountable evidentiary barriers and conceptual mismatches between legal doctrine and technological reality. This outcome is normatively intolerable, for it violates the foundational principles of justice, accountability, and the rule of law.

To remedy this crisis, this article has identified and analyzed a necessary paradigm shift from fault-based liability to risk-based, objective, and strict liability regimes, complemented by mandatory insurance mechanisms. Drawing upon comparative analysis of Algerian, French, and European Union legal frameworks, as well as insights from international tort law scholarship, the article has demonstrated that:

- **Objective liability** (holding developers, manufacturers, or deployers liable for harm caused by their AI systems regardless of fault) aligns liability with those best positioned to prevent harm, internalize risk, and innovate responsibly.
- **Mandatory insurance** (requiring developers and deployers of high-risk AI systems to carry liability insurance or contribute to compensation funds) ensures swift, predictable victim compensation, distributes risk broadly across society, and reduces litigation costs.
- **Evidentiary presumptions** (rebuttable presumptions of causation and fault, as proposed in the EU AI Liability Directive) mitigate the black box problem by shifting the burden of proof to developers, who possess superior knowledge and resources.

These mechanisms collectively prioritize victim protection and systemic risk management over the penalization of individual fault—a shift that is both normatively justified and practically necessary in the context of opaque, autonomous, and socially beneficial technologies.

Legislative Imperatives for Algeria

For Algeria, the analysis reveals an urgent need for comprehensive legislative reform. Specifically:

1. **Enact a dedicated AI Liability Law:** Algeria should promulgate a *sui generis* statute governing liability for AI-generated harms, distinct from the general provisions of the Civil Code. This law should:
 - Establish strict liability for developers, manufacturers, and deployers of high-risk AI systems.

- Define high-risk AI (drawing upon the EU AI Act's risk-based classification).
 - Introduce rebuttable presumptions of causation and fault.
 - Require mandatory liability insurance or contributions to a compensation fund.
 - Prohibit or strictly regulate AI applications posing unacceptable risks (e.g., social scoring, mass surveillance).
2. **Amend the Insurance Law:** Introduce provisions mandating insurance coverage for AI-related liability, with premiums calibrated to risk profiles and regulatory compliance.
 3. **Establish Regulatory Oversight:** Create or designate a national AI regulatory authority (analogous to the EU AI Office) responsible for classification, conformity assessment, incident reporting, and enforcement.
 4. **Judicial Capacity Building:** Provide specialized training for judges, lawyers, and court-appointed experts in AI technology, algorithmic auditing, and digital evidence, enabling the judiciary to adjudicate AI liability disputes effectively.
 5. **Promote Multistakeholder Dialogue:** Engage developers, insurers, civil society, victims' advocates, and academic experts in the co-design of liability and insurance frameworks to ensure they are technically feasible, economically viable, and normatively just.

Broader Implications

The challenges posed by AI liability are not unique to Algeria; they are global and systemic. The solutions developed in Algeria can contribute to the broader international discourse on AI governance, offering models and lessons applicable across the MENA region and beyond. Conversely, Algeria can learn from and adapt best practices emerging from the EU, France, and other jurisdictions at the forefront of AI regulation.

Ultimately, the goal is not to stifle innovation but to ensure that technological progress serves humanity justly and equitably. AI has immense potential to improve lives, but this potential can be realized only within a legal framework that guarantees accountability, protects victims, and aligns incentives toward safety and responsibility.

Final Reflection

The advent of autonomous AI systems challenges us to rethink fundamental legal categories—personhood, agency, causation, responsibility—that have structured our normative order for millennia. This rethinking is not a capitulation to technological determinism but an affirmation of law's adaptability and enduring purpose: to secure justice, protect the vulnerable, and hold power accountable, regardless of whether that power is exercised by human hands or algorithmic code.

The path forward is clear: legislative action, institutional innovation, and unwavering commitment to the primacy of human dignity and justice. Algeria, like all nations navigating the AI revolution, stands at a crossroads. The choices made today will determine whether AI becomes a tool of empowerment and equity—or a source of unchecked harm and systemic injustice.

References

- Abbott, R. (2020). *The reasonable robot: Artificial intelligence and the law*. Cambridge University Press. <https://doi.org/10.1017/9781108631761>.
- Abraham, K. S., & Rabin, R. L. (2019). Automated vehicles and manufacturer responsibility for accidents: A new legal regime for a new era. *Virginia Law Review*, 105(1), 127–171.
- Algerian Civil Code. (1975). *Order No. 75-58 of September 26, 1975, containing the Civil Code, as amended*. Official Gazette No. 78, September 30, 1975. Algeria.
- Alioui, M. Q. (2023). Artificial intelligence: Its development, applications, and challenges. *Lubab for Strategic Studies*, 5(22), 13–35.
- Bertolini, A. (2020). Artificial intelligence and civil liability. *European Parliament Policy Department for Citizens' Rights and Constitutional Affairs*. [https://www.europarl.europa.eu/RegData/etudes/STUD/2020/621926/IPOL_STU\(2020\)621926_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/STUD/2020/621926/IPOL_STU(2020)621926_EN.pdf)
- Burrell, J. (2016). How the machine 'thinks': Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 1–12. <https://doi.org/10.1177/2053951715622512>
- Calabresi, G. (1970). *The costs of accidents: A legal and economic analysis*. Yale University Press.
- Calo, R. (2015). Robotics and the lessons of cyberlaw. *California Law Review*, 103(3), 513–563.
- Debbache, A. A. M. (2019). The elements of civil liability. *Journal of Legal and Social Sciences*, 4(2), Serial No. 14, 26–27.
- European Commission. (2020). *Report on the safety and liability implications of Artificial Intelligence, the Internet of Things and robotics* (COM(2020) 64 final). https://ec.europa.eu/info/sites/default/files/report-on-safety-and-liability-implications-of-ai-iot-robotics-2020_en.pdf
- European Commission. (2022). *Proposal for a Directive on adapting non-contractual civil liability rules to artificial intelligence (AI Liability Directive)* (COM/2022/496 final). <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52022PC0496>
- European Parliament. (2017). *Resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics* (2015/2103(INL)). https://www.europarl.europa.eu/doceo/document/TA-8-2017-0051_EN.html
- European Parliament and Council of the European Union. (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. *Official Journal of the European Union*, L 1689. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689>
- Fishback, P. V., & Kantor, S. E. (2000). *A prelude to the welfare state: The origins of workers' compensation*. University of Chicago Press. <https://doi.org/10.7208/chicago/9780226251462.001.0001>
- Honoré, T. (2010). *Responsibility and fault*. Hart Publishing.
- Hubbard, F. P. (2014). "Sophisticated robots": Balancing liability, regulation, and innovation. *Florida Law Review*, 66(5), 1803–1872.
- Khaloui, N. (2025). Civil liability in the era of artificial intelligence: An analytical study on the features of a liability of a special nature. *Journal of Legal and Economic Studies*, 8(2), 480–510.



- Lohsse, S., Schulze, R., & Staudenmayer, D. (Eds.). (2019). *Liability for artificial intelligence and the internet of things*. Nomos Verlagsgesellschaft mbH & Co. KG. <https://doi.org/10.5771/9783748901570>
- Malouhi, A. A. J. (2009). *Liability for things and its applications to legal persons in particular* (1st ed.). Dar Al-Thaqafa for Publishing and Distribution.
- Mohamed, M. T. I. (2025). Civil liability insurance for the risks of artificial intelligence. *Damietta Journal of Legal and Economic Studies*, 12(12), Part II, 743–787.
- Muawad, N. (n.d.). *Manufacturer liability for aircraft* (2nd ed.). Dar Al-Nahda Al-Arabiya.
- OECD Nuclear Energy Agency. (2018). *The Paris Convention on Third Party Liability in the Field of Nuclear Energy*. https://www.oecd-neo.org/jcms/pl_37264/paris-convention-on-third-party-liability-in-the-field-of-nuclear-energy
- Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press.
- Al-Qousi, H. (2018). The problem of the person responsible for operating the robot: The impact of the human proxy theory on the feasibility of law in the future. *Jil Journal of Legal Research*, (25), Center for Jil Scientific Research, Lebanon Branch, 55–88.
- Russell, S. J., & Norvig, P. (2021). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
- Salama, S., & Abu Qura, K. (n.d.). *Challenges and ethics of the age of robots* (1st ed.). Emirates Center for Strategic Studies and Research.
- Scherer, M. U. (2016). Regulating artificial intelligence systems: Risks, challenges, competencies, and strategies. *Harvard Journal of Law & Technology*, 29(2), 353–400. <https://doi.org/10.2139/ssrn.2609777>
- Selbst, A. D., & Barocas, S. (2018). The intuitive appeal of explainable machines. *Fordham Law Review*, 87(3), 1085–1139.
- Solaiman, S. M. (2017). Legal personality of robots, corporations, and idols, and affirmative obligations for AI injuries. *Artificial Intelligence and Law (Springer)*, vol. 25(2), pp. 155–179.
- Teubner, G. (2006). Rights of non-humans? Electronic agents and animals as new actors in politics and law. *Journal of Law and Society*, 33(4), 497–521. <https://doi.org/10.1111/j.1467-6478.2006.00368.x>
- Viney, G., & Jourdain, P. (2017). *Les conditions de la responsabilité* (4th ed.). L.G.D.J.
- Wagner, G. (2020). Robot liability. In S. Lohsse, R. Schulze, & D. Staudenmayer (Eds.), *Liability for artificial intelligence and the internet of things* (pp. 27–62). Nomos Verlagsgesellschaft mbH & Co. KG. <https://doi.org/10.5771/9783748901570-27>