

Two Years after the Adoption of the EU AI Act (Regulation (EU) 2024/1689, 13 June 2024): Impacts on Legal Practice through Selected Italian Judgments—Stricter Diligence Obligations and the Concept of Pedagogical Damages

Christian Giuseppe Comito¹, Federico Badessi²

¹Department of Legal Sciences, Sapienza University of Rome, Rome, Italy

²Presidency of the Council of Ministers—DAGL, Rome, Italy

Email: christian.comito@uniroma1.it, federicobadessi1@gmail.com

How to cite this paper: Comito, C. G., & Badessi, F. (2026). Two Years after the Adoption of the EU AI Act (Regulation (EU) 2024/1689, 13 June 2024): Impacts on Legal Practice through Selected Italian Judgments—Stricter Diligence Obligations and the Concept of Pedagogical Damages. *Beijing Law Review*, 17, 611-623. <https://doi.org/10.4236/blr.2026.172032>

Received: May 12, 2026

Accepted: June 19, 2026

Published: June 22, 2026

Copyright © 2026 by author(s) and Scientific Research Publishing Inc. This work is licensed under the Creative Commons Attribution International License (CC BY 4.0). <https://creativecommons.org/licenses/by/4.0/>



Open Access

Abstract

Two years after the adoption of Regulation (EU) 2024/1689 of 13 June 2024 (“AI Act”) and shortly after the entry into force of Italian Law no. 132 of 23 September 2025, the first body of Italian case law on the so-called “algorithmic hallucinations”—i.e. fictitious judicial precedents generated by Large Language Models and incorporated, without verification, into pleadings—makes it possible to assess the impact of the new normative framework on the legal profession. This paper systematically reconstructs six rulings issued between March 2025 and March 2026 (Florence, Latina, the Lombardy Administrative Tribunal, Ferrara, Syracuse and Mantova) and reads them against three interpretative coordinates: i) Article 14 of the AI Act, which, *mutatis mutandis*, lays down the principle of effective human oversight over high-risk AI systems and the duty to counter “automation bias”; ii) Article 13 of Italian Law no. 132/2025, which confines the use of AI by intellectual professionals to merely instrumental purposes and preserves the prevalence of human intellectual work; iii) the theoretical framework of punitive—or, more properly, “*pedagogical*”—damages, whose plurifunctional rationale (afflictive, educational, deterrent, compensatory and symbolic) sheds new light on the application of Article 96, paragraphs 3 and 4, of the Italian Code of Civil Procedure as consistently activated by the surveyed case law. The conclusion is that the Italian judiciary, anticipating the gradual application of the AI Act, has already developed a coherent

doctrine of professional diligence which is best understood through the lens of pedagogical, rather than merely punitive, sanctioning.

Keywords

Artificial Intelligence Act, Algorithmic Hallucinations, Human Oversight, Punitive Damages, Article 96 Code of Civil Procedure, Legal Profession

1. Introduction

Two years after the adoption of Regulation (EU) 2024/1689 of 13 June 2024 (hereinafter, “AI Act”) and a few months after the entry into force, on 10 October 2025, of Italian Law no. 132 of 23 September 2025 (“Provisions and Delegations to the Government in Matters of Artificial Intelligence”), the time appears ripe for an initial assessment of the impact of the new normative framework on legal practice. The present contribution does so through a particular lens: that of a cluster of Italian judgments—six between March 2025 and March 2026—confronting the phenomenon of “algorithmic hallucinations”, i.e. fictitious case-law citations generated by Large Language Models and incorporated, without verification, into pleadings.

The phenomenon is the convergent result of three concurrent factors: the widespread, often uncritical, diffusion of generative AI tools such as ChatGPT in law firms (Schwarcz & Choi, 2024); the statistical-probabilistic nature of such systems, which—as the Tribunal of Syracuse has authoritatively recalled—do not constitute legal databases but rather engines for the automatic generation of plausible language (Dahl et al., 2024; Magesh et al., 2025); and, finally, the fragility of a system of professional and procedural filters still largely calibrated on the figure of a lawyer who personally consults primary sources (Curlin, 2025). The result is what may be termed, borrowing a technical expression now widely received by the courts, an “algorithmic hallucination”: the production, by the language model, of formally plausible yet substantively false content, destined to transit into the pleadings through the signature of the professional.

The analysis is organised along three coordinates. First, the AI Act—and in particular Article 14 on “human oversight”—offers the supra-national normative anchor of the principle of centrality of human decision that already underlies the surveyed Italian case law. Second, Article 13 of Law no. 132/2025 translates that principle, at the national level, into a positive duty incumbent on intellectual professionals: the use of AI must remain instrumental and ancillary, never substitutive of human intellectual work. Third, the punitive-damages framework—or, more properly, the framework of “pedagogical damages”—illuminates the way in which Article 96, paragraphs 3 and 4, of the Italian Code of Civil Procedure has been applied by the surveyed courts: not as a merely punitive device, but as a plurifunctional instrument oriented to education, deterrence and the symbolic reaf-

firmation of professional standards.

A preliminary terminological clarification is in order. For the purposes of the present contribution, “*pedagogical damages*” denote a category of judicially-imposed pecuniary sanctions which, while sharing with the classical idea of punitive damages the extra-compensatory dimension, are distinguished by two other criteria. First, by functional plurality: the sanction simultaneously performs afflictive, educational, deterrent and symbolic purposes, rather than the predominantly punitive or retributive aim associated with punitive damages *stricto sensu*. Second, by constructive orientation: the sanction is calibrated not to inflict suffering but to guide the conduct of the addressee and, more broadly, of the professional community to which the addressee belongs (Benatti, 2017; Owen, 1989; Owen, 1994; Scarchillo, 2018).

2. The Normative Framework: Article 14 AI Act and Article 13 of Law No. 132/2025

2.1. Article 14 AI Act: Human Oversight and Automation Bias

Article 14 of the AI Act articulates the principle of “human oversight” as a structural requirement of high-risk AI systems. Paragraph 1 requires such systems to be designed and developed in such a way “that they can be effectively overseen by natural persons during the period in which they are in use”. Paragraph 2 specifies that human oversight aims to prevent or minimise risks to health, safety or fundamental rights “in particular where such risks persist despite the application of other requirements set out in this Section”. Paragraph 4, more precisely, mandates that the system be provided to deployers in such a way that natural persons assigned to oversight may, *inter alia*, “properly understand the relevant capacities and limitations” of the system; “remain aware of the possible tendency of automatically relying or over-relying on the output” (the so-called “automation bias”); and “decide, in any particular situation, not to use the high-risk AI system or to otherwise disregard, override or reverse the output”.

Two analytical points deserve emphasis. First, although the formal scope of Article 14 is limited to “high-risk” systems within the meaning of Article 6 of the AI Act, the rationale it expresses—effective oversight, awareness of model limits, resistance to automation bias—has manifestly broader normative purchase: it crystallises a general principle of trustworthy human-machine interaction. Second, the explicit reference to “automation bias” in Article 14(4)(b) acquires special significance in the legal-professional context. The *bias* consists precisely in the human tendency to over-rely on the apparent fluency of the machine, accepting its outputs without the critical mediation that the structure of the activity would require. As scholarship has observed, the efficacy of liability rules in automated contexts depends structurally on the quality of integration between the human agent and the technical system: where such integration is weak, the rules conceived for an agent in full control of the process lose their grip (Bertolini, 2024; Fink, 2025; Green, 2022; Laux, 2024; Lupo, 2019; Sourdin, 2018).

A methodological caveat on the legal status here ascribed to Article 14 is in order. Article 14 is therefore invoked, throughout this contribution, as a broader interpretive principle: a hermeneutic key. The norm crystallises, at the European level, a general standard of “*trustworthy human-machine interaction*” which—although its operative reach is, in the proper sense, addressed to providers and deployers of high-risk systems—illuminates the rationale underpinning both Article 13 of Law no. 132/2025 and the duty of professional diligence reconstructed by the Italian courts in the cases discussed below.

2.2. Article 13 of Law No. 132/2025: Instrumental Use and Prevalence of Human Work

At the domestic level, Article 13 of Law no. 132/2025—rubricated “Provisions on Intellectual Professions”—provides that “the use of artificial intelligence systems in intellectual professions shall be limited to the exercise of ancillary and supporting activities, with the prevalence of the intellectual work which is the subject matter of the professional service”. Paragraph 2 adds that, in order to safeguard the fiduciary relationship between professional and client, information on the AI systems employed by the professional must be communicated “in clear, simple and exhaustive language”. As the explanatory report makes clear, the requirement of “prevalence” does not refer to quantity but to quality: critical human thought must always retain primacy over the use of AI tools.

Article 13 of Law no. 132/2025 thus accomplishes, on Italian soil, the same systematic operation that Article 14 of the AI Act performs at the European level: it translates into a positive duty the principle that the human professional may not abdicate to the machine. More precisely, it crystallises that principle into two distinct obligations: a substantive one, concerning the prevalence of human work; and an informational one, concerning the duty of disclosure to the client. As shall be argued in what follows, the Italian case law on algorithmic hallucinations has, in substance, anticipated the substantive obligation: well before the entry into force of the new law, Italian courts had already begun to construct, on the basis of Article 96 of the Code of Civil Procedure and general principles of professional diligence, a robust duty of verification of AI-generated outputs.

3. The Inaugural Case: Tribunal of Florence, 14 March 2025

The chronological reconstruction must start from the order of the Tribunal of Florence, Specialised Section for Enterprises, dated 14 March 2025—the first Italian judicial pronouncement to address directly the issue of judicial citations generated by AI and not verified. The case concerned interim relief in trade-mark and copyright law: the owner of a denominative-figurative trade mark sought to enjoin the sale of t-shirts reproducing satirical cartoons. What matters for present purposes is, however, ancillary: in the proceedings on appeal, counsel for one of the defendants had cited certain Supreme Court precedents in support of the absence of bad faith on the part of the retailers—precedents which, upon verification,

proved to be entirely non-existent.

Called to explain, counsel disclosed that the citations resulted from research carried out by an associate by means of “ChatGPT”, of whose use the lead counsel had been unaware. The AI had generated outputs corresponding to the phenomenon of “artificial intelligence hallucinations”: numbers nominally referring to Supreme Court rulings whose actual content bore no relation to the topic at hand. The Tribunal rejected the application for aggravated liability under Article 96 of the Code of Civil Procedure, observing that the defendant’s strategy had been grounded, since the first instance, on the absence of bad faith—a position already accepted in the earlier orders—and that the citation of Supreme Court precedents on appeal was aimed “at reinforcing an already-known defensive apparatus” and not “at resisting in bad faith”. The court nevertheless did not conceal “the disvalue concerning the omitted verification of the actual existence of the rulings resulting from the interrogation of AI”.

The Florentine decision thus represents the most cautious point of balance: it acknowledged the censurability of the conduct but excluded the procedural sanction, on the basis of an interpretation of Article 96 that required a specific functional connection between the act of citation and the strategic aim of misleading the court. Subsequent case law has decisively moved beyond this position.

4. The Turn towards Aggravated Liability: Latina and the Lombardy Administrative Tribunal

4.1. Tribunal of Latina, 23 September 2025, No. 1034

Six months after the Florentine decision, the Tribunal of Latina, in its judgment no. 1034 of 23 September 2025, marked an initial inflection. The case displayed traits of particular gravity: the judge observed that the application before the court—“as well as all the other hundreds of proceedings handled by the same counsel, all drafted in stencil fashion”—had manifestly been generated through AI tools. The diagnosis rested on multiple converging indicators: procedural anomalies, low quality of pleadings, total irrelevance of the arguments, a heap of statutory and case-law citations devoid of logical order. The judgment held that the action, conducted in such manner, was “introduced in bad faith or with gross negligence, such as to justify a sanction under Article 96, paragraph 3, of the Code of Civil Procedure”, and ordered the payment of one thousand euros to the opposing party and a further one thousand euros to the fines fund.

4.2. T.A.R. Lombardy, 21 October 2025, No. 3348: The Centrality of Human Decision

A few weeks later, the Administrative Tribunal of Lombardy, Section V, in its judgment no. 3348 of 21 October 2025, contributed decisively to the systematic articulation of the relevant principles. Although moving in administrative law (the case concerned the failure of a student under an individualised teaching plan to advance to the next class at a Milanese music high school), the decision is of gen-

eral significance. All the cited precedents—of the Council of State and of various administrative tribunals—proved upon inspection to be irrelevant or simply non-existent: judgments concerning urban-planning litigation, denials of amnesty, management of reception centres, public-sector indemnities, denials of clearance for recreational flight—in short, materials wholly unrelated to the subject matter.

At the oral hearing, counsel for the applicant declared on the record that the citations had been obtained through “research tools based on artificial intelligence which generated erroneous results”. The Collegium denied that the circumstance could have any exonerating value. First, it invoked the principle that the signature of pleadings serves to attribute responsibility for the outcome of the written defences to the signatory “regardless of whether the latter has drafted them personally or by relying on the activity of associates or AI tools”: an integral assumption of content that excludes any form of “technical delegation” to the machine. Second, and more innovatively, the court invoked “the principle of centrality of human decision”, referring expressly to the “Charter of Principles for a Conscious Use of Artificial Intelligence Systems in the Legal Profession” adopted by the Milan Bar Council in 2024, and described AI tools as a “possible source of erroneous results commonly qualified as artificial intelligence hallucinations, which occur when such systems invent non-existent results that nevertheless appear consistent with the topic addressed”. The court, in addition to ordering the costs of the proceedings, transmitted a copy of the judgment to the Milan Bar Council for the relevant disciplinary assessments.

It is at this point that the connection with Article 14 of the AI Act becomes manifest. The principle of “centrality of human decision” invoked by the administrative court is, in substance, the same principle that Article 14—with the language of “effective oversight” and resistance to “automation bias”—imposes at the European level. The Italian judiciary, well before the gradual application of the AI Act (which becomes substantially operational only on 2 August 2026), is already enforcing a functional analogue of that principle through the procedural device of Article 96 of the Code of Civil Procedure and the disciplinary apparatus of the Bar.

5. The Systematic Synthesis: Tribunal of Syracuse, 20 February 2026, No. 338

5.1. The Case and the Reasoning by Exclusion

The judgment of the Tribunal of Syracuse, Second Civil Section, no. 338 of 20 February 2026, constitutes—at the current state of the case law—the apex of the Italian elaboration on the topic. Its importance derives less from the substantive matter (a dispute concerning Article 38 of the Civil Code and the applicability of Article 1957 of the Civil Code to those who have acted on behalf of an unincorporated association) than from the rigour of its reasoning. The plaintiff, in opposing the time-bar defence, had quoted in memorial four Supreme Court precedents (Cass. no. 1216/2000, no. 8379/2006, no. 14795/2003, no. 4553/2004), reproducing

passages between inverted commas. The court verified, through consultation of the Supreme Court database, that none of the four contained the quoted passages: the judgments, in their authentic versions, dealt with wholly extraneous matters, and the quoted passages “do not correspond to any existing judicial text”.

The court then proceeded by methodical exclusion. It ruled out, first, the hypothesis of a malfunction of professional legal databases: such tools index authentic decisions and do not generate text, and cannot therefore produce precedents with non-existent numbers, arguments and quotations. It excluded, second, the hypothesis of a mere mnemonic or transcription error: the citations were not a misnumbered or mis-sectioned reference, but “maxims constructed *ex novo*, devoid of any correspondence with the recalled rulings”. It excluded, third, the hypothesis of deliberate fabrication: a professional knowingly inventing four non-existent precedents would expose herself to disproportionately severe disciplinary consequences in light of any conceivable defensive advantage. The only residual hypothesis, fully consistent with the concrete phenomenology of the case, was “that counsel made use of a generative artificial intelligence tool without subjecting its outputs to the requisite verification against primary sources”.

5.2. The Statistical-Inferential Nature of LLMs as Notorious Fact

The reasoning culminates in an argument that constitutes the conceptual core of the entire decision. It is now a notorious fact—the court holds—“acquired by the generality of citizens and certainly required of a professional legal operator”, that generative AI models do not constitute case-law databases from which to extract precedents, but rather automatic language-generation tools founded on statistical and probabilistic inferential mechanisms. Such systems, in other words, do not “know” or “remember” anything, but merely produce sequences of text statistically plausible on the basis of billions of training parameters, without access to any verified or verifiable knowledge base. For this reason they are notoriously subject to the phenomenon of hallucinations: the generation of formally plausible yet substantively false content, including never-rendered judicial citations.

The uncritical use of such tools—without verification of the reliability of outputs against primary sources (legal databases, official repertories, Supreme Court CED)—“integrates the elements of gross negligence”, as errors of such nature can no longer be tolerated given the current state of widely disseminated technical knowledge. Here lies the decisive break with the Florentine precedent: the conduct is censurable “in itself”, regardless of its specific strategic function in the architecture of the defence, because it significantly burdens the activity of the court and the opposing parties, who must verify each citation and respond to non-existent precedents. The Tribunal condemned the plaintiff under Article 96, paragraph 3, of the Code of Civil Procedure to pay a sum equal to the litigation costs (fourteen thousand one hundred and three euros) to the opposing party, and under paragraph 4 to pay two thousand euros to the fines fund—close to half of the statutory ceiling.

6. The Consolidation of the Holding: Ferrara and Mantova

The Syracuse approach found immediate echo in two subsequent decisions. The Tribunal of Mantova, Section I, in its judgment of 24 March 2026, concerned an action under Article 2901 of the Civil Code (revocatory action) and applied the same paradigm: the plaintiff had relied upon nine Supreme Court precedents which, upon verification, did exist but bore on entirely unrelated subject matters (administrative fines on EU subsidies, public employment, nursing labour, defamation damages, inheritance tax, contribution revaluation for asbestos exposure, tax claims, contracts). The court—expressly citing both T.A.R. Lombardy and the Tribunal of Syracuse—held that the plaintiff’s conduct integrated gross negligence and imposed, under Article 96, paragraphs 3 and 4, of the Code of Civil Procedure, two awards of two thousand two hundred and fifty euros each—one quarter of the awarded professional fees.

The Tribunal of Ferrara, in its order of 20 February 2026, encountered the same pattern in an entirely different context: an application for technical expert consultation under Article 696-bis of the Code of Civil Procedure regarding a fatal road accident. The procedural context is not without significance: the duty of verification of AI-generated outputs thus extends to the summary and instructional phases of civil litigation, not only to ordinary cognition proceedings. The subject matter—a fatal road accident—equally confirms that the phenomenon is no longer confined to complex commercial or administrative disputes but has reached the most routine areas of civil practice. The transversality of the phenomenon—from cautionary intellectual-property litigation to school appeals, from revocatory actions to road accidents—is now an established feature of the Italian legal landscape. The brevity of the Ferrara order is itself indicative: the court applied the Syracuse paradigm without elaboration, treating it as settled principle. Ferrara therefore operates as a consolidating link in the jurisprudential chain rather than an independent doctrinal source, and is analysed here accordingly.

7. Article 96 C.P.C. through the Lens of Punitive Damages: A Pedagogical Reading

7.1. The Plurifunctional Rationale of Pedagogical Damages

The way in which the surveyed courts have applied Article 96, paragraphs 3 and 4, of the Code of Civil Procedure invites a doctrinal reading through the lens of punitive damages—or, more properly, of what recent scholarship has termed “pedagogical damages” (Benatti, 2017; Owen, 1989; Owen, 1994; Salvi, 2019; Scarchillo, 2018; Sirena, 2018). The proposed re-denomination is itself revealing: drawing on the Greek etymology of *paidagogía*—“to guide”, “to accompany”—the formula seeks to restore the institutional dignity of an instrument whose function is not merely afflictive but plurifunctional: educational, deterrent, compensatory, and symbolic. As Zig Zigar incisively put it, “punishment is what you do to someone; discipline is what you do for someone”. Punishment generates resentment; disci-

pline aims at transformation.

Applied to the case law in question, this conceptual matrix proves illuminating. The sanctions imposed by the Tribunal of Syracuse (sixteen thousand one hundred and three euros in total) and by the Tribunal of Mantova (four thousand five hundred euros) cannot be adequately explained by reference to the compensatory function alone: the awards under paragraph 4 are paid not to the opposing party but to the State Treasury's fines fund, and bear no relation to any individual harm. They are best understood as performing, simultaneously, an afflictive function (towards counsel), an educational function (towards the legal profession as a whole), a deterrent function (towards future similar conduct), and a symbolic function (the public reaffirmation that algorithmic hallucinations are incompatible with professional standards).

7.2. Calibration of the Sanction: The “Kicker” and the 3:1 Proportion

The scholarly model of pedagogical damages provides, moreover, useful coordinates for the calibration of the sanction. The proposal of a “kicker”—a discretionary increment available to the judge to reflect the gravity of the conduct—together with the recommended cap of a 3:1 ratio between the pedagogical and the compensatory component (Calabresi, 1970; Maggiolo, 2024; Polinsky, 1972) finds a striking empirical anticipation in the Italian case law surveyed. In Mantova, the award under paragraph 3 corresponds to one quarter of the professional fees, and the award under paragraph 4 is set at the same level: in substance, a measured exercise of judicial discretion modelled on the seriousness of the conduct. In Syracuse, the gravity of the inferred conduct (four fabricated quotations) justifies a higher sanction in clear proportion to the exceptional disvalue of the conduct.

Equally significant is the destination of the sanction. The dissertation's model suggests that, beyond the compensatory portion, the surplus should serve social purposes, signalling the public dimension of the sanction. Article 96, paragraph 4, of the Italian Code of Civil Procedure already incorporates a functional equivalent of this proposal: the award could be paid to the fines fund of the State Treasury, that is to say, to a collective recipient. The pedagogical character of the sanction is thus structurally inscribed in the architecture of the rule. The fines do not only enrich the opposing party (whose actual damage is already addressed by paragraph 3), but rather serve the general interest in the integrity of judicial proceedings—an interest of constitutional rank, encompassing fair trial, the effectiveness of adversarial proceedings, and the centrality of the judge in the construction of judicial truth.

8. Systematic Implications and Prospects

Two years after the AI Act and a few months after the entry into force of Law no. 132/2025, the surveyed Italian case law allows three concluding considerations. First, the Italian judiciary has anticipated, through general principles of professional diligence and the procedural lever of Article 96 of the Code of Civil Proce-

cedure, the substantive duties that the new normative framework codifies. The principle of “centrality of human decision” invoked by T.A.R. Lombardy, the statistical-inferential analysis of LLMs developed by the Tribunal of Syracuse, and the qualification of uncritical AI use as gross negligence consolidated by Ferrara and Mantova, all functionally implement—ex ante and through different legal techniques—the obligations that Article 14 of the AI Act and Article 13 of Law no. 132/2025 impose on professionals.

Second, this jurisprudential anticipation is best understood through the lens of pedagogical damages. The sanctions awarded under Article 96, paragraphs 3 and 4, of the Code of Civil Procedure cannot adequately be explained by the compensatory paradigm alone: their afflictive, educational, deterrent and symbolic dimensions are equally constitutive. The destination of the awards under paragraph 4 to the fines fund of the State Treasury—a collective recipient—confirms that the rationale of the sanction transcends private compensation and addresses the public interest in the integrity of judicial proceedings. The doctrinal model proposed in the punitive-damages literature, with its kicker and 3:1 ratio, offers a coherent analytical framework for the proportional quantification of such sanctions.

Third, the gradual entry into operation of the AI Act and the recent entry into force of Law no. 132/2025 are likely to consolidate the existing case-law direction rather than to alter it. The duty of verification of AI-generated outputs, already crystallised in the surveyed Italian judgments, finds normative anchorage in both Article 14 of the AI Act (counter to automation bias) and Article 13 of Law no. 132/2025 (instrumental use only and prevalence of human work). The intersection of supra-national and domestic norms, together with the elaboration of soft-law deontological instruments such as the Milan Charter of Principles, configures a multi-layered system within which the case law on algorithmic hallucinations is destined to find a robust normative anchoring.

The legal profession is thus called to a task that is at once technical and cultural: to integrate generative AI tools into legal practice in modalities that preserve, rather than erode, the dimension proper to qualified protection of rights. The machine as instrument, not as substitute; verification as duty, not as accessory; responsibility as principle, not as exception. These three formulae—to judge by the speed with which they have asserted themselves in Italian case law—are set to become, in the coming years, the load-bearing axis of legal-professional deontology in matters of new technologies. The pedagogical, rather than merely punitive, character of the sanction is the symbolic key that holds the system together: a sanction that is afflictive, but also formative; deterrent, but also constructive; punitive in form, but pedagogical in function (Scarchillo, 2018).

It is, however, appropriate to offer two final caveats. First, the sample on which the present analysis rests is small—six rulings—and temporally compact (March 2025 - May 2026): it constitutes an early, formative phase in the Italian jurisprudential treatment of generative AI in legal practice, and the conclusions reached here must be understood as preliminary, susceptible of revision as the body of

relevant decisions grows and as authoritative interpretation of both the AI Act and Law no. 132/2025 consolidates. The doctrinal model of pedagogical damages developed in Section 7 should accordingly be read as a heuristic framework grounded in an emerging case-law orientation, rather than as a fully stabilised normative finding. Second, and more precisely, all six surveyed cases concern a specific and relatively narrow phenomenology: the citation of non-existent or wholly unrelated judicial precedents—so-called citation hallucinations—generated by Large Language Models deployed in the drafting of judicial pleadings. The surveyed jurisprudence does not yet address the broader spectrum of AI-assisted legal drafting: argumentative construction by AI, the drafting of factual narratives, AI-assisted contract drafting, automated summarisation of judicial materials, or the use of predictive analytics in litigation strategy.

The doctrinal coordinates reconstructed in this paper—the duty of verification, the principle of centrality of human decision, the qualification of uncritical AI use as gross negligence—are, in the present writers’ assessment, presumptively transposable to those further dimensions of professional activity. The empirical confirmation of such transposability, however, must await the development of a more extensive and differentiated body of case law.

Acknowledgements

The author Dr. Comito wishes to express his deepest gratitude to the late Professor Guido Alpa, whose teaching and scholarly guidance have shaped, and shall continue to shape from above, the path of this and every other inquiry into the law. To his memory this contribution is gratefully dedicated.

Author’s Contribution

Although the contribution is the result of shared and coordinated reflections, authorship of the work is attributable, as regards chapters 1, 4, 7, 8, to Dr. Comito, and, as regards chapters 2, 3, 5, 6, to Dr. Badessi.

Declaration

The opinions expressed in this work are attributable to the authors alone and do not in any way commit the administration to which they belong.

Conflicts of Interest

The authors declare no conflicts of interest regarding the publication of this paper.

References

- Benatti, F. (2017). Note sui Danni Punitivi in Italia: Problemi e Prospettive. *Contratto e Impresa*, 1129-1155.
- Bertolini, A. (2024). *Intelligenza Artificiale e Responsabilità Civile. Problema, Sistema, Funzioni*. Il Mulino.
- Calabresi, G. (1970). *The Costs of Accidents: A Legal and Economic Analysis*. Yale Uni-

versity Press.

- Curlin, J. H. (2025). ChatGPT Didn't Write This ... or Did It? The Emergence of Generative AI in the Legal Field and Lessons from Mata V. Avianca. *Arkansas Law Review*, 78, Article 6. <https://doi.org/10.54119/alr.mlrt4575>
- Dahl, M., Magesh, V., Suzgun, M., & Ho, D. E. (2024). Large Legal Fictions: Profiling Legal Hallucinations in Large Language Models. *Journal of Legal Analysis*, 16, 64-93. <https://doi.org/10.1093/jla/laae003>
- Fink, M. (2025). *Human Oversight under Article 14 of the EU AI Act*. AI Governance Library Working Paper.
- Green, B. (2022). The Flaws of Policies Requiring Human Oversight of Government Algorithms. *Computer Law & Security Review*, 45, Article ID: 105681. <https://doi.org/10.1016/j.clsr.2022.105681>
- Laux, J. (2024). Institutionalised Distrust and Human Oversight of Artificial Intelligence: Towards a Democratic Design of AI Governance under the European Union AI Act. *AI & Society*, 39, 2853-2866. <https://doi.org/10.1007/s00146-023-01777-z>
- Lupo, G. (2019). Regulating (Artificial) Intelligence in Justice: How Normative Frameworks Protect Citizens from the Risks Related to AI Use in the Judiciary. *European Quarterly of Political Attitudes and Mentalities*, 8, 75-96.
- Magesh, V., Surani, F., Dahl, M., Suzgun, M., Manning, C. D., & Ho, D. E. (2025). Hallucination-Free? Assessing the Reliability of Leading ai Legal Research Tools. *Journal of Empirical Legal Studies*, 22, 216-242. <https://doi.org/10.1111/jels.12413>
- Maggiolo, M. (2024). Introduzione al Risarcimento Proporzionale. *Danno e Responsabilità*, 30, 691-698.
- Owen, D. G. (1989). The Moral Foundations of Punitive Damages. *Alabama Law Review*, 40, 705-739.
- Owen, D. G. (1994). A Punitive Damages Overview: Functions, Problems and Reform. *Villanova Law Review*, 39, Article 3.
- Polinsky, A. M. (1972). Probabilistic Compensation Criteria. *The Quarterly Journal of Economics*, 86, 407-425. <https://doi.org/10.2307/1880801>
- Salvi, C. (2019). *La Responsabilità Civile*. Giuffrè Francis Lefebvre.
- Scarchillo, G. (2018). La Natura Polifunzionale della Responsabilità Civile. *Contratto e Impresa*, 34, 289-327.
- Schwarcz, D. B., & Choi, J. H. (2024). AI Tools for Lawyers: A Practical Guide. *Minnesota Law Review Headnotes*, 108, 1-39.
- Sirena, P. (2018). Il Riconoscimento dei Danni cd. Punitivi e la Restituzione dell'Arricchimento senza Causa. *Rivista di Diritto Civile*, 64, 531-537.
- Sourdin, T. (2018). Judge V Robot? Artificial Intelligence and Judicial Decision-Making. *University of New South Wales Law Journal*, 41, 1114-1133. <https://doi.org/10.53637/zgux2213>

Appendix

Trib. Firenze, sez. spec. impresa, ord. 14 March 2025, in DeJure.

Trib. Latina, 23 September 2025, no. 1034, in DeJure.

T.A.R. Lombardia, Milano, sez. V, 21 October 2025, no. 3348, in giustizia-amministrativa.it.

Trib. Ferrara, ord. 20 February 2026, in DeJure.

Trib. Siracusa, sez. II civ., 20 February 2026, no. 338, in DeJure.

Trib. Mantova, sez. I, 24 March 2026, in DeJure.