

Volume 1, Issue 1

Research Article

Date of Submission: 15 July, 2026

Date of Acceptance: 11 Aug, 2026

Date of Publication: 19 Aug, 2026

Hallucinated Justice; AI-Generated Case Law and the Attorney's Obligation to Verify

Andrew Ndubisi Ucheomumu* 

B.Sc. Economics and Political Science, US

***Corresponding Author:** Andrew Ndubisi Ucheomumu, B.Sc. Economics and Political Science, US.

Citation: Ucheomumu, A. N. (2026). Hallucinated Justice; AI-Generated Case Law and the Attorney's Obligation to Verify. *World J Law Govern Pub Policy*, 1(1), 01-15.

Abstract

The proliferation of generative artificial intelligence ("AI") in legal practice has produced a novel and troubling category of professional responsibility violations. Courts across the United States have confronted an increasing volume of pleadings and briefs containing legal authorities that either do not exist or that AI systems have materially mischaracterized. Unlike ordinary citation errors attributable to research lapses, AI-generated hallucinations can produce fabricated case names, fictitious reporters, invented quotations, and nonexistent holdings, all rendered with the superficial confidence of authentic legal authority.

Generative AI is an important tool in today's legal practice, and it is neither the first nor the last technological innovation to transform the practice of law. From the printing press to the typewriter, from Westlaw's computerized databases to electronic filing systems, each era's defining technology has enhanced attorney productivity, expanded access to legal services, and reshaped the mechanics of legal work. Generative AI is the latest and most powerful entry in that lineage — an indispensable tool that enables attorneys to research, draft, and analyze with a speed and breadth previously unattainable. Critically, the courts that have imposed sanctions for AI misuse have been uniform on one point: they do not condemn the use of AI in legal research and briefing. From *Mata v. Avianca* to *Mezu v. Mezu*, the judicial message has been consistent — AI is a legitimate and valuable instrument of legal practice. What the courts condemn, without exception, is the abdication of professional judgment that occurs when attorneys submit AI-generated work product without independent verification. The obligation is not to avoid AI, but to use it as every transformative tool before it has always demanded to be used: with skill, with diligence, and with the lawyer's own professional accountability firmly intact.

This Article surveys the rapidly developing body of federal and state case law addressing sanctions for the submission of AI-generated fabricated citations and for AI misrepresentation of authentic precedent. Beginning with the foundational decision in *Mata v. Avianca, Inc.*, the Article traces the emergence of a national sanctions framework across numerous federal circuits and district courts, as well as courts in California, Hawaii, Louisiana, Maryland, and Utah. It examines the distinct but related problem of AI systems that accurately cite real cases but misrepresent their holdings, a subtler error that may escape initial detection yet pose equal risks to the integrity of adjudication.

The Article then analyzes the ethical obligations that existing professional responsibility rules impose on attorneys who use AI tools, including the duties of competence, diligence, and candor toward the tribunal, as amplified by American Bar Association Formal Opinion 512. Against this backdrop, the Article examines Maryland's existing rules and the significance of *Mezu v. Mezu* (2025) as Maryland's first published appellate opinion addressing AI-citation misconduct. The Article further examines the emerging question of law firm governance and institutional responsibility, analyzing the extent to which law firms bear independent obligations to implement AI oversight protocols under existing supervisory ethics rules. The Article concludes with concrete recommendations for Maryland courts, bar authorities, and the legislature: that Maryland should adopt AI-certification requirements analogous to those already implemented in several federal districts, should enforce its existing procedural and ethical rules vigorously in cases involving AI-generated errors, should issue definitive bar guidance to assist practitioners in navigating the responsible use of generative AI tools in legal practice, and should enact statutory authority enabling appellate courts to impose monetary sanctions *sua sponte*, closing the enforcement gap that constrained the court's response in *Mezu v. Mezu*.

Keywords: Artificial Intelligence, Generative AI, AI Hallucination, Attorney Sanctions, Professional Responsibility, Legal Ethics, Maryland, Rule 11, Competence, Candor Toward the Tribunal

Introduction

Legal practice is being transformed by generative artificial intelligence. Attorneys across the country now employ AI-powered platforms to draft memoranda, identify analogous precedents, synthesize statutory schemes, and prepare pleadings. The efficiency gains are real and, in many contexts, substantial. Yet this transformation has exposed the profession to a category of error it has not previously encountered at scale: the confident presentation to a court of legal authority that does not exist.

The phenomenon is neither trivial nor obscure. Since 2023, courts across the country have sanctioned attorneys who submitted briefs citing cases that no legal database contains, featuring judges who never served on the courts identified, and quoting language that no court ever wrote. The tool responsible for these errors, generative AI, produces its outputs through probabilistic text generation, not through the retrieval of stored documents. When an AI system is asked to identify case law supporting a legal proposition, it may generate a response that has the form of an authentic citation while lacking any factual basis. Courts and legal scholars have adopted the term “hallucination” to describe this phenomenon, borrowing from the computer science literature on large language model failures.

The hallucination problem is compounded by a related but distinct failure mode: AI systems that correctly identify real cases but materially misrepresent what those cases hold. An attorney relying on such output may submit a brief that cites genuine authority for propositions that the cited cases actively reject, undercut, or are simply silent about. This category of error may be even more dangerous than outright fabrication, precisely because the case citation itself will survive even diligent verification of existence, concealing the underlying error from the submitting attorney, potentially from opposing counsel, and in some instances from the court itself.

This Article addresses both failure modes. It surveys the leading federal and state decisions that have addressed AI-related citation misconduct, analyzes the professional responsibility framework that governs attorney conduct in this domain, and offers specific recommendations for Maryland courts and the Maryland bar. Maryland’s first published appellate decision on the subject, *Mezu v. Mezu* (2025), establishes a foundation; this Article argues that Maryland should build upon that foundation by adopting affirmative certification requirements and issuing clear, practical ethics guidance to the bar.

Part I provides background on how generative AI produces hallucinations and why the legal research context is particularly susceptible to this failure mode. Part II surveys the federal sanctions cases, beginning with the landmark *Mata v. Avianca* decision and tracing the evolution of sanctions jurisprudence through 2026. Part III examines significant state court decisions. Part IV addresses the distinct problem of AI mischaracterization of authentic precedent. Part V analyzes the ethical obligations that existing professional responsibility rules impose on attorneys who use AI. Part VI examines Maryland’s current framework. Part VII presents concrete recommendations. The Article concludes with a call for a coherent, consistent approach that enables Maryland courts to preserve the integrity of adjudication while accommodating the legitimate place of AI tools in modern legal practice.

Background: Generative AI and the Hallucination Problem in Legal Practice

How Generative AI Works

Generative AI systems of the kind now widely used in legal practice are large language models (“LLMs”) trained on vast corpora of text. These systems learn statistical patterns in language, the probability that a given word or phrase follows what precedes it, and use those patterns to generate new text that appears fluent and coherent. The sophistication of modern LLMs is remarkable: they can produce prose indistinguishable from human writing, summarize complex documents, translate between languages, and engage in what appears to be nuanced reasoning. But LLMs do not retrieve information in the way that conventional legal research databases do. Westlaw and Lexis retrieve actual documents from structured databases; their outputs are grounded in real text. An LLM, by contrast, generates output by predicting what tokens (words, punctuation, formatting) should follow the tokens already produced. When asked to cite a case, an LLM does not search a database of real cases; it generates a sequence of text that looks like a case citation. If the model’s training data contains sufficient information about a real case, it may generate an accurate citation. But if the training data is insufficient, the model will generate a citation that has the structural form of a real citation, to wit, a case name, a reporter, a volume number, page numbers, a court, a year, while corresponding to no actual decision.

This is the hallucination: not a deliberate fabrication, but a structural failure of grounded retrieval. The model has no mechanism for distinguishing between its accurate reproductions of real information and its generation of plausible-sounding fiction. It produces both with equal fluency and apparent confidence.

Why Legal Research is Particularly Vulnerable

Several features of legal citations make them especially susceptible to convincing hallucination. First, citations have a highly regular structure: [Party 1] v. [Party 2], [Volume] [Reporter] [Page] ([Court] [Year]). An LLM trained on legal text will produce outputs that conform to this structure even when the underlying information is fabricated. Second,

legal citations are dense with numerical information such as volume numbers, page numbers, years, that LLMs have difficulty faithfully reproducing even for real cases. Third, the sheer volume of case law means that many real cases are underrepresented in training data, making it difficult for the model to distinguish between cases it actually “knows” and cases it is generating from pattern. Fourth, attorneys and judges routinely encounter unfamiliar citations, making fabrications less likely to be detected by simple recognition.

The hallucination problem is not limited to citations. AI systems may also generate fabricated quotations attributed to real judges, invent procedural histories for real cases, or describe holdings that are the opposite of what the cited case actually holds. In each instance, the error may survive cursory review precisely because the surrounding text looks authentic.

The Scale of the Problem

The problem has grown rapidly. Between 2023 and 2026, courts in numerous federal circuits and in multiple states have imposed sanctions on attorneys for submitting AI-generated hallucinated citations. Law reviews, bar journals, and legal technology publications have documented dozens of additional incidents in which courts raised concerns about suspected AI hallucinations without ultimately imposing formal sanctions. Bar associations have begun issuing guidance, and several federal district courts have adopted local rules specifically addressing AI use in filings. The pace of development makes a systematic survey both timely and necessary.

Federal Cases: Sanctions For AI-Generated Fictitious Citations

Mata v. Avianca, Inc.—the Foundational Decision

The foundational federal decision on AI-generated fictitious citations is *Mata v. Avianca, Inc.*, 678 F. Supp. 3d 443 (S.D.N.Y. 2023). Mata arose from a personal injury action in which the plaintiff alleged that he had been injured when a metal serving cart struck his knee during an international flight. When the defendant moved to dismiss on statute of limitations grounds, plaintiff’s counsel filed an opposition brief citing a number of federal decisions. Defendant’s counsel reported to the court that several of the cited cases could not be located in any legal database. The court ordered plaintiff’s counsel to provide copies of the cited cases; counsel submitted documents purporting to be court opinions, but those documents were subsequently revealed to have been generated by the AI tool ChatGPT.

Judge P. Kevin Castel conducted a thorough inquiry and concluded that six of the submitted cases were entirely fictitious. In a passage that has since been widely quoted, Judge Castel wrote: “[Respondents] abandoned their responsibilities when they submitted non-existent judicial opinions with fake quotes and citations created by the artificial intelligence tool ChatGPT, then continued to stand by the fake opinions after judicial orders called their existence into question.” *Id.* at 448.

The court imposed sanctions pursuant to Federal Rule of Civil Procedure 11, ordering the attorneys, *inter alia*, to pay \$5,000 and to send copies of the sanctions order to their client. Judge Castel emphasized that “[t]echnological advances are commonplace and there is nothing inherently improper about using a reliable artificial intelligence tool for assistance. But existing rules impose a gatekeeping role on attorneys to ensure the accuracy of their filings.” *Id.* The court observed that the attorneys had “abandoned their responsibilities” by failing to conduct independent verification. *Id.*

Mata established several principles that have been reiterated across subsequent AI hallucination sanctions jurisprudence. First, the submission of fabricated or non-existent authorities violates Rule 11’s certification requirements. Second, the use of generative artificial intelligence does not excuse noncompliance with those obligations. Third, monetary sanctions are an appropriate remedy for such violations. Fourth, the duty to verify legal authorities rests with the signing attorney and cannot be satisfied through unverified reliance on third-party sources, including AI tools. The decision attracted national attention and has been widely cited as an early catalyst for judicial and bar association responses to AI-generated hallucinated citations.

Park v. Kim

In *Park v. Kim*, 91 F.4th 610 (2d Cir. 2024), the Second Circuit addressed a case in which an attorney submitted a brief to the court of appeals that cited a fabricated case, apparently generated by an AI tool. The case arose from a wage dispute in which the plaintiff had been sanctioned by the district court for serial discovery failures; on appeal, counsel cited a nonexistent case in support of a legal proposition.

The Second Circuit declined to impose monetary sanctions but directed that the matter be referred to the court’s Grievance Panel and that the panel’s determination be communicated to the attorney’s client. The court articulated a clear statement of the Rule 11 obligation in the AI context: “At the very least, the duties imposed by Rule 11 require that attorneys read, and thereby confirm the existence and validity of, the legal authorities on which they rely.” *Id.* at 615.

The court further stated that the submission of a false statement of law to the court, regardless of its source, “falls well below the basic obligations of counsel.” The Court did not impose any monetary sanction but referred the attorney to the Court’s Grievance Panel and ordered the attorney to furnish a copy of the court order to her client. *Id.*

Park v. Kim is significant for several reasons. First, it established that sanctions need not be monetary to be meaningful: a grievance referral carries substantial professional consequences. Second, the court's statement that attorneys must personally read the cases they cite, not merely verify their existence in a legal database, sets a demanding standard that subsequent courts have endorsed. Third, the decision demonstrates that AI hallucination problems have reached the appellate level, not merely trial courts.

United States v. Hayes

United States v. Hayes, 763 F. Supp. 3d 1054 (E.D. Cal. 2025), involved a federal criminal proceeding in which defense counsel submitted a motion that twice cited and quoted from a case that could not be found in any legal database. When the government's opposition pointed out that the cited case appeared not to exist, counsel initially claimed that the quoted passage had been drawn from a different case, a claim that proved equally unfounded. At a hearing, the court found counsel's explanations not credible and determined that the citation bore all the markings of a hallucinated case created by generative artificial intelligence tools.

The court imposed a \$1,500 monetary sanction and directed the Clerk to serve copies of the order on the State Bar of California, the District of Columbia Bar, and all district and magistrate judges in the district, an unusual step reflecting the court's concern about systemic deterrence. The court stated that "[c]iting nonexistent case law or misrepresenting the holdings of a case is making a false statement to a court." and this conclusion applies irrespective of whether an AI system conveyed the erroneous information. A motion for reconsideration was denied.

Hayes is noteworthy for several reasons. The criminal context underscores that AI hallucination problems are not limited to civil litigation. The court's direction that the matter be referred to disciplinary authorities signals an understanding that sanctions alone may be insufficient deterrence, and that professional discipline may be a necessary complement to court-imposed penalties. The court further made clear that reliance on generative artificial intelligence does not excuse the submission of inaccurate or nonexistent authority—even where the attorney may not have fully understood the nature of the tool's output. This reflects a rigorous application of the duty to verify legal authority that subsequent courts have increasingly adopted.

Gauthier V. Goodyear Tire & Rubber Co.

In Gauthier v. Goodyear Tire & Rubber Co., No. 1:23-cv-00281, 2024 WL 4882651 (E.D. Tex. Nov. 25, 2024), a plaintiff's attorney submitted a response to a motion for summary judgment that cited two nonexistent cases and contained multiple fabricated quotations attributed to other cases. Opposing counsel spent considerable time searching for the phantom authorities and ultimately brought the issue to the court's attention in a reply brief. The submitting attorney failed to address the problem until the court issued a show-cause order.

When the attorney acknowledged before the court that he had used a generative AI tool, specifically identified as "Claude," without verifying its output, the court imposed a \$2,000 monetary sanction and directed the attorney to complete a continuing legal education course on the use of generative AI in legal practice, and to serve copies of the court's order on his client. Judge Marcia Crone determined that the failure to verify the accuracy of the cited case law constituted a breach of Rule 11(b)(2).

Gauthier has been widely cited in subsequent AI sanctions decisions, in part because of the relatively modest monetary sanction but also because of the CLE component, a remedial measure that addresses the underlying competence deficit that AI hallucination cases frequently reveal. The decision also highlights the collateral costs of unverified AI use: opposing counsel wasted substantial time searching for fictitious cases, a cost ultimately borne by the opposing party.

Johnson v. Dunn

Johnson v. Dunn, No. 2:21-cv-1701-AMM (N.D. Ala. 2025), arose from high-stakes civil rights litigation concerning conditions in Alabama state prisons. Attorneys from a major national law firm filed routine discovery motions, a motion for leave to depose an incarcerated witness, and a motion to compel interrogatory responses that contained five fabricated case citations. The errors illustrated the full range of AI hallucination: some citations combined a real case name with a fabricated reporter citation; others cited cases that did not exist in any form.

The court's response was notable. Rather than relying primarily on monetary sanctions, the court imposed significant non-monetary measures, including disqualification of the offending attorneys from further participation in the case. The decision reflects a broader concern with deterrence and the integrity of judicial proceedings, and the court further directed that the matter be referred to appropriate disciplinary authorities. This approach underscores the judiciary's increasing willingness to impose structural and professional consequences in response to AI-related misconduct.

The decision also raises important questions regarding institutional responsibility within law firms. Although the court focused primarily on the conduct of individual attorneys, its analysis implicitly reflects the central role of supervision and internal controls in preventing the submission of inaccurate legal authorities. This suggests that law firms that implement robust governance measures, including training, supervision, and verification protocols, may be better positioned to mitigate risk. Accordingly, while the opinion does not directly address firm-wide liability based on internal policies, it

provides a framework for understanding how such measures could influence future courts' assessments of responsibility. See section VII, below.

Johnson v. Dunn reflects a potential escalation in the severity of sanctions courts may impose in response to litigation misconduct. Disqualification from pending litigation carries consequences for the client that monetary penalties do not, underscoring the seriousness of the court's response. Although the court did not explicitly hold that monetary sanctions are categorically inadequate, its reliance on non-monetary remedies suggests a growing willingness to employ more substantial measures where traditional sanctions may be insufficient to deter noncompliance.

Lexos Media Ip LLC v. Overstock.Com, Inc.

In Lexos Media IP LLC v. Overstock.com, Inc., No. 2:22-cv-02324-JAR (D. Kan. 2026), the United States District Court for the District of Kansas confronted one of the more complex AI hallucination cases in the developing jurisprudence. The case involved a patent infringement action in which plaintiff's counsel filed briefing containing an unusually broad range of AI-generated errors: citations to wholly nonexistent cases, citations to real cases that stood for propositions opposite to what the brief claimed, fabricated quotations attributed to actual judges, and a citation to a nonexistent lawsuit against the City of Topeka.

Judge Julie A. Robinson issued a detailed memorandum and order on February 2, 2026, imposing monetary sanctions totaling \$12,000 against five attorneys involved in the defective filing and revoking the pro hac vice admission of one attorney. The sanctions were apportioned based on each attorney's degree of responsibility: the attorney who drafted the offending brief using unverified AI output received the highest individual penalty, while supervisory attorneys who approved the filing without verifying the cited authorities received lesser sanctions reflecting their independent obligations.

The court's treatment of supervisory responsibility is an important contribution to the doctrine. It establishes that a signing attorney cannot escape Rule 11 liability by attributing the error to a colleague or to a legal research service. The obligation to verify the legal authorities cited in a filing runs to every attorney whose signature appears on the document. This principle that citation verification is non-delegable, has been stated in various formulations across the AI sanctions jurisprudence, and Lexos Media provides its most thorough elaboration.

Gardner v. Combs

Gardner v. Combs, No. 2:24-cv-07729 (D.N.J. Dec. 15, 2025), arose from a civil action in which a plaintiff alleged sexual assault against Sean [Diddu] Combs. The plaintiff's attorney submitted a responsive brief relying on a nonexistent case and fabricated legal propositions that the attorney acknowledged were generated by AI.

Judge Leo Gordon imposed a \$6,000 sanction on the attorney and required him to self-report the court's order to bar licensing authorities in New Jersey and New York, to serve copies on his client, and to discuss the implications of the order with the client. The court's order noted that the attorney "confirmed that he indeed used AI that provided him with hallucinated case law and propositions that he then incorporated into his response brief and that he failed to verify on subsequent review." The court's emphasis on the attorney's post-submission failure to verify is significant: even if one accepts that the initial AI-assisted drafting created some risk of error, the failure to review the submission before relying on it in subsequent proceedings compounded the original breach.

This case attracted considerable public attention because of the high-profile nature of the underlying litigation. Its inclusion in a body of case law otherwise dominated by commercial disputes underscores the universality of the hallucination problem: AI tools do not calibrate their propensity to hallucinate based on the significance or public profile of the litigation.

Mid Central Operating Engineers Health & Welfare Fund v. Hoosiervac LLC

Mid Central Operating Engineers Health & Welfare Fund v. HoosierVac LLC, No. 2:24-cv-00326-JPH-MJD, 2025 WL 574234 (S.D. Ind. Feb. 21, 2025), is remarkable for the pattern of conduct it documented: the responsible attorney filed briefs containing AI-generated fictitious citations on three separate occasions in the same case. As the attorney himself admitted, and as the court expressly found, he had relied on "generative artificial intelligence tools that produced fictitious case citations" without making "any attempt to verify the existence of the generated citations."

A magistrate judge recommended a \$15,000 sanction; the district court, on de novo review in May 2025, reduced the sanction to \$6,000 after crediting the attorney's corrective measures and acknowledging the reputational harm already suffered from the proceedings. The court nonetheless ordered referral to the Indiana attorney disciplinary commission.

The multiple-offense aspect of HoosierVac sets it apart from the typical AI sanctions case, which involves a single submission. That an attorney persisted in filing AI-generated, unverified citations after the first instance had been challenged suggests either a fundamental misunderstanding of the problem or a willingness to continue the practice despite known risks. The court's decision to reduce the sanction in light of remedial steps, while maintaining the disciplinary referral, reflects the dual function of AI-related sanctions: deterrence and rehabilitation.

Picon-Diaz v. Bondi

In *Picon-Diaz v. Bondi*, No. 25-9530 (10th Cir. Feb. 13, 2026), the Tenth Circuit admonished counsel for citing a fictitious Tenth Circuit ruling in an immigration brief, a citation attributed to AI-generated research that counsel failed to verify prior to submission. The court declined to impose monetary sanctions, but issued a formal admonishment and directed counsel to take corrective action. The case is one of several in which the Tenth Circuit addressed AI hallucination issues in immigration proceedings in the period from late 2025 through early 2026, reflecting the particular vulnerability of high-volume immigration practice to the risks of unverified AI use.

Evolving Federal Responses: Local Rules and AI Certification Requirements

In response to the proliferating sanctions cases, numerous federal district courts have adopted local rules, standing orders, or judicial practice requirements specifically addressing the use of AI in court filings. These approaches vary considerably in their scope and specificity, but most share a core requirement: attorneys must either certify that no AI-generated content was used in preparing the filing, or certify that any AI-generated content has been independently verified for accuracy before submission.

Among the more detailed implementations, Judge Brantley Starr of the Northern District of Texas requires attorneys to certify that any AI-drafted language has been "verified with traditional legal research tools." The District of New Jersey's approach, adopted by Judge Evelyn Padin, requires disclosure of the specific AI tool used, identification of the portions of the filing affected, and certification of human attorney review. Judge Scott L. Palk of the Western District of Oklahoma requires both disclosure of AI use and affirmative certification of accuracy for all citations and legal authority.

These requirements represent an important complement to the sanctions jurisprudence. Certification requirements shift the verification obligation to the filing stage, before any error reaches the court, and create a clear record that allows the court to assess whether a subsequent hallucination constituted a knowing violation of a specific court order. They also provide attorneys with a concrete, checklist-style reminder of their verification obligations at the moment of filing.

State Court Responses

California: *Noland v. Land of the Free, L.P.*

California's first published appellate opinion addressing AI hallucinations in court filings, *Noland v. Land of the Free, L.P.*, 336 Cal. Rptr. 3d 897 (Cal. Ct. App. 2025), arose from an appeal in which plaintiff's counsel had submitted opening briefs containing fabricated legal quotations on a significant scale. An investigation revealed that of twenty-three case quotations in the opening brief, twenty-one were fabricated, generated by AI tools and never appearing in any published decision.

The California Court of Appeal, Second Appellate District, Division Three, imposed a \$10,000 sanction on plaintiff's counsel, ordered counsel to serve the court's opinion on his client, and directed the Clerk to notify the State Bar. The court elected to

publish its opinion as "a warning" to California attorneys, notwithstanding the absence of novel legal questions on the merits, reflecting a determination that the hallucination problem warranted statewide notice and guidance.

The court stated that "[n]o brief, pleading, motion, or any other paper filed in any court should contain any citations, whether provided by generative AI or any other source, that the attorney responsible for submitting the pleading has not personally read and verified." This standard, personal reading and verification, is more demanding than mere existence-checking, and aligns with the Second Circuit's standard in *Park v. Kim*.

In re Domestic Partnership of Campos & Munoz

In *In re Domestic Partnership of Campos & Munoz*, No. D085584 (Cal. Ct. App., 4th Dist., Mar. 5, 2026), the California Court of Appeal addressed a novel procedural twist in the AI hallucination context: a party who challenged a trial court order on the ground that it relied on fictitious cases was found to have forfeited that very challenge because he had drafted and submitted the order himself.

The underlying dispute arose from a request for shared custody of a pet dog following the dissolution of a domestic partnership. Respondent's counsel cited two nonexistent cases drawn from a Reddit post about pet custody in support of the position that stability and emotional wellbeing should govern the outcome. When the court commissioner directed the appellant to prepare the proposed order, he incorporated the fictitious citations without objection; the court adopted the order as submitted.

The appellant later appealed, challenging the order on the ground that it rested on fabricated authority. The Court of Appeal, in a unanimous published opinion authored by Justice Martin N. Buchanan, held that although the trial court "erred by citing and relying in material part on fictional cases in its written order," the appellant had forfeited the objection by preparing and submitting the order and failing to alert the court to the problem in a timely manner. The court nonetheless imposed \$5,000 in sanctions on respondent's counsel, Roxanne Chung Bonar, who had first insisted the cited cases were real, producing additional fictitious citations in defense of her position, before eventually

acknowledging during oral argument that she had used AI tools and online resources, including Reddit, to conduct her legal research.

Hawaii: Keaau Development Partnership LLC v. Lawrence, Jr.

Keaau Development Partnership LLC v. Lawrence, Jr., No. CAAP-24-0000494 (Haw. Ct. App. 2025), addressed a case in which an attorney cited a nonexistent case without attempting to read it or confirm its holding. After a show-cause proceeding under Hawaii Rules of Civil Procedure Rule 11(c)(1)(B), the court determined that the citation of a nonexistent case without prior examination was unreasonable and constituted a violation of Rule 11(b)(2).

The court imposed a \$100 sanction personally on counsel, a modest amount that reflected the limited scale of the misconduct in this particular case, but noted that federal courts have imposed sanctions ranging from \$1,000 to \$5,000 in analogous circumstances. The court's opinion explicitly juxtaposed its analysis with federal AI-citation precedents, including *Mata*, *Hayes*, *Gauthier*, and related decisions, providing a comparative framework for Hawaiian practitioners.

Keaau is notable for the breadth of its doctrinal survey despite the relatively minor sanction imposed, and for the Hawaii Supreme Court's subsequent determination that existing rules provide adequate safeguards against AI misuse without requiring new regulatory intervention, a conclusion that reflects continuing judicial disagreement about whether supplementary AI-specific rules are necessary.

Louisiana: In re Sanctions Order of Kenney, Kerry

In re Sanctions Order of Kenney, Kerry, No. 25-C-389 (La. Ct. App. 5th Cir. Oct. 23, 2025), involved an attorney who submitted filings that cited cases with genuine reporter and volume numbers but fictitious page numbers, a subtler form of hallucination in which existing publications were corrupted with fabricated internal references. Defense counsel, unable to locate the cited material on Westlaw or on the courts' own websites, contacted the First and Third Circuit Courts of Appeal directly; those courts confirmed they had no record of the referenced decisions.

The Louisiana Fifth Circuit Court of Appeal affirmed the lower court's sanctions order, imposed costs on counsel, mandated three hours of continuing legal education, and reported counsel to the Louisiana disciplinary authorities. The court's opinion included an unusually direct statement of the attorney's verification duty: "Read the law you cite. Read the Code. Read the statutes. Read the cases. Citation to fake or fabricated cases in pleadings is prima facie evidence that the attorney who signs the pleadings has failed her duty. This is a responsibility that cannot be outsourced."

The court's emphasis on the non-delegable nature of citation verification, articulated with particular rhetorical force here, has been widely quoted and reflects the emerging consensus across jurisdictions.

Maryland: Mezu v. Mezu

Maryland's first published appellate opinion addressing AI hallucinations in legal filings, *Mezu v. Mezu*, No. 361, Sept. Term 2025 (Md. App. Ct. Oct. 29, 2025), arose from a domestic relations proceeding. An attorney collaborated with a law clerk who used ChatGPT and a separate AI legal research tool to draft the brief. The attorney did not read the cases cited; nor, by the clerk's own acknowledgment, had the clerk read them.

The Appellate Court of Maryland found that the brief contained citations to multiple fictitious cases and citations to real cases that did not support the propositions for which they were cited. Counsel's conduct "required this Court to take time to try to find the fake cases cited in Mother's brief, and then research how other courts have dealt with situations involving similar attorney misconduct, diverting judicial resources from other pressing work."

The court declined to impose a monetary sanction. In a significant observation in footnote 7, the court noted that neither party had sought monetary sanctions and that it lacked citation to any authority permitting the court to award sanctions sua sponte. The court accordingly referred the responsible attorney to the Attorney Grievance Commission for potential violations of Maryland Rules 19-301.1 (Competence), 19-303.1 (Meritorious Claims and Contentions), and 19-305.3 (Responsibilities Regarding Nonlawyer Assistance). The court further noted in footnote 7 that the General Assembly and the Rules Committee "may want to address whether a more specific statute or rule permitting this Court to award sanctions in cases such as the present case is warranted" — an express recognition that Maryland lacks the statutory sanctions authority that California, Georgia, and Utah already possess, and that this gap constrained the court's remedial response.

Mezu v. Mezu is significant as the first Maryland appellate case to address this issue and for the court's express invitation to legislative and rulemaking action. It establishes that the existing Maryland ethical framework is sufficient to support disciplinary consequences for AI hallucination misconduct, while signaling institutional openness to more specific regulatory guidance.

Utah: Garner v. Kadince, Inc.

Garner v. Kadince, Inc., No. 20250188-CA (Utah Ct. App. 2025), arose from a contract dispute in which the petitioner's attorney filed an appellate petition citing multiple cases that either did not exist or were not applicable to the propositions

asserted. The brief had been drafted by an unlicensed law clerk using ChatGPT; the supervising attorney signed the petition without independently checking its accuracy.

The Utah Court of Appeals imposed a compound sanction: the signing attorney was ordered to reimburse his client for fees associated with the defective petition, to pay opposing party's attorneys' fees incurred in responding to the petition, and to make a \$1,000 contribution to "and Justice for all," a Utah nonprofit organization promoting equal access to justice. The court stated that the use of AI tools is not inherently improper, but that Utah's Rules of Civil Procedure require attorneys to ensure the accuracy of their filings, a responsibility that survives the delegation of research and drafting to support staff or AI platforms.

Garner introduces the element of charitable contribution as a sanctioning mechanism, complementing more conventional monetary and reputational sanctions. It also underscores the supervisory dimension of the problem: the attorney who signed the petition, rather than the law clerk who drafted it, bore the ultimate legal responsibility for its accuracy.

The Distinct Problem: AI Misrepresentation of Authentic Case Law

The cases surveyed in Parts II and III all involve, at their core, citations to nonexistent legal authority. But Generative AI presents a related and arguably more insidious problem: the accurate citation of real cases accompanied by a material mischaracterization of what those cases hold.

This form of AI error, sometimes described as "citation accuracy" hallucination, in contrast to "citation existence" hallucination, occurs when an AI system retrieves a real case name and citation but generates a description of the holding, a quotation from the opinion, or a summary of the facts that is inconsistent with the actual text of the decision. The attorney who submits such a brief may verify that the cited case exists, find the citation in Westlaw or Lexis, and yet fail to detect that the proposition attributed to the case is not supported, or is actively contradicted by the opinion's actual language.

Several features of the legal research process make this form of error particularly difficult to detect. First, citations to real cases are less likely to trigger skepticism than citations to wholly fictitious ones. An attorney who verifies that a citation appears in a legal database may conclude, without reading the opinion, that the citation is reliable. Second, in complex appellate matters, the volume of cited authority may make comprehensive review of each source impractical. Third, AI-generated case summaries are often presented with apparent confidence, using the linguistic markers of reliable legal authority ("the court held," "in a leading decision," "the Second Circuit has consistently ruled") in a manner that discourages further inquiry.

Courts have noted this problem in several of the cases discussed above. In *Lexos Media*, the court found that some of the defective citations pointed to real cases that held the opposite of what the brief asserted, a form of misrepresentation distinguishable from outright fabrication but equally inconsistent with Rule 11's certification requirements. In *United States v. Hayes*, the court noted that "... misrepresenting the holdings of a case is making a false statement to a court," no less than citing a nonexistent case.

The solution to the citation-accuracy problem is the same as the solution to the citation-existence problem: attorneys must read the cases they cite. Existence verification alone, confirming that a case appears in a legal database, is a necessary but not sufficient component of the attorney's obligation. Rule 11 and its analogs require that each cited authority actually support the proposition for which it is cited, and that verification cannot be delegated to an AI tool or satisfied by database existence checks. As the Second Circuit stated in *Park v. Kim*, the duties of Rule 11 "require that attorneys read and thereby confirm the existence and validity of, the legal authorities on which they rely."

This standard has implications for how attorneys integrate AI tools into their research workflows. An AI-generated case summary, however confident in its presentation, should be treated as a research starting point rather than a research conclusion. The attorney's obligation is not merely to confirm that the AI pointed to a real case but to confirm, by reading the opinion, that the case says what the AI claims it says.

The Ethical Framework Governing Attorney Use of AI Tools

The sanctions cases surveyed in Parts II and III are grounded in procedural rules, primarily Federal Rule of Civil Procedure 11 and its state-court analogs. But the conduct those rules address also engages the professional responsibility obligations that govern attorney conduct more broadly. Courts and disciplinary authorities have consistently recognized that the use of AI tools does not modify the professional duties that attorneys owe to their clients, to opposing parties, and to the tribunal.

Competence

The foundational ethical obligation in the AI hallucination context is competence. Model Rule 1.1 of the ABA Model Rules of Professional Conduct requires that a lawyer "provide competent representation to a client," which "requires the legal knowledge, skill, thoroughness and preparation reasonably necessary for the representation." Comment 8 to Rule 1.1 further specifies that "[t]o maintain the requisite knowledge and skill, a lawyer should keep abreast of changes in the

law and its practice, including the benefits and risks associated with relevant technology.”

This “technological competence” requirement is directly implicated by AI use in legal practice. An attorney who uses AI-powered legal research tools without understanding their limitations, including their propensity to hallucinate, is not providing the “thoroughness and preparation reasonably necessary” for competent representation. Competence in the AI era requires not merely technical familiarity with AI tools, but a working understanding of when and how those tools fail, and the implementation of verification protocols that account for those failures.

Diligence

Model Rule 1.3 requires that a lawyer “act with reasonable diligence and promptness in representing a client.” Diligence in the context of AI-assisted legal research requires the attorney to implement a verification workflow sufficient to catch hallucinations before they reach the court. The submission of a brief containing AI-generated fictitious citations, without independent verification, is a failure of diligence regardless of the attorney’s subjective intent.

The courts that have imposed sanctions in AI hallucination cases have uniformly treated the failure to verify as a breach of the reasonable care standard that diligence requires. In *Mata*, the court found that the attorneys “abandoned their responsibilities” by filing without verification. In *Park v. Kim*, the court found that the attorney made no inquiry, much less the reasonable inquiry required by Rule 11 and long-standing precedent. These findings track the diligence requirement as closely as the competence requirement; the two duties reinforce each other in this context.

Candor Toward the Tribunal

Model Rule 3.3 prohibits a lawyer from knowingly making “a false statement of fact or law to a tribunal.” The AI hallucination cases raise questions about how the candor obligation interacts with conduct that the attorney may not have known was false at the time of submission. Courts have generally resolved this question by holding that an attorney who submits a pleading without exercising the diligence necessary to verify its accuracy is constructively aware of the falsity, that the failure to verify is itself the ethical breach, rather than a mitigating circumstance.

Moreover, several courts have found that conduct in the post-submission phase, when the attorney discovered or should have discovered the error, constitutes an independent violation of the candor duty. In *Hayes*, the court found that the attorney’s denial, when challenged by opposing counsel and the court, compounded the original violation. Rule 3.3(a)(1)’s prohibition on false statements is ongoing; once an attorney knows that a submission contains a false statement of law, the duty to correct it attaches.

Supervisory Responsibilities

Model Rule 5.1 requires supervising attorneys to make reasonable efforts to ensure that subordinate lawyers conform to the Rules of Professional Conduct, and Rule 5.3 imposes equivalent obligations with respect to nonlawyer assistants, including paralegals and law clerks. The hallucination cases demonstrate that both provisions are engaged whenever AI-assisted research or drafting is delegated to subordinate personnel.

In *Garner v. Kadince*, the supervising attorney’s failure to review the law clerk’s AI-generated brief before signing it constituted a breach of the supervisory duty under Rule 5.3. In *Lexos Media*, the court found that supervisory attorneys who approved a filing without verifying the cited authorities violated their independent Rule 11 obligations. These decisions establish that the existence of a subordinate who used AI does not insulate the supervising attorney from liability; supervision entails verification.

ABA Formal Opinion 512

In July 2024, the American Bar Association Standing Committee on Ethics and Professional Responsibility issued Formal Opinion 512, the ABA’s first comprehensive ethics guidance on the use of generative AI in legal practice. Formal Opinion 512 addresses, among other topics, the duties of competence, confidentiality, and candor as applied to AI tools. On the competence question, the Opinion states that Model Rule 1.1 requires attorneys to understand the benefits and risks of AI tools used to deliver legal services and to implement reasonable verification procedures for AI-generated outputs.

Formal Opinion 512 cautions that AI-generated legal research is particularly susceptible to errors that can be difficult to detect without careful review and that attorneys must verify AI-generated citations through reliable sources before relying on them in any client matter or court submission. The Opinion confirms that AI tools cannot replace the judgment, experience, and ethical obligations of the lawyer and that the submission of an AI-generated document without adequate review constitutes a failure of professional responsibility regardless of the AI tool’s sophistication.

Formal Opinion 512 does not resolve every question about the ethical use of AI, but it provides an authoritative statement of national consensus: AI is a permissible tool in legal practice, but its use does not alter or diminish the professional duties that govern attorney conduct. The attorney who uses AI remains responsible for the accuracy of every word submitted to the court.

Procedural Rules: Federal Rule of Civil Procedure 11

Federal Rule of Civil Procedure 11(b)(2) requires that an attorney's signature on a pleading certify that "the claims, defenses, and other legal contentions therein are warranted by existing law or by a nonfrivolous argument for the extension, modification, or reversal of existing law or the establishment of new law." This certification standard is the bedrock of the sanctions jurisprudence surveyed in this Article.

Compliance with Rule 11(b)(2) requires that the certifying attorney conduct a reasonable inquiry into the legal authorities cited before filing. As courts have consistently held, the hallucination-vulnerable nature of AI-generated legal research means that reliance on AI output without independent verification does not constitute a reasonable inquiry for Rule 11 purposes. The attorney must do more than ask an AI tool for case law supporting a proposition; the attorney must verify each cited case by reading the actual opinion. State procedural rules analogous to Rule 11 impose the same requirements. Maryland Rule 1-311, for example, governs the certification of filings in Maryland courts and establishes standards substantially equivalent to those of federal Rule 11.

Maryland's Current Legal Framework Maryland Rules of Professional Conduct

Maryland has adopted ethics rules that are substantially based on the ABA Model Rules of Professional Conduct. The rules most directly implicated by AI hallucination misconduct are:

Maryland Rule 19-301.1 (Competence) tracks Model Rule 1.1, requiring that a lawyer "provide competent representation to a client" through the exercise of the "legal knowledge, skill, thoroughness and preparation reasonably necessary for the representation." Comment 8, as adopted in Maryland, includes the admonition that lawyers should remain current on changes in legal practice, including technology-related developments—a provision that requires Maryland attorneys to be aware of both the utility and the limitations of AI tools.

Maryland Rule 19-301.3 (Diligence) requires attorneys to "act with reasonable diligence and promptness in representing a client." As in the ABA Model Rules, this obligation encompasses the verification of legal authorities cited in any court submission.

Maryland Rule 19-303.3 (Candor Toward the Tribunal) prohibits attorneys from knowingly making "a false statement of fact or law to a tribunal" and requires disclosure of directly adverse authority in the controlling jurisdiction. The submission of AI-hallucinated citations implicates this rule both at the initial submission stage and at any subsequent stage at which the attorney becomes aware of the falsity.

Maryland Rule 19-305.3 (Responsibilities Regarding Nonlawyer Assistance) requires supervising attorneys to ensure that work product generated by nonlawyer assistants conforms to the Rules of Professional Conduct. This provision directly addresses the supervisory failures documented in *Garner*, *Lexos Media*, and similar cases.

Maryland Rule 19-308.4 (Misconduct) prohibits conduct involving "dishonesty, fraud, deceit, or misrepresentation." Depending on the circumstances, the submission of fabricated AI-generated citations may constitute misconduct under this provision, particularly where the attorney knew or should have known that the cited authority did not exist.

Maryland Rule 1-341

Maryland Rule 1-341 provides a procedural mechanism for the imposition of sanctions in civil actions, authorizing the court to require a party or its attorney to pay costs and reasonable attorneys' fees to the opposing party if the court finds that the conduct of the sanctioned party "was in bad faith or without substantial justification." Maryland courts have characterized Rule 1-341 as an "extraordinary remedy" to be invoked in "rare and exceptional cases," with a "patently frivolous" standard for sanctionable legal argument.

The submission of AI-hallucinated citations would plainly satisfy these standards in most circumstances. A legal argument based on nonexistent authority is, by definition, not "warranted by existing law;" it is, in the most literal sense, an argument based on law that does not exist. Where the attorney knew or should have known that the cited cases were fictitious, and the attorney's failure to read the cases will almost always satisfy the "should have known" standard the conduct is without substantial justification.

Rule 1-341 may be particularly useful in AI hallucination cases where the opposing party has incurred costs in searching for phantom authorities, a cost that several of the surveyed cases recognize as a concrete and compensable harm. Courts applying Rule 1-341 can tailor monetary sanctions to reflect the actual costs imposed on the opposing party, providing a more precise compensatory remedy than flat sanctions imposed by order of the court.

Mezu v. Mezu (2025) as Maryland's First AI Citation Precedent

Mezu v. Mezu establishes several principles of particular importance for Maryland practitioners and courts. First, the existing Maryland Rules of Professional Conduct are sufficient to capture and address AI hallucination misconduct. The Appellate Court's referral of the responsible attorney to the Attorney Grievance Commission under Rules 19-301.1, 19-303.1, and 19-305.3 confirms that no new rule is required to impose professional consequences for this conduct.

Second, AI-related misconduct consumes judicial resources even where no monetary sanction is imposed. The court expressly noted that counsel's conduct "required this Court to take time to try to find the fake cases cited in Mother's brief, and then research how other courts have dealt with situations involving similar attorney misconduct, diverting

judicial resources from other pressing work.” These are real costs to the administration of justice — costs that, when aggregated across the many cases in which courts have been compelled to conduct analogous inquiries, represent a significant and growing burden on the judicial system.

Third, the court’s observation in footnote 7 that the General Assembly and the Rules Committee “may want to address” whether a statute or rule permitting the court to award sanctions without a party’s request “is warranted” reflects a gap in Maryland’s existing enforcement framework. Unlike California and Georgia, Maryland has no rule expressly authorizing an appellate court to impose monetary sanctions sua sponte, which constrained the court’s response in this case notwithstanding the severity of counsel’s misconduct.

Fourth, *Mezu v. Mezu* joins a growing body of precedent establishing that the attorney’s verification obligation runs to the actual content of each cited case, not merely to its existence. An attorney who submits a brief without reading the cited authorities has not satisfied the duty of competence regardless of whether those authorities exist.

Firm Liability and AI Governance

The emerging jurisprudence on AI-generated hallucinated citations has focused primarily on the conduct of individual attorneys. Decisions such as *Mata v. Avianca, Inc.* and *United States v. Hayes* emphasize the personal responsibility of the signing lawyer to ensure that all cited authorities exist and accurately reflect the law. Yet these cases also raise a broader and increasingly important question: to what extent may law firms themselves bear responsibility for failures arising from the use of generative artificial intelligence? We will answer this question in Section A-F, *infra*.

The Doctrinal Baseline: Individual Responsibility with Institutional Dimensions

Under existing doctrine, responsibility for filings rests most directly with the attorney who signs the pleading. Rule 11 requires that counsel conduct a reasonable inquiry into both the facts and the law before submission, and courts have consistently held that this obligation cannot be satisfied through blind reliance on third-party sources, including AI tools. In *Mata*, the court stressed that existing rules impose a gatekeeping role on attorneys to ensure the accuracy of their filings. 678 F. Supp. 3d 443, 448 (S.D.N.Y.2023). Similarly, in *Hayes*, the court made clear that citing nonexistent or misrepresented authority constitutes a false statement to the tribunal. 763 F. Supp. 3d 1054, 1060 (E.D. Cal. 2025).

At the same time, these decisions implicitly recognize that legal practice is not purely individual. Attorneys operate within institutional environments, such as law firms that shape research practices, supervision, training, and technological adoption. The increasing use of AI tools, therefore, introduces not only questions of individual diligence but also issues of organizational competence and supervision.

Supervisory Duties Under Existing Ethical Rules

The framework for assessing firm-level responsibility already exists within the Rules of Professional Conduct. In most jurisdictions, including Maryland, supervisory rules require law firms and managing attorneys to ensure that all lawyers conform to professional obligations. These duties include implementing reasonable measures to ensure compliance with the rules governing competence, diligence, and candor.

When applied to AI tools, these supervisory obligations take on heightened significance. If generative AI systems are known to produce fabricated or inaccurate legal authorities, then a failure to implement safeguards may raise questions about whether the firm has exercised reasonable oversight. In this respect, AI does not create a new category of ethical duty; rather, it amplifies existing ones.

The Emerging Role of AI Governance Frameworks

Although courts have not yet developed a comprehensive doctrine of firm liability for AI-related misconduct, the early cases suggest that institutional practices will become increasingly relevant. As courts confront repeated instances of hallucinated citations, they are likely to look beyond individual error and examine whether adequate systems were in place to prevent such failures.

This points toward the growing importance of AI governance frameworks within law firms. Such frameworks may include:

- Training attorneys on the limitations of generative AI tools;
- Requiring independent verification of all AI-generated citations;
- Establishing review protocols for filings that incorporate AI-assisted research; and
- supervising the use of AI tools in high-risk contexts, such as dispositive motions and appellate briefs.

While no reported decision has yet held that the presence or absence of such policies is dispositive of firm liability, the logic of existing cases strongly suggests that these factors could influence future courts’ assessments of responsibility.

Mitigation, Deterrence and Institutional Accountability

The sanction decisions to date reflect a broader concern with deterrence. Courts have imposed monetary sanctions, disqualification, and disciplinary referrals to address the risk that AI-generated errors may proliferate if left unchecked.

Within this framework, firm-level practices may become relevant not only to liability, but also to mitigation.

A law firm that can demonstrate robust training, supervision, and verification procedures may be better positioned to argue that an individual attorney's failure represents a deviation from established practice rather than a systemic deficiency. Conversely, the absence of such safeguards could support a finding that the misconduct was foreseeable and preventable.

Thus, even in the absence of explicit doctrinal development, courts are likely to view AI governance as part of the broader inquiry into whether counsel acted reasonably under the circumstances.

Toward a Coherent Doctrine of Firm Responsibility

The current case law reflects an early stage of doctrinal development. Courts have focused on the most immediate problem, submission of fictitious or inaccurate authorities, but have not yet fully articulated the role of institutional responsibility. Nevertheless, the trajectory is clear. As AI tools become more deeply integrated into legal practice, courts will increasingly confront situations in which responsibility cannot be understood solely at the individual level.

A coherent doctrine of firm liability in this context would likely build on three principles. First, the duty to verify legal authority remains personal to the signing attorney. Second, law firms have an independent obligation to implement reasonable systems of supervision and control. Third, the adequacy of those systems may influence both liability and the severity of sanctions.

Implications for Practice

For law firms, the implications are immediate. Generative AI should not be treated as an ordinary research tool without risk. Instead, it should be approached as a technology that requires structured oversight. Firms that proactively adopt governance measures—training, supervision, and verification protocols, will not only reduce the likelihood of error, but also strengthen their position if courts later evaluate the reasonableness of their conduct.

For courts, the challenge will be to balance deterrence with fairness. Sanctioning individual attorneys remains necessary, but it may not be sufficient where institutional practices contribute to the risk of error. As the jurisprudence develops, courts may increasingly consider whether failures in AI use reflect individual negligence, systemic shortcomings, or both.

Recommendations for Maryland Courts and the Maryland Bar Affirm that Appropriate AI Use in Legal Practice is Permitted

Maryland courts and bar authorities should begin from a clear and explicit premise: the use of generative AI as a drafting and research tool is a legitimate component of modern legal practice. Attorneys are not required to avoid AI; they are required to use it responsibly. This premise acknowledges the genuine utility of AI in improving access to legal services and reducing the cost of representation. It avoids overregulation that would disadvantage Maryland practitioners relative to those in other jurisdictions. And it focuses regulatory attention on the actual problem, unverified use, rather than on AI use per se.

This affirmative stance toward responsible AI use should be stated clearly in any guidance or rule that Maryland adopts, consistent with the approach taken by the ABA in Formal Opinion 512 and by courts such as the Utah Court of Appeals in *Garner*, which stated that “the use of reliable AI tools is not improper in itself.”

Adopt AI Certification Requirements Analogous to Those in Leading Federal Districts

Maryland courts should promulgate a rule or standing order requiring attorneys who use AI tools in preparing any filing to certify either (1) that no portion of the filing was generated by AI, or (2) that all AI-generated content has been independently verified for accuracy, including that every cited authority has been personally read and verified by the signing attorney or a supervising attorney who has read the opinion. This certification requirement should be adopted as a uniform statewide practice, rather than being left to individual judicial officers.

A certification requirement has several advantages over a purely sanctions-based enforcement approach. It brings the verification obligation to the attorney's attention at the moment of filing, when it can most easily be satisfied. It creates a clear record of claimed compliance that courts can evaluate in subsequent sanctions proceedings. It provides attorneys with a concrete framework for thinking about their verification obligations. And it signals to the bar that Maryland courts take the issue seriously, encouraging law firms to develop internal AI governance policies and verification protocols.

Enforce Existing Procedural Rules Vigorously

Maryland courts should make clear that the existing procedural framework—Maryland Rule 1-311 and Maryland Rule 1-341 is fully applicable to AI hallucination misconduct and that its provisions will be enforced with the same rigor as any other sanctions-rule violation. Courts should not hesitate to exercise their inherent authority to impose sanctions, require disclosure, and issue show-cause orders when the circumstances warrant, and should explicitly invoke *Mezu v. Mezu* as Maryland authority for the proposition that AI-generated fictitious citations are sanctionable under existing rules.

Courts should also consider the full range of available sanctions, not merely monetary penalties. The cases surveyed in this Article demonstrate that monetary sanctions have not, by themselves, eliminated the problem; courts in other jurisdictions have found it necessary to impose disqualification, disciplinary referral, and remedial education requirements to address persistent noncompliance. Maryland courts should be prepared to employ equally diverse remedies when the circumstances warrant.

Issue Definitive Bar Association Guidance

The Maryland State Bar Association and the Attorney Grievance Commission should issue comprehensive ethics guidance specifically addressing the use of AI in legal practice. Such guidance should address: (1) the competence obligations that attach to AI use, including the requirement to understand AI limitations; (2) the verification procedures that constitute reasonable care in the context of AI-generated legal research; (3) the supervisory obligations that arise when subordinates use AI in preparing client work product; (4) disclosure obligations to clients regarding AI use; and (5) best practices for implementing AI governance policies within law firms of varying sizes.

Such guidance would complement the *Mezu v. Mezu* decision by providing practical instruction to practitioners who understand their obligations in the abstract but need concrete guidance on how to satisfy them. Several state bar associations have already issued comparable guidance; Maryland should join them in providing the profession the practical framework it needs to use AI responsibly.

Consider Mandatory Continuing Legal Education

The Maryland Court of Appeals should consider requiring continuing legal education credits in technology competence, including coverage of AI tools and their limitations, as part of the mandatory CLE requirements for active Maryland attorneys. Several of the courts in the cases surveyed in this Article imposed CLE requirements as part of the sanctions order, reflecting a recognition that the underlying problem is often a competence deficit that remedial education can address. A proactive statewide requirement would address that deficit before it produces sanctionable conduct.

Reiterate the Non-Delegable Nature of Citation Verification

Whatever form Maryland's response takes, courts, bar authorities, and rules should reiterate the principle that citation verification is a non-delegable professional duty of the signing attorney. An attorney cannot satisfy this obligation by delegating it to a law clerk, a paralegal, an AI tool, or any combination of the three. The attorney who signs the filing is responsible for verifying every legal authority cited in it—by reading it, not merely by confirming its existence in a legal database.

This principle, articulated with particular force in *In re Kenney* ("this is a duty that cannot be delegated"), *Lexos Media, Noland, and Park v. Kim*, is the most important lesson of the AI sanctions jurisprudence for practicing attorneys. Its clear articulation in Maryland authority will provide an essential foundation for the enforcement of the other recommendations set forth in this Part.

Conclusion

The hallucination problem in AI-assisted legal research is not a temporary anomaly that improving technology will soon eliminate. AI systems may become more reliable over time, but they will continue to produce confident-sounding errors; the architecture of large language models does not support grounded retrieval in the way that conventional legal databases do. The profession must therefore adapt its verification practices to the capabilities and limitations of the tools it has adopted.

The courts that have addressed this challenge have done so with considerable consistency. From *Mata v. Avianca* in 2023 through the decisions surveyed in Parts II and III of this Article, courts across the country have delivered a clear and coherent message: the use of AI does not diminish the attorney's professional obligations, and the attorney who submits AI-generated content to a tribunal without independent verification will face sanctions, disciplinary referral, or both.

Maryland has entered this jurisprudence with *Mezu v. Mezu* (2025), and the foundation that decision provides is sound. But foundation is not structure. The hallucination problem has already reached Maryland courts and will continue to do so as AI tools become more deeply integrated into legal practice. Maryland courts and the Maryland bar should act now, through certification requirements, vigorous enforcement of existing rules, comprehensive ethics guidance, and a clear articulation of the non-delegable nature of citation verification, to ensure that the promise of AI-enhanced legal practice is realized without compromising the integrity of the adjudicative process.

The stakes are high. Courts depend on the accuracy of the legal authorities submitted to them. Opposing parties depend on the truthfulness of adversary submissions. Clients depend on the competence of the attorneys who represent them. AI hallucinations threaten all three dependencies. The profession's response must be equal to the challenge.

* The author thanks many of his colleagues and friends for their contributions to this Article. All errors remain the author's own. This article is written using, in part, Generative AI in the "proper" form, meaning all the cases cited herein were verified, and all the quotations are verbatim from the actual cases.

Footnotes

- ¹ The term “hallucination” in the AI context was adopted from cognitive science and refers to outputs generated by AI systems that are fluent and confident but factually unfounded. See Emily M. Bender, Timnit Gebru, Angelina McMillan-Major & Shmargaret Shmitchell, *On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?*, in *Proc. of the 2021 ACM Conf. on Fairness, Accountability, and Transparency* 610 (2021).
- ² See Maura R. Grossman, Paul W. Grimm & Daniel G. Brown, *Is Guidance for Generative AI in Legal Practice Appropriate?*, 107 *Judicature* 10 (2023).
- ³ See Damien Charlotin, *AI Hallucination Cases Database*, <https://www.damiencharlotin.com/hallucinations/> (last visited Mar. 31, 2026) (cataloging AI hallucination cases in legal proceedings across multiple jurisdictions).
- ⁴ See Paul Friedman & Lauren Mayer, *Common Issues That Arise in AI Sanction Jurisprudence and How the Federal Judiciary Has Responded to Prevent Them*, *ABA Bus. L. Today* (Sept. 2024).
- ⁵ See generally RIPS Law Librarian Blog, *When AI Hallucinates Courts Hammer Attorneys* (Feb. 28, 2025), <https://ripslawlibrarian.wordpress.com/2025/02/28/when-ai-hallucinates-courts-hammer-attorneys/>.
- ⁶ *Mata v. Avianca, Inc.*, 678 F. Supp. 3d 443, 447–49 (S.D.N.Y. 2023).
- ⁷ *Id.* at 448.
- ⁸ *Id.* at 466–67.
- ⁹ *Park v. Kim*, 91 F.4th 610, 611–12 (2d Cir. 2024).
- ¹⁰ *Id.* at 613.
- ¹¹ *Id.*
- ¹² *United States v. Hayes*, 763 F. Supp. 3d 1054, 1058 (E.D. Cal. 2025).
- ¹³ *Id.* at 1062.
- ¹⁴ *Id.* at 1059.
- ¹⁵ *United States v. Hayes*, No. 2:24-cr-0280-DJC, 2025 WL 1067323 (E.D. Cal. Apr. 9, 2025).
- ¹⁶ *Gauthier v. Goodyear Tire & Rubber Co.*, No. 1:23-cv-00281, 2024 WL 4882651, at *2–3 (E.D. Tex. Nov. 25, 2024).
- ¹⁷ *Id.* at *4–5.
- ¹⁸ *Johnson v. Dunn*, No. 2:21-cv-1701-AMM (N.D. Ala. May 19, 2025).
- ¹⁹ *Id.*
- ²⁰ *Id.* The court’s treatment of firm-level responsibility and the mitigating effect of responsible AI governance policies is analyzed in EDMR, *When AI Policies Fail: The AI Sanctions in Johnson v. Dunn and What They Mean for the Profession* (July 2025), <https://edrm.net/2025/07/when-ai-policies-fail-the-ai-sanctions-in-johnson-v-dunn-and-what-they-mean-for-the-profession/>.
- ²¹ *Lexos Media IP LLC v. Overstock.com, Inc.*, No. 2:22-cv-02324-JAR (D. Kan. Feb. 2, 2026). Practitioners should verify the current citation through official court databases.
- ²² *Id.*; see also *Kansas Federal Judge Penalizes 5 Lawyers for AI-Generated Fake Legal Citations*, *Minn. Law.* (Feb. 6, 2026).
- ²³ *Id.*
- ²⁴ *Gardner v. Combs*, No. 2:24-cv-07729 (D.N.J. Dec. 15, 2025).
- ²⁵ *Id.*
- ²⁶ *Mid Cent. Operating Eng’rs Health & Welfare Fund v. HoosierVac LLC*, No. 2:24-cv-00326-JPH-MJD, 2025 WL 574234, at *3 (S.D. Ind. Feb. 21, 2025).
- ²⁷ *Id.*; see also *Mid Cent. Operating Eng’rs Health & Welfare Fund v. HoosierVac LLC*, No. 2:24-cv-00326-JPH-MJD (S.D. Ind. May 28, 2025) (district court order reducing sanction to \$6,000).
- ²⁸ *Picon-Diaz v. Bondi*, No. 25-9530 (10th Cir. Feb. 13, 2026).
- ²⁹ See 10th Circuit Orders Lawyer to Pay \$1,000 for Faulty AI Citations, *Colo. Pol.* (Feb. 9, 2026).
- ³⁰ See generally *What You Need to Know: AI Disclosure Rules in Legal Filings*, *Eve Legal* (2025), <https://www.eve.legal/blogs/what-you-need-to-know-ai-disclosure-rules-in-legal-filings>.
- ³¹ See N.D. Tex., *Standing Order Regarding Artificial Intelligence* (Judge Brantley Starr).
- ³² See D.N.J., *Individual Practices of Judge Evelyn Padin*.
- ³³ See W.D. Okla., *Standing Order Regarding Use of AI in Court Filings* (Judge Palk).
- ³⁴ *Noland v. Land of the Free, L.P.*, 336 Cal. Rptr. 3d 897 (Cal. Ct. App. 2025); see also *California Appellate Court Issues \$10K Sanctions in State’s First Published Opinion on AI-Hallucinated Case Citations*, *McGuireWoods* (2025), <https://www.mcguirewoods.com/>.
- ³⁵ *Noland*, 336 Cal. Rptr. 3d at 897.
- ³⁶ *Id.*
- ³⁷ *Keaau Dev. P’ship LLC v. Lawrence, Jr.*, No. CAAP-24-0000494 (Haw. Ct. App. 2025).
- ³⁸ *Id.*
- ³⁹ *In re Sanctions Order of Kenney, Kerry*, No. 25-C-389 (La. Ct. App. 5th Cir. Oct. 23, 2025).
- ⁴⁰ *Id.*
- ⁴¹ *Mezu v. Mezu*, No. 361, Sept. Term 2025 (Md. App. Ct. Oct. 29, 2025), <https://www.mdcourts.gov/data/opinions/cosa/2025/0361s25.pdf>.
- ⁴² *Id.*
- ⁴³ *Garner v. Kadince, Inc.*, No. 20250188-CA (Utah Ct. App. 2025).
- ⁴⁴ *Id.*; see also *Utah Attorneys Sanctioned by Appeals Court for Using AI to Write Legal Brief*, *Deseret News* (June 5, 2025).

- ⁴⁵ See AI Hallucinations Are Real—and How to Avoid Them, ABA Litigation News (2024), <https://www.americanbar.org/groups/litigation/resources/newsletters/appellate-practice/ai-hallucinations-real-how-avoid-them/>.
- ⁴⁶ See Lexos Media IP LLC v. Overstock.com, Inc., No. 2:22-cv-02324-JAR (D. Kan. Feb. 2, 2026).
- ⁴⁷ United States v. Hayes, 763 F. Supp. 3d 1054, 1059 (E.D. Cal. 2025).
- ⁴⁸ Park v. Kim, 91 F.4th 610, 613 (2d Cir. 2024).
- ⁴⁹ Model Rules of Prof'l Conduct R. 1.1 (Am. Bar Ass'n 2023).
- ⁵⁰ Id. cmt. 8.
- ⁵¹ Model Rules of Prof'l Conduct R. 1.3 (Am. Bar Ass'n 2023).
- ⁵² Park v. Kim, 91 F.4th at 613.
- ⁵³ Model Rules of Prof'l Conduct R. 3.3(a)(1) (Am. Bar Ass'n 2023).
- ⁵⁴ See Mata v. Avianca, Inc., 678 F. Supp. 3d at 461–62; see also United States v. Hayes, 763 F. Supp. 3d at 1059.
- ⁵⁵ Model Rules of Prof'l Conduct R. 5.1, 5.3 (Am. Bar Ass'n 2023).
- ⁵⁶ Garner v. Kadince, Inc., No. 20250188-CA (Utah Ct. App. 2025); Lexos Media IP LLC v. Overstock.com, Inc., No. 2:22-cv-02324-JAR (D. Kan. Feb. 2, 2026).
- ⁵⁷ ABA Standing Comm. on Ethics & Prof'l Responsibility, Formal Op. 512 (July 2024) [hereinafter ABA Formal Op. 512].
- ⁵⁸ Id. at 4–5.
- ⁵⁹ Fed. R. Civ. P. 11(b)(2).
- ⁶⁰ See Mata v. Avianca, Inc., 678 F. Supp. 3d at 462–63; Gauthier v. Goodyear Tire & Rubber Co., 2024 WL 4882651, at *4; Mid Cent. Operating Eng'rs Health & Welfare Fund v. HoosierVac LLC, 2025 WL 574234, at *4.
- ⁶¹ Md. Rule 1-311.
- ⁶² Md. Rule 19-301.1.
- ⁶³ Md. Rule 19-301.3.
- ⁶⁴ Md. Rule 19-303.3.
- ⁶⁵ Md. Rule 19-305.3.
- ⁶⁶ Md. Rule 19-308.4.
- ⁶⁷ Md. Rule 1-341.
- ⁶⁸ See generally Md. Rule 1-341 cmt.; see also marylandlegalmalpracticeattorney.com (summarizing Rule 1-341 standards).
- ⁶⁹ Garner v. Kadince, Inc., No. 20250188-CA (Utah Ct. App. 2025).