

AI Ethics in Legal Practice: A Comprehensive Case Law and Governance Research Report (2023–April 2026)

Attorneys who use AI without verification face escalating consequences — from \$5,000 fines to full disqualification — while the profession grapples with whether AI tools destroy attorney-client privilege. This report compiles every significant U.S. court case involving AI-related sanctions, the landmark privilege-waiver rulings, current platform data practices, and the regulatory landscape as of April 2026. The research covers 20+ sanctions cases totaling over \$500,000 in penalties, two opposing federal privilege rulings issued on the same day, and an evolving patchwork of 300+ judicial standing orders, proposed federal rules, and state bar opinions. For Ohio practitioners, this report includes Ohio-specific guidance and flags the state’s notable absence of a formal bar ethics opinion on AI.

Part 1: The sanctions explosion — every major AI case from Mata to today

The pace of AI-related sanctions has accelerated dramatically since the landmark *Mata v. Avianca* decision in June 2023. Courts have moved from modest fines to attorney disqualification, and penalties have grown from \$5,000 to nearly \$100,000. What began as a cautionary tale about ChatGPT has expanded to encompass sanctions involving proprietary firm AI tools, Westlaw’s CoCounsel, Google Gemini, and vLex. **No attorney or firm is immune** — sanctioned parties include Butler Snow (400+ attorneys), K&L Gates (one of the nation’s largest firms), Morgan & Morgan, and Goldberg Segalla.

Mata v. Avianca, Inc. — the case that started it all

Mata v. Avianca, Inc., 678 F. Supp. 3d 443 (S.D.N.Y. 2023), Case No. 1:22-cv-01461 (PKC), decided June 22, 2023, before Judge P. Kevin Castel. Attorney Steven Schwartz used ChatGPT to research a personal injury brief, producing at least **six entirely fabricated case citations** (including *Varghese v. China Southern Airlines*, *Shaboon v. EgyptAir*, and others). When opposing counsel challenged the citations, Schwartz asked ChatGPT to confirm their existence — and the AI doubled down.  Schwartz then prepared an affidavit attaching fabricated “decisions” generated by ChatGPT. Co-counsel Peter LoDuca filed the

documents without reading them. [Association of Corporate Counsel](#) **Sanction: \$5,000** jointly and severally against both attorneys and their firm, Levidow, Levidow & Oberman P.C., [Wikipedia +2](#) plus mandatory notification letters to every judge falsely identified as an author of the fake opinions. [Justia](#) Judge Castel stated: "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." [Hvbba](#)

People v. Crabill — the first attorney suspension

People v. Crabill, Case No. 23PDJ067, 2023 WL 8111898 (Colo. O.P.D.J. Nov. 22, 2023). Colorado attorney Zachariah Crabill, [Uscourts](#) approximately two years out of law school, used ChatGPT to draft a motion containing [Legal Ethics Lawyer](#) "dozens" of fabricated citations. [Georgia State University Law School](#) He discovered the fabrications before a hearing but did not inform the court, and initially blamed an intern. [Ficlaw](#) **Sanction: one-year-and-one-day suspension** (90 days served, remainder stayed on two-year probation), [Joannhoffman](#) \$2,000 fine, and mandatory ethics training. [WebProNews](#) This was the first formal attorney disciplinary action for AI-generated fabrications.

Park v. Kim — the first federal appellate referral

Park v. Kim, 91 F.4th 610 (2d Cir. 2024), [American Bar Association](#) No. 22-2057-cv, decided January 30, 2024. Attorney Jae S. Lee filed an appellate reply brief [Uscourts](#) in a medical malpractice case [Justia News](#) citing *Matter of Bourguignon v. Coordinated Behavioral Health Services, Inc.*, 114 A.D.3d 947 (3d Dep't 2014) — a case that does not exist. Lee admitted she generated it using ChatGPT because she had difficulty finding supporting precedent. [Joseph Hage Aaronson](#) **Sanction: referral to the Court's Grievance Panel** [Uscourts](#) under Local Rule 46.2. [FindLaw](#) The Second Circuit held: "Attorney Lee's submission of a brief relying on non-existent authority reveals that she failed to determine that the argument she made was 'legally tenable.'" [Reason.com](#) This was the **first federal appellate case** to refer an attorney for discipline over AI-fabricated citations.

2024 sanctions cases

Kruse v. Karlen, No. ED111172 (Mo. Ct. App., E.D., Feb. 13, 2024). Pro se appellant hired an online "consultant" who used AI to prepare a brief in which **22 of 24 case citations were fictitious**. [American Bar Association +2](#) Sanction: **\$10,000** in damages toward respondent's attorneys' fees; appeal dismissed as frivolous. [STLPR](#) [LawSites](#) First Missouri appellate AI sanctions case. [Jefferson City News Tribune](#)

Smith v. Farwell (Norfolk Cty. Superior Court, Mass., Feb. 12, 2024). **\$2,000 sanction** against an attorney whose memoranda cited fabricated cases; the attorney blamed recent law school graduates. [LawSites](#)

Gauthier v. Goodyear Tire & Rubber Co., No. 1:23-CV-281, 2024 WL 4882651 (E.D. Tex. Nov. 25, 2024). **\$2,000 sanction** for AI-generated fake citations. [Clarkcountybar](#)

2025: The year sanctions became structural

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). Texas attorney Rafael Ramirez relied entirely on generative AI for legal research, filing three briefs with fictitious citations. When confronted, he claimed he did not know AI could fabricate cases. [Gstexlaw](#) Magistrate Judge Mark J. Dinsmore recommended \$15,000 (\$5,000 per brief); [RIPS Law Librarian Blog](#) **the final sanction was \$6,000.**

Wadsworth v. Walmart Inc., 348 F.R.D. 489 (D. Wyo. 2025), No. 2:23-cv-118-KHR, decided February 24, 2025, by Judge Kelly H. Rankin. Attorney Rudwin Ayala of Morgan & Morgan drafted motions in limine citing **8 of 9 nonexistent cases** [Bloomberg Law](#) — generated by the firm's proprietary AI program, not public ChatGPT.

[Esquire Deposition Solutions](#) **Sanctions: \$3,000 (Ayala, plus revocation of pro hac vice admission), \$1,000 each against two other attorneys.** The firm was not sanctioned because it took immediate remedial steps.

Lacey v. State Farm / Ellis George & K&L Gates (C.D. Cal.), sanctions order May 6, 2025, by Special Master (former Magistrate Judge) Michael R. Wilner. Attorney Trent Copeland used CoCounsel, Westlaw Precision, and Google Gemini to create a brief outline, then sent it to K&L Gates without disclosing AI origins. K&L Gates incorporated the material without verifying citations. [ABA Journal](#) [LawSites](#) **9 of 27 citations were wrong; at least 2 were entirely nonexistent.** [ABA Journal](#) A "corrected" brief still contained six AI errors.

[Reason.com](#) **Sanction: \$31,100** jointly and severally [LawSites](#) (\$26,100 in ADR fee reimbursement plus \$5,000 for costs). [Complete AI Training](#) [Spellbook](#) Wilner wrote: "I read their brief, was persuaded (or at least intrigued) by the authorities that they cited, and looked up the decisions to learn more about them — only to find that they didn't exist. That's scary." [LawSites](#)

Coomer v. Lindell/MyPillow, No. 1:22-cv-01129-NYW (D. Colo., July 7, 2025). Attorneys for Mike Lindell/MyPillow produced **nearly 30 faulty citations** with AI assistance. **Sanction: \$3,000 each** against attorneys Christopher Kachouroff and Jennifer DeMaster. [Newzzy](#)

[NPR](#)

Johnson v. Dunn — the disqualification watershed

Johnson v. Dunn, No. 2:21-cv-1701-AMM (N.D. Ala., July 23, 2025), 2025 U.S. Dist. LEXIS 141805, before Judge Anna M. Manasco. Plaintiff Frankie Johnson, an Alabama inmate stabbed multiple times, sued former DOC Commissioner Jefferson Dunn. Butler Snow LLP (paid \$42M+ by the state since 2020 for prison cases) represented defendants. [AIC](#)

Attorney Matthew Reeves used ChatGPT to find citations and inserted them without verification; false citations appeared in two separate filings. [Theoutpost](#) **Sanction: public reprimand and disqualification of all three attorneys from further case participation**, plus referral to the Alabama State Bar. [CBS 42](#) No monetary sanction — Judge Manasco declared monetary penalties are inadequate: “If fines and public embarrassment were effective deterrents, there would not be so many cases to cite. And in any event, fines do not account for the extreme dereliction of professional responsibility that fabricating citations reflects.” [EDRM](#) The opinion was ordered published in the Federal Supplement. [Naturalandartificiallaw](#) [Esquire Deposition Solutions](#) Butler Snow itself was not sanctioned because the firm “acted reasonably in its efforts to prevent this misconduct.” [Theoutpost](#) [Esquire Deposition Solutions](#) This case established **disqualification as a proportionate sanction** for AI fabrication. [Corporatcomplianceinsights](#)

ByoPlanet v. Johansson — the largest verified AI sanction

ByoPlanet International, LLC v. Johansson, 792 F. Supp. 3d 1341 (S.D. Fla. 2025), Case No. 0:25-cv-60630, decided July 17, 2025 (fee order August 1, 2025), by Judge David S. Leibowitz. Attorney James Martin Paul used AI-generated hallucinated citations across **at least eight related cases** (four federal, multiple state court). Despite being put on notice, Paul continued submitting fabricated filings — including in a response to a show cause order. [Relativity](#) [VinciWorks](#) **Sanction: approximately \$85,568** in attorneys’ fees, dismissal of all four federal cases, and referral to the state bar. [VinciWorks](#) Judge Leibowitz opened his sanctions order with a fabricated Scalia quote generated by his own ChatGPT query to illustrate the point. [Relativity](#) At the time, this was the **largest AI-hallucination sanction on record**. [VinciWorks](#)

Additional notable 2025 sanctions

Goldberg Segalla / Chicago Housing Authority (Cook County Circuit Court, Ill., Dec. 2025, Judge Thomas Cushing). Attorney Danielle Malaty used ChatGPT to draft a post-trial motion in a \$24 million lead-paint verdict, citing the nonexistent *Mack v. Anderson*. Malaty was fired. [The Real Deal](#) Plaintiffs identified 14 fabricated or misrepresented citations. [Darin Klemchuk](#) **Sanction: approximately \$60,000** (~\$50,000 against firm, ~\$10,000 against individual attorney). [Segalmccambridge](#)

Noland v. Land of the Free, L.P. (Cal. Ct. App., 2d Dist., Sept. 2025). Attorney Amir Mostafavi filed a brief in which **21 of 23 quotes were fabricated by ChatGPT**. [CalMatters](#) **Sanction: \$10,000** plus referral to the California State Bar. [LawSites](#) Notably, the court declined to award fees to opposing counsel partly because opposing counsel failed to detect the fabrications — potentially the first case addressing opposing counsel’s duty regarding AI hallucinations. [LawSites](#)

Mavy v. Commissioner of Social Security Administration (D. Ariz., Aug. 14, 2025, Judge Alison Bachus). **12 of 19 cases cited were fabricated, misleading, or unsupported.** Non-monetary sanctions including mandatory notification to judges in all pending cases.

United States v. Farris (6th Cir. 2025/2026). Court suspicions triggered by the file name "CoCounsel Skill Results" on uploaded documents. Three citations contained fabricated quotations. **Sanctions: denial of CJA compensation, removal from representation, referral for disciplinary proceedings.** [Speaker Law Firm](#)

Oregon Court of Appeals / Ghiorso (Dec. 2025). Established a **per-item sanction formula**: \$500 per fake citation, \$1,000 per false quotation. Initial calculation: \$16,500, **capped at \$10,000.** [Gizmodo](#)

Southern District of Ohio (2025, Senior Judge Walter H. Rice). **\$7,500 collective sanction** against two attorneys; found in contempt; referred to Ohio Supreme Court's Office of Disciplinary Counsel for "most egregious violations of Rule 11." [ABA Journal](#)

2026 cases: Sanctions continue climbing

Whiting v. City of Athens, Tennessee, Nos. 24-5918/5919/25-5424 (6th Cir., March 13, 2026), Judges Stranch, Bush (author), and Murphy. Attorneys filed briefs containing **more than two dozen fake citations.** [PlatinumIDS Blog](#) [Reason.com](#) When asked whether they used AI, the attorneys refused to answer and accused the court of a "vast conspiracy."

[Local3News.com](#) [Ediscoverytoday](#) **Sanction: \$15,000 each in punitive fines (\$30,000 total),** [PlatinumIDS Blog](#) plus reimbursement of all appellees' reasonable fees, double costs, and referral for disciplinary proceedings. [Noah News](#) Both attorneys had prior discipline for lack of candor. [PlatinumIDS Blog](#) [LawSites](#)

Rivera v. Triad Properties Corp. (N.D. Ala., ~April 2026, Judge Anna M. Manasco — the same judge from *Johnson v. Dunn*). Attorney Joshua Watkins used AI to produce "both misleading and outright fabricated statements of law." The firm was also sanctioned for "apparent lack of internal controls and guardrails surrounding its attorneys' use of artificial intelligence — indeed, the very AI the firm pays for and encourages its attorneys to use." [Itbehere](#) **Sanction: approximately \$47,057 in total,** [Itbehere](#) plus public reprimand of attorney and firm. [Reason.com](#)

The ~\$96,000 sanctions case

! Verification caveat: Research identified a case reportedly involving attorney Stephen Brigandi in the District of Oregon, with an opinion filed approximately April 4, 2026, by U.S. Magistrate Judge Mark Clarke. The reported sanction is **\$96,000** (\$15,500 in disciplinary penalties plus \$80,500 in opposing counsel's fees), with co-counsel liable for an additional \$14,200. Three filings allegedly contained **23 fabricated citations and 8 false quotations.**

The judge reportedly found “persuasive evidence” that the client, not the attorney, drafted the briefs using AI. [The Ethics Reporter](#) **This case was primarily sourced from secondary legal blogs and could not be independently verified through primary court documents, established legal databases, or major legal news outlets (ABA Journal, Bloomberg Law, Reuters).** It should be independently verified before inclusion in a published practitioner’s guide.

Part 2: Does AI destroy privilege? Two federal courts disagree

On **February 10, 2026**, two federal courts issued the first rulings on whether using AI tools affects attorney-client privilege and work product protection — and reached opposite conclusions. [Proskauer Rose LLP](#) These cases are essential reading for every practitioner using AI.

United States v. Heppner — privilege waived

United States v. Heppner, No. 1:25-cr-00503-JSR, 2026 WL 436479 (S.D.N.Y. Feb. 17, 2026), Judge Jed S. Rakoff. ⚠️ *Note: This case was referenced in the research request as “Heppner v. Turnberry Towers Corp” — the correct case name is United States v. Heppner, a criminal prosecution with no “Turnberry Towers” party.*

Bradley Heppner, a Dallas financial services executive indicted on securities fraud charges, independently used the **consumer (free-tier) version of Anthropic’s Claude** to generate approximately **31 documents** containing defense strategy analyses. [Gibson Dunn](#) He inputted information he had learned from Quinn Emanuel defense counsel into Claude.

[Chapman and Cutler LLP](#) When the FBI executed a search warrant, these AI-generated materials were seized from his devices. [Gibson Dunn](#)

Judge Rakoff ruled the documents were **protected by neither attorney-client privilege nor work product doctrine**, calling it “a question of first impression nationwide.” [Ogletree](#)

The reasoning addressed three key points.

First, **Claude is not an attorney**. All recognized privileges require “a trusting human relationship” with “a licensed professional who owes fiduciary duties and is subject to discipline.” [Gibson Dunn](#) Because “no such relationship exists, or could exist, between an AI user and a platform such as Claude,” that alone disposed of the privilege claim. [Paul, Weiss](#)

Second, **there was no reasonable expectation of confidentiality**. The court analyzed **Anthropic’s consumer Claude privacy policy** (dated February 19, 2025), finding it: collects data on users’ “inputs” and Claude’s “outputs”; uses that data to “train” its model; and “reserves the right to disclose such data to a host of ‘third parties,’ including

'governmental regulatory authorities'" — even absent a subpoena. [Debevoisedatablog](#) The court concluded: **"The policy clearly puts Claude's users on notice"** that data may be disclosed "in connection with claims, disputes, or litigation." [Paul, Weiss](#)

Third, the documents were **not prepared at counsel's direction**. Heppner acted "of his own volition." [Brooks Pierce](#) Claude's terms of service expressly disclaim providing legal advice. [O'Melveny](#)

Critically, Judge Rakoff left open a narrow path: **had counsel directed Heppner to use Claude**, "Claude might arguably be said to have functioned in a manner akin to a highly trained professional" under the *Kovel* doctrine (*United States v. Kovel*, 296 F.2d 918 (2d Cir. 1961)). [O'Melveny](#) The distinction between **consumer and enterprise AI platforms** was therefore implicitly central to the ruling.

Warner v. Gilbarco, Inc. — work product protected

Warner v. Gilbarco, Inc., No. 2:24-cv-12333, 2026 WL 373043 (E.D. Mich. Feb. 10, 2026), Magistrate Judge Anthony P. Patti. [Proskauer Rose LLP](#) ⚠️ *Note: This case was referenced as "Warner v. Flanders" — the correct name is Warner v. Gilbarco, Inc.*

Pro se plaintiff Sohyon Warner used **ChatGPT** as a "drafting interface" to prepare litigation filings in an employment discrimination case. Defendants moved to compel production of all AI-related materials, arguing AI use waived privilege and work product protections.

[Proskauer Rose LLP](#)

[Fisher Phillips](#)

Judge Patti **denied the motion** and upheld work product protection. [Proskauer Rose LLP](#) The critical holding: **"[Generative AI programs] are tools, not persons, even if they may have administrators somewhere in the background."** [Paul, Weiss](#) [Proskauer Rose LLP](#) The court drew a sharp line between privilege waiver standards and work product waiver standards. While voluntary disclosure to a third person "will generally suffice to show waiver of the attorney-client privilege, it should not suffice in itself for waiver of the work product privilege." [Paul, Weiss](#) Work product waiver requires **disclosure to an adversary** or in a manner likely to reach an adversary — and using ChatGPT does not meet that standard.

[Jones Walker LLP](#)

[Herbert Smith Freehills Kramer](#)

Judge Patti warned that accepting the defendants' theory "would nullify work-product protection in nearly every modern drafting environment, a result no court has endorsed."

[Minerva26](#)

Unlike *Heppner*, the court **did not analyze ChatGPT's privacy policy or terms of service**, finding platform terms irrelevant to work product analysis. [Proskauer Rose LLP](#)

The court also characterized the defendants' pursuit as a "fishing expedition" and warned them that their "preoccupation with Plaintiff's use of AI needs to abate." [Hrlegalist](#)

Morgan v. V2X, Inc. — a third ruling sides with Warner

Morgan v. V2X, Inc., No. 25-cv-01991-SKC-MDB, 2026 WL 864223 (D. Colo. Mar. 30, 2026), Magistrate Judge Maritza Dominguez Braswell. This ruling sided with *Warner*, holding that a pro se litigant's use of AI to prepare for litigation is protected work product. [SERGENIAN LAW](#) However, the court ordered disclosure of the **name** of the AI tool used (finding this does not reveal mental impressions) and amended the protective order to **bar inputting confidential information into AI platforms** unless the provider contractually prohibits training on inputs, prohibits third-party disclosure, and allows deletion upon request. [Herbert Smith Freehills Kramer](#) [NatLawReview](#)

What these cases mean for practitioners

The tension between *Heppner* and *Warner* creates an immediate practice imperative. The key variables are: (1) whether privilege or work product is at issue (different waiver standards apply); (2) whether the AI platform is consumer or enterprise tier (privacy policies matter for privilege); (3) whether AI use was at counsel's direction (the *Kovel* pathway); and (4) jurisdiction. [NatLawReview](#) [Bond, Schoeneck & King](#) **Until circuit courts resolve this split, the safest approach is using enterprise-tier AI platforms with contractual no-training guarantees and treating all AI interactions as potentially discoverable.**

Part 3: Consumer versus enterprise AI — the data practices that determine privilege

The *Heppner* court's reliance on Anthropic's consumer privacy policy makes platform data practices a **privilege-determinative factor**. The differences between consumer and enterprise tiers are stark and directly relevant to the duty of confidentiality under Model Rule 1.6.

OpenAI/ChatGPT

Consumer tiers (Free, Plus, Pro) train on user data by default. Users can opt out via Settings → Data Controls → "Improve the model for everyone," [OpenAI](#) but the opt-out is **not retroactive** — prior data may already be in training pipelines. [OpenAI Help Center](#) Content may be reviewed by human reviewers. [Nightfall AI](#)

Business and Enterprise tiers contractually prohibit training: [ChatGPT](#) "We do not use data from ChatGPT Enterprise, ChatGPT Business, ChatGPT Edu, or our API platform — including inputs or outputs — for training or improving our models." [OpenAI](#) [IntuitionLabs](#) Enterprise features include AES-256 encryption, customer-controlled encryption keys, custom data retention, and SOC 2 compliance. [OpenAI](#) [Gradually AI](#)

OpenAI API has not trained on user data since March 1, 2023. [Openai](#) Zero Data Retention

(ZDR) is available for qualifying enterprise customers — data processed in-memory only.

Medium

Anthropic/Claude

A critical policy shift occurred in **October 2025**. Anthropic moved from “never trains on consumer data” to a mandatory opt-in system. [Protecto AI](#) [Bitdefender](#) **Consumer users (Free, Pro, Max)** who opted in have data retained in de-identified form for **up to 5 years** — a 60× increase from the prior 30-day policy. Users who opted out retain the 30-day policy.

[Fello AI](#) [Char](#) Safety-flagged conversations may be used regardless of settings.

[claude +2](#) **This is the policy Judge Rakoff analyzed in *Heppner*.**

Claude for Work (Business/Enterprise) and **Anthropic API** are contractually excluded from training. Commercial customers are governed by separate Commercial Terms. [Char](#) Critically, **Claude Pro is classified as a consumer product** — paying \$20/month does not purchase enterprise-level protections. Only Claude for Work, API, or cloud marketplace deployments (Amazon Bedrock, Google Vertex AI) provide contractual no-training guarantees. [AMST Legal](#)

Google Gemini

Consumer Gemini (including paid Gemini Advanced) trains on user data by default via “Keep Activity” settings. [Gayaone](#) Disabling this opt-out removes chat history entirely — users cannot keep history and opt out simultaneously. [Tom's Guide](#) [Drainpipe](#) **Paying for Google One AI Premium does not change training policies.** Google Workspace Gemini explicitly states: “Workspace does not use customer data for training models without customer’s prior permission.” [Google Support](#) The Gemini API does not train on user data (55-day abuse monitoring retention). [Char](#)

Microsoft Copilot

Microsoft 365 Copilot (Enterprise) provides the strongest contractual protections: “Prompts, responses, and data accessed through Microsoft Graph aren’t used to train foundation LLMs.” [Metomic](#) [Microsoft Learn](#) Data is not shared with OpenAI despite using OpenAI models. [Microsoft Learn](#) Covered under Microsoft’s Data Protection Addendum. [Metomic](#) [Microsoft Learn](#) **Consumer Copilot** data can be used for training with user consent. [buckleyPLANET](#)

The “paid ≠ private” trap

The single most important finding for legal practitioners: **paying for an individual AI subscription does not provide enterprise-level data protections.** ChatGPT Plus, Claude Pro, and Gemini Advanced are all classified as consumer products that train on user data by

default (or require active opt-out). [Drainpipe](#) Only enterprise tiers — ChatGPT Business/Enterprise, Claude for Work, Workspace Gemini, Microsoft 365 Copilot — offer contractual no-training guarantees suitable for confidential legal work.

Part 4: The regulatory and ethics landscape reshaping AI in law

Over 300 judges now have AI standing orders

The claim of 300+ judicial standing orders is **verified by multiple sources**, including the Ropes & Gray tracker (the most-cited primary source), Bloomberg Law, and the National Law Review. The trajectory: from **1 order** (Judge Brantley Starr, N.D. Texas, May 30, 2023) to approximately 36 by mid-2024 to **over 300 by late 2025**. [Hintyr](#) Approaches vary: some require disclosure of AI use, some require disclosure plus certification that citations were verified, and some impose outright bans on generative AI in filings. [NatLawReview](#)

[OBLIC](#) In Ohio specifically, **three federal judges** in the Southern District have entirely banned AI use, [American Bar Association](#) and Judge Christopher A. Boyko (N.D. Ohio) issued a standing order banning AI with sanctions ranging from striking pleadings to dismissal.

[Uscourts](#)

Illinois says disclosure should not be required

The **Illinois Supreme Court** took a contrary position effective **January 1, 2025**, following the IJC AI Task Force report. [HeplerBroom LLC](#) The policy explicitly states: **"Disclosure of AI use should not be required in a pleading"** and that AI use "should not be discouraged, and is authorized provided it complies with legal and ethical standards." [2Civility](#) Illinois treats AI as "just another tool" governed by existing ethical rules rather than requiring special AI-specific provisions. [Eve](#) [Darrow](#) This approach is **unique among states** and directly contrasts with the 300+ judges who have imposed disclosure requirements. However, tension exists within Illinois: four Cook County state judges and multiple Illinois federal judges have independently imposed their own AI disclosure mandates despite the Supreme Court's position.

DOJ AI Litigation Task Force

President Trump signed **Executive Order 14365** on December 11, 2025, directing creation of an AI Litigation Task Force. [Lexology](#) Attorney General Pam Bondi formally announced it on **January 9, 2026**. [Lexology](#) The Task Force's mandate is to challenge state AI laws inconsistent with federal policy of maintaining "a minimally burdensome national policy framework for AI" — primarily on Dormant Commerce Clause and preemption grounds.

[FinancialContent](#) Expected targets include the Colorado AI Act, New York's RAISE Act, and

California's SB 53. A bipartisan coalition of **36 state attorneys general** has opposed a federal moratorium on state AI laws. [Consumer Financial Services Law M](#)

Proposed Federal Rule of Evidence 707

The Advisory Committee on Evidence Rules proposed **Rule 707** ("Machine-Generated Evidence") in May 2025, approved for public comment in an 8-to-1 vote on June 10, 2025. The public comment period closed **February 16, 2026**; the rule is **not yet enacted**. The proposed text: "When machine-generated evidence is offered without an expert witness and would be subject to Rule 702 if testified to by a witness, the court may admit the evidence only if it satisfies the requirements of Rule 702(a)-(d)." [GT TechVenture Views](#) This would prevent proponents from evading reliability requirements by offering AI output directly without a human expert. A related proposed amendment to Rule 901(c) addressing deepfakes was tabled. [Drug & Device Law](#)

ABA Formal Opinion 512 — the ethical framework

The ABA Standing Committee on Ethics and Professional Responsibility issued **Formal Opinion 512** ("Generative Artificial Intelligence Tools") [American Bar Association](#) on **July 29, 2024** [American Bar Association](#) — the ABA's first formal opinion on generative AI. [Darrow](#) [Washington University School of La](#) It establishes duties across six areas: competence (Rule 1.1 — must understand AI capabilities and limitations), [American Bar Association](#) confidentiality (Rule 1.6 — must review vendor data-handling practices), communication (Rule 1.4 — may need to disclose AI use to clients), candor to tribunal (Rules 3.1, 3.3 — must verify all citations), supervision (Rules 5.1, 5.3 — must supervise AI-generated work product), and fees (Rule 1.5 — cannot bill for time saved by AI, cannot charge clients for learning general AI competence). [Steno](#) [Unc](#)

The ABA Task Force on Law and Artificial Intelligence issued its **Year 2 Report** on December 15, 2025, [American Bar Association](#) finding AI has moved "from experiment to infrastructure" and that "AI adoption has surpassed understanding." [Above the Law](#) The report notes: "As of February 2025, no known GenAI tools have fully resolved the hallucination problem."

[Dirittoue](#)

Over 35 states have issued AI guidance

Key formal ethics opinions include: Florida Opinion 24-1 (January 2024), [Smcba](#) Pennsylvania Joint Formal Opinion 2024-200 (2024), North Carolina 2024 Formal Ethics Opinion 1 (2024), [Spellbook](#) Texas Ethics Opinion No. 705 (February 2025), [Steno](#) Oregon Formal Opinion No. 2025-205 (2025), [Justia](#) and New York Formal Opinion 2025-6 (2025). [Spellbook](#) Common requirements across jurisdictions: attorneys must understand AI before use, must not input confidential data into unsecured tools, must verify all AI-

generated output, may need to disclose AI use to clients, and cannot inflate billing for AI-saved time. [Texas Bar Blog](#) Louisiana is the only state with a **statute** (Act No. 250, effective August 2025) addressing AI-generated evidence. [Hintyr](#)

Ohio's current posture — resources without mandates

Ohio has **not issued a formal bar ethics opinion** on attorney AI use. [Steno](#) The Supreme Court of Ohio maintains an "Artificial Intelligence Resource Library" with guidance keyed to Ohio Rules of Judicial Conduct, but explicitly states it is "in no way endorsing the use of AI." [Ohio Supreme Court](#) The Ohio Bar Liability Insurance Company (OBLIC) published guidance in October 2024 covering ABA guidance and practical recommendations. An article by D. Allan Asbury in *Columbus Bar Lawyers Quarterly* (April 2025) discusses AI in the context of Ohio Rules of Professional Conduct. Notably, a **Southern District of Ohio case** before Senior Judge Walter H. Rice resulted in \$7,500 in collective sanctions and referral to the Ohio Supreme Court's Office of Disciplinary Counsel [ABA Journal](#) — demonstrating that Ohio federal courts are actively enforcing AI-related duties.

Part 5: The Stanford study proves legal AI tools still hallucinate

The Stanford RegLab study is now peer-reviewed and published: **Magesh, V., Surani, F., Dahl, M., Suzgun, M., Manning, C.D., and Ho, D.E. (2025). "Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools." *Journal of Empirical Legal Studies*, [LegalAIWorld](#) 22: 216–242. DOI: 10.1111/jels.12413.**

The study — the first preregistered empirical evaluation of AI-driven legal research tools [Stanford](#) [Stanford Law School](#) — tested over 200 legal queries across Lexis+ AI, Ask Practical Law AI, and Westlaw AI-Assisted Research in May 2024. **Lexis+ AI hallucinated at a rate exceeding 17%** [LegalAIWorld](#) (with 65% accurate responses and 18% incomplete/refused). **Westlaw AI-Assisted Research hallucinated at approximately 33%** — nearly twice [LegalAIWorld](#) Lexis+ AI's rate. Ask Practical Law AI hallucinated at >17% but refused to answer 62% of queries, producing accurate responses only 18% of the time. For comparison, **GPT-4 hallucinated at approximately 43%**.

The study identified two hallucination types: "incorrect" (describes law incorrectly) and "misgrounded" (describes law correctly but cites sources that do not support the claim). Thomson Reuters initially denied Stanford researchers access to Westlaw AI-Assisted Research three separate times; after the preprint was published and generated public pressure, Thomson Reuters provided access. The updated analysis confirmed the higher hallucination rate. A predecessor study by Dahl et al. (2024) found general-purpose LLM hallucination rates of **58–88%** on legal questions, and a separate finding by Xu et al.

(2024) formally proved that eliminating hallucination in LLMs is **mathematically impossible** given their fundamental architecture.

Part 6: Building an AI governance policy — available templates and key components

Published model policies

The **Virginia Bar Association** produced the most comprehensive standalone “Model AI Policy for Law Firms” (Version 1.0, May 2024), directly designed as an adoptable template covering training, confidentiality, validation, and AI in decision-making. Additional published templates include: AllRize’s comprehensive model policy (covering governance, risk assessment, standards of care, training, and compliance); Jaffe’s fill-in-the-blank template (with data tiers for Public/Confidential/Highly Confidential and vendor due diligence requirements for SOC 2/ISO 27001); Fisher & Phillips LLP’s published sample “Acceptable Use of Generative AI Tools” policy; Clio’s law firm AI policy template; and resources from the American Inns of Court and the Illinois ARDC.

Essential policy components

Based on synthesis of all available templates and guidance, a comprehensive law firm AI policy should address fifteen areas:

- **Purpose, scope, and definitions** (AI, GenAI, enterprise AI, covered tools, covered personnel)
- **Approved versus prohibited tools** (whitelist approach; prohibition on consumer AI for client data)
- **Data classification** (Public/Confidential/Highly Confidential tiers; prohibition on inputting PII/PHI into unapproved tools)
- **Human oversight and verification** (mandatory attorney review; citation verification before filing)
- **Ethical compliance** (mapping to ABA Model Rules 1.1, 1.4, 1.5, 1.6, 3.3, 5.1, 5.3; state-specific requirements)
- **Vendor due diligence** (security certifications; training policies; contractual protections)
- **Client communication and disclosure** (engagement letter updates; court disclosure requirements)

- **Billing practices** (cannot bill hourly for AI-saved time; cost of AI tools as overhead)
- **Governance structure** (AI committee; designated AI officers; incident reporting)
- **Training and competence** (mandatory onboarding; ongoing CLE)
- **Monitoring and audit logging** (prompt/output logs; quarterly reviews)
- **Incident response** (error reporting; breach protocols)
- **Regular review cycle** (at least annual policy updates)
- **Employee acknowledgment** (signed understanding of policy)

The adoption gap remains significant: according to the 2025 Clio Legal Trends Report, **79% of legal professionals utilized AI tools** but **44% of law firms had not implemented formal governance policies**.

Conclusion

The arc of AI ethics enforcement in legal practice bends sharply toward greater accountability. Three developments define this moment. First, **sanctions have evolved from symbolic fines to career-altering consequences** — *Johnson v. Dunn* established disqualification as proportionate, *ByoPlanet* pushed monetary sanctions past \$85,000, and *Whiting* showed appellate courts will impose punitive fines plus fee-shifting. The pattern is clear: courts treat cover-ups and refusal to take responsibility far more harshly than initial errors.

Second, the *Heppner/Warner* split on AI and privilege creates **immediate, jurisdiction-dependent risk for every attorney using consumer AI tools**. The safest reading of current law: consumer AI platforms may destroy attorney-client privilege (per *Heppner*) even if work product protection survives (per *Warner*). Enterprise platforms with contractual no-training commitments are the only defensible choice for confidential legal work, and the *Morgan v. V2X* protective order points toward courts mandating this distinction.

Third, the regulatory patchwork — 300+ judicial standing orders, 35+ state bar opinions, proposed FRE 707, the DOJ Task Force, and Illinois's dissent — has created a compliance landscape where no single approach satisfies all jurisdictions. Ohio practitioners face a particular challenge: the state has no formal ethics opinion, but Ohio federal courts have imposed some of the strictest responses (outright AI bans, \$7,500 sanctions with disciplinary referrals). The absence of state-level guidance does not mean absence of liability — it means practitioners must self-govern using the ABA framework, their jurisdiction's standing orders, and the rapidly developing case law compiled in this report.

The fundamental lesson from every sanctions case since *Mata*: AI is a permissible tool, but the attorney remains the guarantor of accuracy. No technology shifts that burden.