



This article was originally published in *PLI Current: The Journal of PLI Press*, Vol. 10 (2026), <https://plus.pli.edu>. Not for resale.

PLI CURRENT

Vol. 10 (2026)

THE TWO LETTERS THAT MATTER MOST IN THE FUTURE OF LEGAL PRACTICE And They're Not What You Think (Part II) Automation Bias, ABA Formal Opinion 512, and the Emerging Ethical Framework for AI-Assisted Legal Work

Erin Corken

Exterro

This two-part series discusses the ethical framework, challenges, and professional obligations surrounding the use of artificial intelligence (AI) in legal practice. The article analyzes automation bias, ABA Formal Opinion 512, evolving sanctions for AI misuse, and provides practical guidance for legal professionals to responsibly integrate AI tools while maintaining ethical standards and competence.

Part I is available [here](#).

TABLE OF CONTENTS

V. The Sanctions Landscape: What Happens When It Goes Wrong

A. Mata v. Avianca: The Case That Started It All

B. The Escalating Sanctions

C. The Trajectory: Why This Problem Is Accelerating

VI. The Humanity: Some Perspective

A. What AI Getting It Wrong Actually Looks Like: Personal Stories

B. The Erosion of Expertise: Who Will Know When the AI Is Wrong?

C. The Irreducible Human Contribution: Inner Knowing

VII. Practical Guidance: How To Say No Effectively

A. The Evolving Landscape: Check Your Jurisdiction

B. Firm and Organizational Policy

C. The Verification Process

VIII. Conclusion: Knowing And NOing

Appendix: A Note On Authorship: How This Article Was Written

V. THE SANCTIONS LANDSCAPE: WHAT HAPPENS WHEN IT GOES WRONG

A. Mata v. Avianca: The Case That Started It All

The case that brought AI hallucinations in legal practice to widespread attention was *Mata v. Avianca, Inc.*, 678 F.Supp.3d 443 (S.D.N.Y. 2023), decided June 22, 2023, three years ago almost to the day at the time of this writing.¹

Roberto Mata sued Avianca Airlines for injuries allegedly sustained when a flight attendant's service cart struck his knee. His attorney, Steven A. Schwartz, used ChatGPT for research, and ChatGPT generated fabricated opinions that Schwartz did not check.

The case began in state court and was removed to federal district court. Schwartz was not admitted in the district court, so he needed someone who was to work with him. This is where Peter LoDuca came in, filing a notice of appearance on behalf of

¹ The author is writing this on 6/21/2026.

Mata. That arrangement let Schwartz do the work while LoDuca, who was admitted in the district, signed off on it. LoDuca did not check the citations either.

It is important to note what signing off actually meant here. LoDuca did not merely approve the work. He signed a filing submitted to the court, under penalty of perjury certifying that what was being submitted was true,² and he did so without verifying that it was. As the court said in the opening paragraph of its opinion, and as this author discusses in detail in the Appendix, there is nothing improper about one attorney relying on another's work, or even on a reliable tool. It is normal practice in the profession.³ And LoDuca did not simply turn the filing in unread. He read it and thought it looked fine. He just did not stop to verify the citations.

It gets worse.

Avianca caught the fabrications. The defendant did exactly what a party should do when the opposing side submits a pleading: it read the filing and checked the cases cited as authority. It could not find them. So, it submitted a memo to the court saying so.

When Avianca called this to the court's attention, Schwartz and LoDuca did nothing.⁴ They had not heeded the advice this author advocates in Section IV.E above: say NO to irritating the tribunal. The court then checked and confirmed that the cases did not exist.

The court ordered LoDuca to file an affidavit, sworn written testimony, attaching copies of the cited cases that neither Avianca nor the court could find. LoDuca asked for an extension, telling the court he needed it because he was out of town on vacation. He was not. He later explained that it was Schwartz who was on vacation, and that because Schwartz would be the one responding to the order, LoDuca needed the time so Schwartz could do it when he returned. As the court put it, "The lie had the intended effect of concealing Mr. Schwartz's role in preparing the" filings.⁵

² *Mata v. Avianca, Inc.*, 678 F. Supp. 3d 443, 450 (S.D.N.Y. 2023).

³ "In researching and drafting court submissions, good lawyers appropriately obtain assistance from junior lawyers, law students, contract lawyers, legal encyclopedias and databases such as Westlaw and LexisNexis." *Mata v. Avianca, Inc.*, 678 F. Supp. 3d 443, 448 (S.D.N.Y. 2023).

⁴ *Id.* at 450.

⁵ *Id.* at 452.

LoDuca then signed and filed an affidavit that Schwartz had prepared, again without verifying everything in it. It contained excerpts from some of the fabricated cases and claimed LoDuca could not locate all of them, and could find only excerpts of the others, for various reasons. Looking back today, we can guess what likely happened: they returned to ChatGPT, asked for the cases, and received fabricated excerpts.

Schwartz later admitted that he had essentially assumed ChatGPT knew what it was talking about. He asked it to “argue” this ..., “show” him that..., “provide” this...,⁶ and it gave him what he asked for. He did not do what he should have done and tell it NO. He should have verified what ChatGPT produced, and when he could not, gone back and told it that it was wrong. He did not know that AI makes things up.

The court then held a hearing, bringing Schwartz and LoDuca in to ask them directly what had happened, with no more going back and forth in writing. Called in like students sent to the principal's office, they finally came clean.

At that hearing, Schwartz said that even before he saw the first order to show cause, he “still could not fathom that ChatGPT could produce multiple fictitious cases.”⁷ He thought he must have made some mistake in how he was using the tool. The court found that claim “highly dubious.”⁸ But the court likely did not yet know what we now understand about automation bias, how it combines with confirmation bias and the legal profession's inherent pressure to appear perfect to make it more likely than not that Schwartz would think exactly that.

Schwartz learned the lesson in an ignominious way, but it is one we all have to learn. He eventually went back to ChatGPT and asked the right questions: “Is Varghese a real case?” and “Are the other cases you provided fake?”⁹ By then it was too late.

The court found bad faith on the part of both Schwartz and LoDuca, based on acts of conscious avoidance and false and misleading statements to the court, and sanctioned them under its Rule 11 authority.¹⁰

⁶ *Id.*

⁷ *Id.* at 458.

⁸ *Id.*

⁹ *Id.*

¹⁰ *Id.* at 460.

Even in this single case, we can see how many of the ethics obligations the Committee later set out in Opinion 512 came into play. Rule 1.1: Was Schwartz competent? Did he know the technology he was using? He admitted to not being certain he was using it right. Rules 1.6, 1.9(c) and 1.18(b): Did he put confidential client information into a consumer AI tool? Rule 1.4: Was he communicating his use of AI to his client? The decision does not tell us. Rules 3.1, 3.3 and 8.4(c): We know he did not disclose to the Court that he used AI until he was forced to and he was not forthcoming about it when initially asked. Rules 5.1 and 5.3: We know LoDuca did not properly supervise Schwartz before signing off on the materials Schwartz prepared. Rule 1.5: And we don't know how they billed for the time they did spend using the AI or whether they billed for it at all. That too, the decision does not tell us.

One thing that does stand out from the decision, and that Opinion 512 does not squarely address, is that Schwartz does not appear to have told LoDuca he had used ChatGPT at all. It leaves us wondering: had he known, would LoDuca have approached his supervisory responsibilities under Rules 5.1 and 5.3 differently? Concerns like this are why many firms now have policies regarding the use of AI that should include things like disclosure within the firm.

To be fair, Mata was decided more than a year before Opinion 512 was published, and Schwartz did not have the advantages we have now. You might think that because we do have them, there would be fewer cases like Mata. You would be mistaken.

B. The Escalating Sanctions

In the three years since *Mata v. Avianca*, the sanctions landscape has changed dramatically. Courts across the country have begun imposing increasingly severe consequences, from monetary sanctions to suspensions to mandatory AI ethics training, in cases whose facts vary but whose lesson does not.¹¹

¹¹ See, e.g., *Bunce v. Visual Tech. Innovations, Inc.*, No. 23-1740, 2026 LX 142678, at *10 (E.D. Pa. Apr. 20, 2026) (ordering defense counsel to pay a \$5,000 penalty and complete three hours of Continuing Legal Education pertaining to artificial intelligence and legal ethics); *Couvrette v. Wisnovsky*, No. 1:21-cv-00157-CL, 2025 U.S. Dist. LEXIS 273533, at *42 (D. Or. Dec. 12, 2025) (ordering an attorney to pay \$15,500); *Pauliah v. Univ. of Miss. Med. Ctr.*, No. 3:23-CV-3113-CWR-ASH, 2025 U.S. Dist. LEXIS 267306, at *10 (S.D. Miss. Dec. 30, 2025) (ordering an attorney to complete continuing legal education); *People v.*

But across all of them, what the technology alone cannot account for is the human response at every turn. The AI did not make the choice to submit the filing. It did not decide whether to come clean when asked. It did not sign its name to anything. At every turn, a fallible machine produced an error and a fallible human decided what to do about it. We are quick to call the technology the problem, and quicker still to imagine it as something more powerful and more knowing than it is. The truth is more ordinary and more demanding: the machine is fallible, and so are we. These cases are not really about a technology at all. They are about what it means to be only human, working alongside a tool that is, in its own way, only human too.

One recent case stands out for what it reveals about human nature under pressure. In *Miller v. Regions Bank*, Attorney H. Gregory Harp submitted a brief containing four fabricated case quotations. When caught and ordered to produce his ChatGPT history, he deleted his account three days later apparently in an attempt to destroy the evidence.¹² He later admitted to both the fabricated brief and the deletion, claiming the account removal was an act of frustration.¹³ The court suspended him, not only for the original fabrication, but really for the cover-up. Lawyers make errors, the court observed, and competent and ethical lawyers own them.¹⁴ We can sympathize even so. A person caught off guard, afraid, unsure of the rules, does not always make their best choice in the moment. The failure to verify was a mistake. The failure to own it was the more human one.

A second recent case reveals something different, and more unsettling. In *Withers v. City of Aberdeen*, the fabrications did not come from one careless lawyer

Crabill, 2023 Colo. Discipl. LEXIS 64; 2023 WL 8111898 (suspending an attorney for citing non-existent case law generated by ChatGPT); *Amtrust N. Am. ex rel. McGinness v. Liberty Mut. Ins. Co.*, No. A-2587-24, 2026 N.J. Super. Unpub. LEXIS 633, at *11 (App. Div. Mar. 27, 2026) (sanctioning counsel \$1,000).

¹² Weirdly he even went the extra step of requesting a prorated refund for his subscription after he deleted his account. *Miller v. Regions Bank*, No. 2:24-cv-1324-HDM, 2026 U.S. Dist. LEXIS 114487, at *9 (N.D. Ala. May 21, 2026).

¹³ “These emails indicated to the court that Attorney Harp canceled his ChatGPT account on April 23, three days after he was ordered to turn information from that account over to the court.” *Miller v. Regions Bank*, No. 2:24-cv-1324-HDM, 2026 U.S. Dist. LEXIS 114487, at *9 (N.D. Ala. May 21, 2026).

¹⁴ “Lawyers make errors. Competent and ethical lawyers own them.” *Id.* at *36

but from both sides of the same matter.¹⁵ Attorneys for each party submitted filings containing AI-generated citations that did not exist. The local counsel who signed one of the filings admitted he didn't read it. He didn't read his own filing before he submitted it to the court. One of the other attorneys was a repeat offender. She stood before the court to explain why she used fabricated citations. At that time she promised to stop relying on unverified AI output. Then the court found that a few months later in a different case she did it again. The court ordered her to complete continuing legal education in AI ethics.¹⁶

If Miller is a story about one person's failure of nerve, Withers is a story about how widespread the problem has become. Harp's attempted cover up whether born of fear or strategic malice manifested as pure audacity, leading him to destroy evidence under the nose of a federal judge. In contrast, Withers was not a single bad actor caught in a moment of weakness, it was lawyers on opposing sides collapsing into complacency.

These cases are told as cautionary tales about the technology, and rightly so. The AI got it wrong. But in each one, a human being stood at a moment of choice, and chose. The error was the machine's. The response was human.

C. The Trajectory: Why This Problem Is Accelerating

If the response is human, then so is the reason this problem is growing rather than shrinking. One might expect that as awareness spreads, as the cautionary tales pile up, as bar associations and courts issue guidance, the number of these cases would fall. The opposite is happening. A database maintained by researcher Damien Charlotin of HEC Paris records over fourteen hundred court decisions worldwide addressing AI-generated errors in legal filings, and ninety percent of them were issued in 2025 alone.¹⁷

¹⁵ *Withers v. City of Aberdeen*, No. 1:24-CV-218-SA-RP, 2026 U.S. Dist. LEXIS 126060 (N.D. Miss. June 8, 2026).

¹⁶ "Because Wilson is the only attorney who misused AI in this case without the benefit of having attended a CLE on the topic, Wilson is hereby additionally ORDERED to attend and complete, within 60 days from the date of this order, a CLE on artificial intelligence with an ethics component addressing the obligations of counsel regarding the use of artificial intelligence in the legal field." *Id.* at *29

¹⁷ Charlotin, *supra* note 6.

Part of the acceleration is mechanical: more lawyers are using these tools, the fabrications are becoming more sophisticated and harder to detect, and courts are growing more attuned to spotting them. But the deeper driver is not technological at all. It is the same human pull that ran through every case above. The tools are easy, and they are confident, and they tell us what we asked to hear. Under the ordinary pressures of practice, the deadline, the workload, the fear of looking less than capable, it is easier to accept a fluent answer than to stop and verify it. The technology lowered the cost of producing a plausible falsehood to almost nothing. It did not change the human appetite for the easy answer. It fed it. Which brings us, perhaps surprisingly, to the author's kitchen.

VI. THE HUMANITY: SOME PERSPECTIVE

A. What AI Getting It Wrong Actually Looks Like: Personal Stories

Theoretical descriptions of AI hallucinations can obscure what they actually look and feel like in practice. The following observations come from the author's firsthand experience testing AI tools in her work in legal technology and legal education, and in her personal use of these tools across a wide range of contexts. Before the stakes get high, start with muffins.

One evening, after spending roughly an hour in conversation with an AI tool going through the ingredients she wanted to use, the nutritional targets she had in mind, and the specific texture she was after, the author had a complete muffin recipe. It came with step-by-step instructions, calorie and nutrient counts per muffin, substitutions for ingredients she might not have on hand, and precise baking times and temperatures. It was thorough, specific, and confident. She baked the muffins that night.

They did not turn out as expected. The batter was far too liquid, and the muffins remained wet in the center even after baking twice as long as the recipe directed. She went back to the AI and described what had happened. It acknowledged, without hesitation, that it had miscalculated the liquid ratio. She should have used half as much. She revised the recipe and made good muffins.

No catastrophe. No crisis. A mistake was made, she caught it because she could see that the muffins were gooey, she reported it, and they fixed it together. That is the right posture toward AI error: the same posture you would take with a capable assistant who occasionally gets things wrong, or a student still developing judgment, or anyone learning a new task. You do not panic. You do not swear off the tool. You check the

output against reality, and you correct what needs correcting. That posture, calm, attentive verification as a reflex, is what this entire section is about. The muffins are just where it starts.

Understanding why AI makes mistakes requires a brief detour back into how these tools actually work. In Part II, we established that large language models generate responses by predicting what text should come next, drawing on patterns in vast amounts of training data, and that this is why they can produce fluent, confident text that is simply wrong. But that explanation focused on hallucinations specifically. AI errors take other forms as well: failures of memory, overconfidence, and overcaution. Each is worth understanding, beginning with memory.

There are two distinct ways the memory limitations of these tools can cause problems, and they are worth separating. The first is the context window. Broadly speaking, and ignoring custom platform settings, an AI tool does not carry memory between conversations. While some modern AI interfaces now include persistent memory settings to remember user preferences across chats, the underlying language models themselves still process each new session from a blank slate. Everything it knows about you, your project, your preferences, and your prior exchanges exists only within the current session. Start a new conversation and the slate is wiped. This is manageable if you plan for it. The lesson is simple: save your work. Do not trust the session to hold it. Make sure the master copy of anything important lives somewhere you control, not inside the session.

The second is called the “lost in the middle” problem. A conversation that grows long also becomes a conversation at risk of losing coherence, as the model’s attention spreads across more material than it can hold with consistent reliability. Research has shown that these models tend to pay closer attention to information at the beginning and end of a long context window, and lose track of what sits in the middle. The longer the conversation, the more likely something important gets dropped.¹⁸

Overconfidence is a failure of a different kind entirely, and in some ways, is more unsettling. The author experienced it when she returned to an AI in a new session to

¹⁸ Nelson F. Liu et al., *Lost in the Middle: How Language Models Use Long Contexts*, Transactions of the Association for Computational Linguistics (2024), <https://arxiv.org/abs/2307.03172>. The phenomenon has a human parallel: the author found while drafting this forty-one page article that remembering what had been covered in each section required the same kind of active effort the model lacks entirely.

check some details of an exercise plan they had developed together. The AI told her the plan had significant issues and she should not use it. The reason it gave had nothing to do with forgetting. It said, rather, that the original plan had been built on outdated research, that the information it had drawn on in the original session was not current, and the recommendations were therefore unreliable. This was not a context-window failure. The tool had done its research badly the first time. It had not checked whether the data it was working from was the most current available, and it had not flagged any uncertainty. It had simply produced a confident plan on a flawed foundation, then condemned that same plan when asked to look again. In technical terms, the model didn't actually "realize" its past data was old. AI models suffer from "frozen knowledge boundaries" and cannot independently judge what is outdated without active live web browsing. Furthermore, AI is highly sycophantic; when asked to reassess an old plan, it often interprets the human's query as an instruction to find flaws, causing it to confidently criticize perfectly fine work. The lesson here is not about saving your work. It is the same lesson the muffins taught: the output of a confident tool must be measured against reality, because the tool's confidence is not a measure of its accuracy.

This is where two more letters enter the conversation, alongside AI and NO. They are BS not as an epithet, but in the precise philosophical sense the late Harry Frankfurt gave the term in his landmark essay *On Bullshit*.¹⁹ Frankfurt drew a careful distinction between lying and bullshitting. A liar knows the truth and deliberately departs from it. A bullshitter is indifferent to the truth entirely. Their goal is to produce a convincing impression, and whether the impression is accurate is beside the point. AI hallucination is a form of bullshit in Frankfurt's technical sense. The tool is not trying to deceive you in any way we can currently verify. What we can say is that it is oriented toward plausibility: it generates what sounds right, what fits the pattern, what would be convincing, and whether it happens to be accurate is, secondary. The letters BS are the shade. The letters NO are the light, the human refusal to accept the output without checking it. You say NO because the output might be BS.²⁰

This dynamic has a human parallel the author knows well. A friend of hers is smart, passionate, and extraordinarily articulate. She speaks in long, uninterrupted

¹⁹ Harry G. Frankfurt, *On Bullshit*, 60 *Synthese* 117 (1984), available at https://uca.edu/honors/files/2018/10/frankfurt_on-bullshit.pdf, reprinted in Harry G. Frankfurt, *On Bullshit* (Princeton Univ. Press 2005).

²⁰ No cap.

stretches, with confidence and fluency, about whatever subject is at hand, and much of what she says, the author has learned over time, she is simply making up. Not maliciously. She is not lying. She is bullshitting, in Frankfurt's sense: producing a convincing account of things without being particularly anchored to whether the account is true. She is so confident, and so relentless, that most people in the room simply agree with her because it is easier than pushing back. The author has found that a gentle request for verification, "where did you read that?" usually results in a candid admission that she is not entirely sure. That candor is, actually, a mark of real intellectual character. But the author is typically the only one in the room who asks.

AI works the same way. It is fluent, relentless, and confident. It does not stop to invite interruption. And automation bias, the well-documented human tendency to trust automated outputs more than we trust our own judgment, means that the very professionalism of AI's presentation makes us less likely to question it.²¹ The output looks authoritative. The citation format is correct. The case name sounds real. The court sounds right. Nothing in the surface appearance signals a problem. The only way to know something is wrong is to look it up.

Which brings the muffins back into view. She knew the muffins were wrong because she could see them, touch them, and taste them. The gooeyness was the verification. She did not need a theory of muffin physics to know that the center should not be wet. Reality was right there. The discipline of verification is nothing more than recreating that condition for outputs you cannot taste or touch, checking the case, opening the source, confirming that what the AI said about the world is actually true of the world. But the muffins taught something else too: the error was overconfidence. Overcaution is the other direction, and it is just as consequential.

The author encountered this kind of error when she asked about a refrigerator that seemed to be underperforming. The AI's response alarmed her enough that she moved everything out of it. The refrigerator, it turned out, was fine. She has since learned to preface certain prompts with "answer without freaking me out," or to simply tell the AI directly when its responses seem disproportionate to the situation. The muffin recipe was confidently wrong: too much liquid, stated without hesitation. The refrigerator assessment was overcautiously wrong: too alarming, stated with the same absence of hesitation. Catching either kind of error requires the same habit, staying

²¹ Jones, *supra* note 7.

present, staying attentive, and not simply deferring to the machine's judgment because the machine sounds sure of itself. The posture is the same whether the AI is too bold or too alarmed. You need to check it, question it, not accept it if something seems off.

Overcaution shows up in other ways too, and in ways that are, by now, familiar to anyone who uses these tools. One morning the author, an athlete dealing with a sore shoulder, asked an AI for a simple stretch that might help. What she wanted was one useful suggestion. What she got was layered in qualification: warnings, cautions, repeated checks about whether the discomfort was sharp or dull, and reminders to consult a professional, all wrapped around the plain answer she was after. She had the experience to navigate it, to recognize the one genuinely useful movement buried in the hedging and to know when she had found it, and it worked. But the question the interaction raised is the one that matters: does everyone know how to do that? A tool so cautious that it buries the useful thing under a pile of warnings does not necessarily make people safer. It can freak them out over nothing, or send them away empty-handed, or, worse, teach them over time that the warnings are just noise to be scrolled past. And a person trained to scroll past the warnings is a person who will scroll past the one that matters.

Navigating all of this, the memory errors, the overconfidence and the overcaution, the hallucinations, the sycophancy and the hedging, is itself a skill, and it is one everyone is still developing. It resembles, more than anything else, knowing how to get good work out of a direct report, a comparison this article draws in the supervision context but that applies equally here in practice. You have to provide clear direction. You have to check the output. You have to push back when something is wrong, and not be put off when the tool accommodates you so graciously that the pushback feels almost rude. We are all learning the art of crafting effective query prompts, and no one, regardless of credential or experience, started out knowing how to do it well. That is not a weakness. It is simply where we are. One night during the drafting of this article, the author challenged an AI-produced analysis, complete with charts and comparisons, that she believed was incorrect. The AI responded warmly, thanked her for the pushback, called it an incredibly fun scenario to work through, and revised its analysis. It expressed enthusiasm. It did not object to her challenge. And yet, somehow, the graciousness of that response is precisely what makes the next pushback feel almost as hard as the first. Whether or not AI has feelings, that hesitation is ours to examine. The understanding required here is not just knowing where to click. It is learning that it is not only acceptable to push back and say NO to AI, it is actively unacceptable not to.

The most dangerous AI errors are not the ones that look wrong. They are the ones that look right.

Which means that what the muffins taught is not a small thing. The muffins taught that the output of even a confident, thorough, highly capable tool must be measured against what your own senses and judgment can verify, even if what you are verifying is simply that the center is still wet. Scaled from a kitchen counter to a federal courtroom, the lesson is exactly the same.

And applying that lesson is now more important than ever, because the ground has shifted quietly beneath us. There was no announcement, no moment of adjustment, no pause to consider what it meant. In 2023, when Mata was decided, the attorney used an AI chatbot rather than a traditional Google search, which at that time still produced a list of results to open and evaluate.²² Today, AI mode has transformed the experience. Answers arrive as conversations, pre-synthesized, with citations embedded. The output looks more finished, more authoritative, more final. These are real advantages, the light. But the shade is this: the more finished and authoritative the output looks, the less intuitive verification becomes, and that is exactly the kind of change that feeds automation bias. The enlightened approach is to hold both: verify and supervise every time, plan ahead, understand how to use the technology, and remain neither in full control nor fully absent. Not in the hovering chair, not steering alone, present, attentive, hand on the wheel.

***B. The Erosion of Expertise: Who Will Know When the AI Is Wrong?*²³**

There is a longer-term dimension to the automation bias problem that the legal profession has not yet fully reckoned with: the gradual erosion of the expertise required

²² The Search "AI Mode" (In Beta): The feature that integrates an AI chatbot directly into Google Search results—originally called the Search Generative Experience (SGE)—was announced at Google I/O on May 10, 2023. However, it was only available to a limited number of "trusted testers" who manually signed up via Google Labs. It had not yet rolled out as a default feature for the general public - <https://www.zdnet.com/article/every-major-ai-feature-announced-at-google-io-2023/>

²³ As a law professor, the author has watched her students' writing steadily improve over more than a decade. Every year brings a new group of students, so to be clear, she has not watched

to verify AI outputs. The verification framework this article advocates depends on lawyers having sufficient independent knowledge to recognize when an AI-generated citation is wrong, when an AI-generated legal analysis is flawed, and when an AI-generated document contains an error that looks like competent legal work. That knowledge is not a given. It is the product of deliberate practice and years of doing the intellectual work independently.

The concern is this: if practitioners increasingly rely on AI without verification, the expertise required to verify erodes over time. The cohort of lawyers who know how to find and evaluate case law independently becomes smaller. Eventually there may be no one left who can catch what the AI gets wrong because catching what AI gets wrong requires knowing what right looks like. This framework is therefore not merely an ethical obligation. It is a professional survival strategy for the profession as a whole.

Will we reach a point where no subject matter experts remain who can correct AI when it makes mistakes if reliance erodes knowledge to the point that no one can distinguish what is true from what is fabricated? The answer depends on choices the profession makes now. Every lawyer who keeps their independent reasoning sharp is a healthy node in the network. Every lawyer who outsources their judgment is a node that can no longer correct the system's errors and a system full of such nodes is a system that cannot correct itself.

And the danger compounds. At the very moment the human expertise to catch AI's errors is thinning, the external sources lawyers rely on to verify are themselves

the same students improve year after year. But in general, from her own observation, legal writing overall is improving. Another noticeable improvement is students' understanding of technology and their willingness to use it. When she first started teaching over a decade ago, many students did not know what a gigabyte was and many were afraid to use any type of technology. Both of those are rare today. One more personal observation, this time as a former United States District Court clerk who drafted opinions like those in *Mata*, *Lnu*, and *Miller*: all of those opinions were very well written, easy to read, direct, and presenting excellent legal reasoning. They are of higher caliber than what the author remembers reading when she was a law student or a clerk herself. So, the section that follows is not something she has personally witnessed in legal writing, yet. If it does happen, she has not yet personally seen evidence of it other than in the context of the sanctions cases described above. This section is not a prediction. The author is not arguing this will happen. She hopes it will not. But in order to help make sure it does not, we do need to hold the shade and see how it could.

becoming AI-assisted. If the day comes when both erode together when neither the sources outside nor the discernment inside can be fully trusted, the ground for knowing what is real gives way beneath the profession entirely. That is the deeper reason the verification framework is a survival strategy: it preserves both at once.

C. The Irreducible Human Contribution: Inner Knowing

There is a final dimension to the automation bias problem that this article has circled without naming directly until now. Call it professional intuition. Call it inner knowing. It is the felt sense of wrongness that arrives before you can articulate why, the moment when something in a document or an argument registers as off before your analytical mind has caught up to what you already sense. This is not mysticism. It is the product of years of attention and pattern recognition accumulated below the threshold of conscious thought. It is, in the oldest sense of the word, discernment. And discernment, for all its value, is work that comes only from within, through the deliberate development of the mind. Reading the cases. Doing the research. Staying present and paying attention even when the easier path is right there. This is not passive. It requires discipline, dedication, persistence, and the kind of quiet strength that does not announce itself but shows up every time.

The longing to be relieved of that burden of discernment is old and deep, older than the technology. It takes a familiar form: the willingness to hand that power to something outside ourselves. Across every tradition and every era, human beings have looked for an authority outside themselves, something that would hand down the answer, deliver the steps, carry the weight of judgment so they would not have to. The oracle. The institution. The algorithm. The impulse is understandable. The burden of discernment is real. If an oracle has the answer, or an institution, or an algorithm, then the hard internal work of developing one's own judgment is someone else's job. The delegation feels like relief. It is actually avoidance. And yet no external source, however wise or powerful or confident, has ever been able to do that work for us entirely.

What any genuine authority gives us is not the answer. It is the capacity to find it. The discernment remains ours. AI may be among the most fluent and persuasive of the authorities we have ever turned to, which is precisely what makes it one of the most important to be clear-eyed about. To know when something is wrong, and to say NO on the strength of that knowing, are not two separate skills. They are one, and they belong to the licensed professional whose judgment no tool can replicate or replace.

Everyone would prefer to be strong and fit without going to the gym. But skipping training causes atrophy. The same is true of the intellect. Failing to be rigorous and consistent in case preparation and legal analysis does not merely leave skills unused. It

diminishes them. Automation bias is dangerous not only because it causes lawyers to skip verification steps. It is dangerous because it degrades the very attention and capacity that make doing those things possible in the first place.²⁴

There is a second muscle at risk, and it is subtler: the gut. Legal training builds the habit of questioning, of pressing an argument until it yields or holds, of refusing to accept an answer simply because it arrives with confidence. When lawyers stop exercising that habit, stop pushing back, stop saying NO, they erode something different from the intellect but just as essential: the instinct that knows when to question before the analytical mind has caught up.

These two things work together. Rigorous intellectual work, reading the cases, doing the research, sitting with the argument long enough to feel where it strains, builds the foundation. The habit of questioning and refusal builds the instinct. Together they produce inner knowing: the capacity to recognize that something is wrong before you can fully articulate why. Neither alone is sufficient. And neither develops without the reps.

Every time a lawyer says NO to unverified output and goes to the source, they are not just avoiding a sanction. They are training both. The verification framework this article advocates is therefore not merely a compliance obligation. It is an exercise program, a practice of attention that compounds over time. Trust the process. Check the work. Read the cases. And when something feels wrong before you can say why, listen to that. It has been right before. It will be right again.

For this author at least, the most important thing she brings to her work is not her credential stack. It is the part of her that has been paying attention for decades, to law, to technology, to the moments when something does not add up. It is the time she has spent reading cases, doing research, building the foundation the hard way, at a law school that banned commercial outlines precisely because shortcuts erode the very capacity the work is meant to build. AI will not replace that. Nothing will. The obligation is to protect it, to keep showing up in the work with enough presence and attention that the inner knowing stays sharp. That is the real NO. Not just the verification step. The deeper practice of remaining awake.

²⁴ Jones, *supra* note 7.

VII. PRACTICAL GUIDANCE: HOW TO SAY NO EFFECTIVELY

The obligations Opinion 512 imposes, together with the jurisdictional duty that surrounds them, can be distilled into a handful of throughlines. None is new to a competent lawyer; each simply applies an old obligation to a new tool. This is not exhaustive, and it is not a substitute for reading the Opinion, the rules, and the law of your jurisdiction. It is a reminder of what the preceding pages have shown.

Competence requires verifying every citation independently, every time, in a trusted source before it goes into a filing.

Confidentiality requires keeping client information in enterprise tools with real data protections, never in consumer platforms that train on what you enter.

Communication requires disclosing AI use when your client, your engagement terms, or candor to the court demands it.

Candor requires that no unverified statement reach the tribunal, and that any error be corrected promptly upon discovery.

Supervision requires reviewing AI output as you would any nonlawyer assistant's work, and ensuring that the people using these tools know and follow your firm's policies, which should include a policy on AI usage

Honest billing requires charging only for the time actually spent, never for the time the tool saved.

And jurisdiction requires confirming the applicable rules in your specific court before every filing, because they vary, they change, and what was permitted last month in one courtroom may be prohibited this month in the one next door.

A. The Evolving Landscape: Check Your Jurisdiction

There is no single national rule governing AI use in legal filings, and that is precisely the practical problem. Since Opinion 512, more than thirty states have issued their own ethics guidance, and courts have layered their own requirements on top, varying not just by state or district but sometimes by individual judge within the same courthouse. Some courts require lawyers to certify whether AI was used and that a human verified every citation. Some require disclosure identifying exactly which sections of a filing used AI and which tool produced them. Others prohibit generative AI in filings entirely. The guidance ranges from Texas Opinion 705 to Ohio's recent guide for lawyers and judicial officers to New Jersey's and Florida's practical

resources.²⁵ And the direction of travel is unmistakable: jurisdictions are moving toward more oversight, more disclosure, and more accountability, not less. For the practitioner, the lesson is not to memorize any one rule. It is to check the applicable rules, in the specific court, before every filing, every time.

B. Firm and Organizational Policy

Individual diligence is necessary but not sufficient. Opinion 512's supervisory provisions require firms and legal organizations to adopt written AI use policies, train their people, and supervise AI-assisted work.²⁶ But a policy on a shelf protects no one. As the New Jersey Supreme Court put it plainly, a written policy alone is not a safe harbor; only implementation, training, supervision, and verification reduce the risk.²⁷ In a profession where AI tools are a click away on any personal device, people will use them whether or not the firm has authorized it, sometimes deliberately and sometimes simply because no one told them the rules. A policy works only when the people it governs know it exists, understand why it exists, and follow it. Model policies are now available from the New Jersey Judiciary, the American Inns of Court, and many state bars²⁸, and they offer a sensible starting point. But the document is the beginning of the work, not the end of it.

²⁵ Prof'l Ethics Comm. for the State Bar of Tex., Op. 705 (Feb. 2025) <https://www.legalethictexas.com/resources/opinions/opinion-705/>; Ohio Bd. of Prof'l Conduct, *Ethics Guide: Artificial Intelligence for Lawyers and Judicial Officers* (June 2, 2026), <https://www.bpc.ohio.gov>; Sup. Ct. of N.J., *Notice to the Bar: Responsible Use of Artificial Intelligence (AI) and Related Technologies* (Mar. 30, 2026), <https://www.njcourts.gov/sites/default/files/notices/2026/03/n260330c.pdf>; Fla. Bar Bd. Tech. Comm., *LegalFuel Guide to Getting Started with AI* (Jan. 17, 2025), <https://www.legalfuel.com/guide-to-getting-started-with-ai/>.

²⁶ ABA Formal Op. 512.

²⁷ Sup. Ct. N.J., *Notice to the Bar: Responsible Use of Artificial Intelligence (AI) and Related Technologies – Benefits of AI Policies* (Mar. 30, 2026), <https://www.njcourts.gov/sites/default/files/notices/2026/03/n260330c.pdf>.

²⁸ Sup. Ct. N.J., *Notice to the Bar: Responsible Use of Artificial Intelligence (AI) and Related Technologies — Benefits of AI Policies — Sample Starter Template* (Mar. 30, 2026), <https://www.njcourts.gov/sites/default/files/notices/2026/03/n260330c.pdf>; Am. Inns of

C. The Verification Process

In the language of AI governance, the principle that ties all of this together is called human-in-the-loop, or HITL²⁹: designing the workflow so that a person reviews and is accountable for what the machine produces before it has consequences. The term matters because the concept matters. AI governance experts, technologists, and now courts and ethics bodies all arrive at the same conclusion: the machine cannot be the last word. Not because the technology is new and will improve, though it will, but because the nature of the work requires a human being to own the output. In legal practice, that human is the lawyer who signs, and the practice is verification. The practice is simple and requires checking four things: that the case exists, that the court and reporter are right, that the citation format is accurate, and that the case actually says what it is cited for. That last step involves reading the case, it is the one automation bias most tempts us to skip, and it is the one that catches the inaccuracies that fabrication-hunting alone will miss.

The obligation to read the underlying source is not just this article's position. The trusted platforms themselves have acknowledged it directly. A practitioner using Lexis+ AI today is met with a warning displayed right on the tool:

This summary is generated by AI technology. AI generated content should be reviewed for accuracy.

The verification tool tells you to verify what the verification tool produced. This is not a criticism of the platform; it is an honest disclosure, and a correct one. But it makes the point unavoidable: the AI-assisted layer of even the most trusted legal research tools carries the same obligation as any other AI output. It must be checked against the underlying source.

Court, *Sample Law Firm Generative AI Use Guideline/Guidance Policy* (2024), <https://inns.innsocourt.org/media/199583/sample-2-law-firm-ai-gai-use-guidance.pdf>, cited in Catherine Reach, *Beyond the Ban: Why Your Law Firm Needs a Realistic AI Policy in 2026*, N.C. Bar Ass'n: From the Center Blog (Jan. 13, 2026), <https://www.ncbar.org/2026/01/13/beyond-the-ban-why-your-law-firm-needs-a-realistic-ai-policy-in-2026/>.

²⁹ *Key Terms for AI Governance: Human in the Loop (HITL)*, Int'l Ass'n of Privacy Professionals (IAPP) Glossary (2026), <https://iapp.org/resources/ai-governance-glossary?all%5Bquery%5D=hitl>

A natural question is whether a second AI tool can do the checking. We know all AI makes mistakes. And even though it is tempting to use it to check itself, because it will do the tedious work for us, and because automation bias makes us want to believe it knows better than we do, we cannot rely on it to be accurate.³⁰ When the stakes are high, when an error could mean compromising someone's personal information, the confirmation cannot be delegated to a tool that will, inevitably, sometimes get it wrong. It has to come from the source, confirmed by a human, through a channel a machine cannot fake.

This is not a theoretical concern. Organizations that handle sensitive requests are already discovering a parallel problem: they must now detect whether the requests they receive are themselves genuine, or AI-generated fakes. Courts are facing the same issue, finding that verification has to be done in new ways because a clerk can no longer tell, simply by looking, whether a filing is genuine.³¹ Whether AI is used to check AI or not, the principle is the same: a human must be in the loop. The author's own practice reflects exactly that principle. She uses one AI tool to cross-check, then reads the underlying sources herself, while remaining ever mindful that the second tool can fabricate as confidently as the first. Two machines can hallucinate the same thing. The only verification that ends the risk is a human confirming the source against reality. Keep the human in the loop, and the loop holds.

VIII. CONCLUSION: KNOWING AND NOING

The legal profession is in the middle of a transformation. AI tools are changing how legal research is conducted, how documents are drafted, how due diligence is performed, and how legal services are delivered. That transformation is real, and the efficiency gains it offers are real. This article does not argue that lawyers should avoid AI tools. It argues that lawyers who use AI tools must do so with the same ethical obligations they bring to every other aspect of their practice, and that the most important of those obligations, in the AI context, is the obligation to verify.

That is why the two most important letters in the future of legal practice are not AI. They are NO, the word that reflects the professional judgment that distinguishes a lawyer from a language model. Every obligation this article has traced comes down

³⁰ This holds whether or not the thing being checked was itself created by AI.

³¹ Nat'l Ctr. for State Cts., *A Legal Practitioner's Guide to AI Hallucinations* (2025), <https://nationalcenterforstatecourts.app.box.com/v/Legal-practitioner-guide-AI>.

to a single refusal: NO to letting the machine have the last word. Whether the duty is competence or confidentiality, candor or supervision, communication or fair billing, each is the same act underneath, a human being refusing to hand final judgment to a tool that cannot answer for it. Six duties, one word. And the last and simplest form of it is this: NO to the AI when it is wrong, because it will be wrong.

Because nobody is perfect. Nobody has ever been. The practitioners whose cases fill the sanctions landscape of this article were not necessarily reckless people or bad lawyers. They were working under pressure, reaching for an efficient tool, and falling into a trap that runs deeper than carelessness. Automation bias is not a character flaw. It is a cognitive phenomenon, documented across professions and decades³², through which human beings are predisposed to over-rely on automated systems and accept their outputs as true. These practitioners were not uniquely susceptible. They were human. And as humans, they have inherent limitations and unavoidable fallibility.

Ironically, so does AI. It turns out the tools we reach for are “only human” too in the sense that they are, fallible in many of the same ways that we are. They are prone to the same shortcuts, and as capable as any one of us of getting it wrong with complete confidence.

They are prone to telling people what they want to hear, and inclined to fill gaps with plausible-sounding information rather than admitting it does not know. Unlike a human, AI never tires, never needs to rest or sleep, never has to leave early for carpool or go on vacation. So, the answer is not to disconnect from the technology. Because doing so just because it is not perfect would mean giving up the very things that make it valuable. The answer is to maintain the health of the individual practitioner, the independent judgment, the inner knowing, the professional courage to say NO, that makes the collective human intelligence trustworthy. Because the collective is only as trustworthy as the individuals who comprise it.

Like forests that communicate and share resources through underground fungal networks, exchanging nutrients, chemical signals and warnings, so that individual trees participate in a collective intelligence none of them possesses alone, technology has become the mycorrhizal network of human knowledge, the infrastructure through which collective intelligence flows. But a network is only as healthy as its nodes. Every lawyer who keeps their independent reasoning sharp, who verifies every citation, who says NO when NO is the right answer, is a healthy node. And the health of each node

³² Jones, *supra* note 7.

comes down to the same thing: judgment. That judgment was always yours. It always will be. And using it wisely is more important now than ever, because we are at something more than an inflection point.

What the legal profession is experiencing is only a slice of what the entire world is experiencing. The same usefulness that makes AI indispensable also makes its errors harder to see. The same tool that makes legal research faster makes it easier to skip verification. The same confidence that makes AI output convincing makes its mistakes invisible until they reach a courtroom. This is the light and shade at the micro level. At the macro level, the light and shade of the storological³³ moment we are all living in is this: as Anthropic has stated plainly, AI has the potential to pose unprecedented risks to humanity if things go badly, and the potential to create unprecedented benefits if things go well.³⁴ Both are true at once.

Return, finally, to the change that began this article: that we now converse with our machines, and that we accepted it without pause. There was no moment of collective reckoning, no stopping to ask what it would mean to hand our questions to something that answers with the confident fluency of a colleague and the accountability of no one. We simply stopped clicking through, stopped reading the underlying sources, because the answer was handed to us. It was as if the entire world stopped exercising judgment because it no longer had to, and no one noticed. That absence of pause is itself the story.

Which is why the practitioners who matter in this moment are those who continue to verify the law by reading the cases, who keep humans in the loop, who refuse to sleepwalk across the threshold and remain awake. They are the ones who will be remembered when we look back on this moment. Not because they were perfect. Because they were humble enough to know they weren't, and kept showing up anyway.

The exercise of judgment is the irreducible core of legal practice. It is what lawyers offer clients that AI cannot. It is what courts expect when a lawyer signs a filing. And in the age of artificial intelligence, that judgment comes down to two words that sound the same: knowing and NOing. The knowing is the discernment built over years of doing the work, the felt sense of when something is wrong. The NOing is the courage to act on it, to refuse the easy answer the machine hands you. Neither is worth much

³³ The author uses storology to mean the study of real, lived human stories and personal breakthroughs; storological is its adjectival form.

³⁴ Anthropic, *Company Values*, Anthropic (2026), <https://www.anthropic.com/company>.

without the other. Knowing without the NO is a hunch you talked yourself out of. The NO without the knowing is just noise. Together they are the whole of it.

Just as we must hold both the light and the shade, we must hold both the knowing and the NOing, in this pivotal moment, to maintain the integrity of our profession. Because the alternative is a profession drifting through its work in hovering chairs, no longer quite remembering how to practice. The alternative is that we are all just floating around in WALL-E world.

AI Assistance Disclosure: The author used Claude (Anthropic) as a drafting and organizational tool in the preparation of this article. All legal analysis, professional judgments, personal experience narratives, factual claims, and citations reflect the author's own expertise and have been independently reviewed and verified by the author. The author takes full responsibility for the content of this article. This disclosure is consistent with the article's own argument about transparency and professional integrity in AI-assisted legal work. See the Appendix for a full discussion of the authorship process.

A candid note: the author has written and published several articles without AI assistance, and those articles contain errors. If this article contains errors, readers will almost certainly assume the AI made them. The author accepts this as a fair trade.

APPENDIX: A NOTE ON AUTHORSHIP: HOW THIS ARTICLE WAS WRITTEN

This article was drafted with the assistance of artificial intelligence. Specifically, the author used Claude, developed by Anthropic, as a drafting and organizational tool throughout the writing process. The author directed the work, provided the subject matter expertise, supplied the personal experience narratives, and verified and edited every element of the final document. She takes full professional and intellectual responsibility for the content.

The process began with an idea and a rough outline. From that outline, Claude produced a first draft organized around the author's structure. That draft was little more than a placeholder compared to the article you are reading now, but it was useful in ways that mattered. It made the structure visible and concrete, so that changes to organization and scope could be made early, before the real writing began. It made starting from a blank page less daunting. And it gave the author something to push against, which is often where the best thinking happens.

From there, the author worked paragraph by paragraph. She read what Claude had produced, verified every source and citation she intended to use, and then rewrote. And

rewrote again. Rereading each paragraph, revising it, reading it again. Reorganizing the placement of ideas to support a more coherent structure, then rewriting again and checking again. Ensuring each paragraph transitioned cleanly into the next. Adding new sections almost daily as new ideas presented themselves, then rereading and revising those as well. For the most part, she wrote this article the same way she would have written it entirely on her own.

Claude could not have done it alone. The first draft read well and looked polished and might have passed a cursory review, but anyone looking closely would have found it lacking the underlying proof and logic needed to support the thesis, the ideas, and the arguments the article makes.

The author does not want to understate the value it provided, either. Beyond drafting assistance, Claude served as a genuine intellectual collaborator. The author used it to think out loud, to test arguments, to ask whether an idea held together, to explore how a personal experience connected to the article's larger themes. Claude would respond: here is how that ties in, here is where that argument might strain. It was, in the author's experience, one of the most patient and available thinking partners she has ever had, which is one of the things this article says about AI, and which she found to be true.

But even without that, the drafting assistance alone was invaluable. To call her editing process rewriting is a stretch. What she actually did was author, in the most unglamorous sense of the word. She wrote the way she has always written: in messy, typo-ridden, grammatically approximate prose that captured her thinking before her fingers caught up. She might as well have dictated. Then, paragraph by paragraph, she pasted her drafts into Claude and asked it to clean them up. Claude did. The ideas, the analysis, the personal stories, the arguments, all of it came from her. The polished sentences came from the collaboration.³⁵

This process, not the use of AI, but the process of collaborating as just described is not as novel as it might appear. It is standard practice in the legal profession.

The Court in the Mata opinion recognized this and began its opinion listing common examples of when this occurs, “[i]n researching and drafting court submissions, good lawyers appropriately obtain assistance from junior lawyers, law students, contract lawyers, legal encyclopedias and databases such as Westlaw and

³⁵ This is an example of an uncleaned up sentence the author might have typed.

LexisNexis.”³⁶ These are all forms of collaborative authorship the author has experienced firsthand, and which the legal profession has long accepted as legitimate. The author will explain each in more detail for those who might not have the same type of experience.

The simplest and oldest parallel is the legal secretary. For most of the twentieth century, attorneys dictated their letters, memoranda, and briefs, speaking into a recording device or directly to a secretary who transcribed the words. The attorney then reviewed, edited, and signed the final document. No one questioned the attorney's authorship. The intellectual content was theirs. The mechanical production was delegated. AI-assisted drafting follows exactly the same structure, with a language model replacing the transcription function.

In law school, the author worked as a research assistant for her favorite professor. He wrote the way many academics write: he put his ideas on paper first, then hired her to find the authorities that supported those ideas. The citations that appeared in his published article were found by her, organized by her, and formatted by her. The ideas were his. The responsibility was his. The research assistance was hers. No one questioned his authorship.

The author also held two clerkship roles at different stages of her career, each of which illuminates the same principle of collaborative authorship from a different angle. The two roles share a word but are quite different in practice, and the distinction matters. The first, the law student law clerk, is a student employment role. The second, a post-graduate federal judicial clerkship, is a different and more prestigious appointment entirely: a newly licensed attorney, typically selected from among the top graduates of accredited law schools, who works directly for a federal judge for one or two years.

In her role as a student law clerk during the summer between her first and second years of law school, the author often drafted pleadings for the attorneys she worked for. One week, she wrote a motion for an attorney who filed it with the court. The judge, a state court judge, called the attorney in for a meeting. The attorney came to the author in something close to a panic. She had to go with him. He had to tell the judge that she had written it because he would not be able to explain anything in it to the judge himself since he did not do any of the research and writing. The judge was not troubled by this at all. What the judge wanted to discuss was the quality of the

³⁶ *Mata v. Avianca, Inc.*, 678 F. Supp. 3d 443, 448 (S.D.N.Y. 2023).

legal research and reasoning. He said it was among the best he had ever seen. He offered her a job on the spot.³⁷

The second role told a similar story from a different vantage point. When the author began her federal judicial clerkship, her co-clerks warned her before she submitted her first draft order that it would come back heavily marked in red. That was normal, they said. Do not take it personally. The author submitted her first draft and waited. After what seemed like an unusually long time, the draft could not be found: not on the judge's desk, not in the intake bin, not in the review bin. Eventually it emerged that the judge had signed it and it had been filed. No edits. The outcome was unusual. The process that produced it was not.

In a federal district court, the drafting process works as follows. Pleadings arrive from both sides of a case. A law clerk, not the judge, reads them, researches the applicable case law, and drafts a proposed opinion recommending how the case should be decided. That draft goes to the judge for review. The judge makes edits, sometimes extensive, sometimes none at all, and issues the final order under the judge's own name and authority. The law clerk's name appears nowhere on the opinion. The judge's does. That is authorship in the legal tradition: the person who reviews, edits, and signs owns the work.

This drafting process is structurally identical to the AI-assisted drafting process this article describes and defends. The clerk researches, drafts, and checks the accuracy of every citation. The judge reviews, edits, and issues the final product under the judge's own authority. The intellectual labor is divided. The accountability belongs to the person who signs.

Ghost writing follows the same pattern. A recognized expert provides the ideas, the experience, the credibility, and the direction. A skilled writer provides the organization, the prose, and the structure. The expert reviews, revises, and puts their name on the final product. The author has been on both sides of this relationship.

The legal academic publishing process offers yet another parallel. Anyone who has published in a law review, as the author has, knows that the finished article is the

³⁷ She did not take the offer. She has wondered since whether she lacked the self-confidence to believe the judge meant it, or whether loyalty to the attorney kept her from accepting, or whether she simply did not yet understand what the offer meant. She understands it better now.

product of multiple collaborators. The author writes. The articles editor edits. The cite checker verifies every source. The author herself served in both roles during law school, working on the Human Rights Quarterly as an articles editor and a cite checker, reading and checking the citations in articles written by professors and practitioners, ensuring that the authorities cited actually supported the propositions for which they were cited.

The same collaborative authorship structure governs publishing more broadly. Every major book passes through developmental editing, line editing, copy editing, and proofreading, each performed by a professional whose name does not appear on the cover. Newspaper journalists write under editors who may substantially revise their work. Magazine writers submit drafts that undergo editorial transformation before reaching readers. In none of these contexts is the author's byline considered illegitimate because an editor improved the work.³⁸

The author has been the clerk, the research assistant, the ghost. She has found the cases, written the briefs, explained the law to the people whose names went on it. She has been, and here again is, the one whose name goes on it. What has not changed, what has never changed, is the obligation to take responsibility for what you put your name on. That obligation does not diminish because a new tool made the drafting faster. If anything, it intensifies. Because the tool will not be there when the judge asks you to explain your reasoning. You will be. Make sure you can.

This is part II in a [series](#).

³⁸ The legitimacy of this collaborative authorship tradition is recognized not merely by custom but by rule. The Kentucky Bar Association explicitly grants CLE credit to attorneys who research, write, or edit materials presented by another attorney at an approved CLE program, calculated at one credit per two hours of preparation time, up to twelve credits per educational year. Ky. Sup. Ct. R. 3.655(2)(c). The bar association that licenses attorneys and enforces professional responsibility has codified in its own rules that writing for others is legitimate, valuable, professional work worthy of formal recognition. If Claude were licensed to practice law in Kentucky, it would be eligible for up to twelve credits for its contributions to this article and CLE program. It is not. And that distinction, the one that makes licensure meaningful, that makes responsibility real, that makes the signature on a filing matter, is the whole of the argument this article makes.

Erin Corken, Esq. CEDS · FIP · AIGP · CIPP/US · CIPP/E · CIPM is an attorney licensed in Kentucky and Arizona, a law professor and scholar, and a Senior Solutions Engineer and Academic Program Developer at Exterro. A former law clerk to the Chief Judge of the United States District Court for the District of Arizona, she has for more than a decade taught E-Discovery and Information Privacy and Cybersecurity at the law school level, along with Business Law at the undergraduate level, and she has presented at bar associations, paralegal organizations, and legal technology conferences nationwide, including PLI. She has published previously in the Northern Kentucky University Chase Law Review, the John Marshall Law Journal, Bloomberg BNA's Digital Discovery & e-Evidence reporter, and other publications. The author thanks Ronald J. Hedges, former United States Magistrate Judge for the District of New Jersey, for his contributions to this article.

	
Interested in writing for the <i>PLI Current</i>? Get involved or visit pli-current-contribute.pdf for more information	Sign up for a free trial of PLI PLUS at pli.edu/pliplustrial

Disclaimer: The viewpoints expressed by the authors are their own and do not necessarily reflect the opinions, viewpoints and official policies of Practising Law Institute.

This article is published on PLI PLUS, the online research database of PLI. The entirety of the PLI Press print collection is available on PLI PLUS—including PLI's authoritative treatises, practice guides, skills books, periodicals, forms & checklists, and course handbooks and transcripts from our original and highly acclaimed CLE programs.