

**A COURT-READY AI RESOURCE**

The Cross-Examination Test: 12 Questions Opposing Counsel Will Ask About AI in Your Casework

If AI has touched any part of your casework, someone on the other side is preparing to ask you about it under oath. This guide gives you the 12 questions the way counsel will actually phrase them, explains why each one gets asked, and shows you what a strong answer sounds like next to a weak one. Read it with your last report in front of you. Score yourself honestly at the end.

By Eric L. Waldrep, MCFE, certified UAS/Drone forensics examiner, U.S. State Department ATA Cyber Mentor. 27 years in law enforcement. Digital forensics since 2007. Digital evidence expert witness.

Why this test exists

The courtroom record on AI is no longer thin. By mid-2026, roughly 1,500 court decisions worldwide involving AI-fabricated citations had been tracked, with sanctions reaching five figures and, in one case, a license suspension. In *Kohls v. Ellison*, a federal court in Minnesota excluded an AI expert's declaration because his AI-assisted filing cited sources that did not exist, and the court refused to let him fix it. And in May 2026, a federal court ordered an expert to produce the AI prompts he used in his workflow as discoverable methodology. That order is stayed pending review, but the demand has now been made in a federal courtroom, and it will be made again.

I have spent 27 years in law enforcement, 11 of them as a detective, and I have sat in the witness chair enough times to know that cross-examination does not test what you did. It tests what you can prove you did. AI questions are no different, and counsel can hurt you without understanding transformer models: one question from this list that you cannot answer with a document is enough.

None of what follows is legal advice, and answering all 12 questions well guarantees nothing about how a court will rule. What it does is remove the easy wins. Here they are.

The 12 questions



QUESTION 1

“Did you, or anyone working under your direction, use artificial intelligence at any point in your work on this case?”

This is the opening question because it sets the trap for everything after it. Many courts now have standing orders on AI use in filings, and ABA Formal Opinion 512 applies candor duties to generative AI use, so a false or hedged "no" is checkable. In the Watson Grinding litigation, opposing counsel obtained an expert's ChatGPT conversations directly. If use surfaces after you denied it, counsel will spend the rest of the cross on your credibility instead of the technology.

A STRONG ANSWER SOUNDS LIKE

“Yes. AI tools were used for specific administrative and drafting-support tasks, each one permitted under our written AI-use policy. Every use is recorded in the AI-use disclosure that accompanies my report, and I can walk you through each entry.”

A WEAK ANSWER SOUNDS LIKE

“Not really. I may have used a chatbot to clean up some wording, but nothing that mattered.”

QUESTION 2

“Which tool, which version, and whose account?”

Counsel asks this because the answer separates a managed tool from a personal habit. Account type determines who administered the settings, what the terms of service said, and whether anyone in your organization could even see the activity. "My personal account" opens an entire line of questioning about where case-adjacent material went and who agreed to what terms on the organization's behalf.

A STRONG ANSWER SOUNDS LIKE

“It was [product], on our organization's business plan, under an account our agency administers. I recorded the version and the date of use in the case file, and the account configuration is documented.”

A WEAK ANSWER SOUNDS LIKE

“Whatever the free version was at the time. It was on my phone.”



QUESTION 3

“What information from this case did you put into that tool?”

Typing case material into a chat window can amount to disclosing it to a third party, depending on the terms and configuration, and ABA Formal Opinion 512 puts confidentiality squarely on the professional using the tool. This question is where privilege problems start, with protective-order and victim-privacy problems right behind them. Counsel is hoping you never thought about the boundary at all.

A STRONG ANSWER SOUNDS LIKE

“None. Our policy classifies data into permitted and prohibited categories, and case data, evidence, and personal identifiers are prohibited from entering any AI tool. I used generalized and synthetic text only, and I can produce the classification rules I followed.”

A WEAK ANSWER SOUNDS LIKE

“Nothing sensitive, as far as I remember.”

QUESTION 4

“Who authorized you to use that tool in your casework?”

In the 2025 Clio Legal Trends Report, 53% of legal professionals said their firm has no AI policy or they were unaware of one. Counsel knows those numbers and is betting you are in the majority. An approval chain shows an organization exercising judgment. Its absence lets counsel paint every AI touchpoint in your workflow as freelancing.

A STRONG ANSWER SOUNDS LIKE

“Our written policy names the approved tools and the approving authority. My use falls under it, and the approval is on file. I did not make that call alone.”

A WEAK ANSWER SOUNDS LIKE

“Nobody had to approve it. It's a tool, like a calculator.”



QUESTION 5

“Can you produce the exact prompts you gave the tool and the exact outputs it returned?”

In Matter of Weber, a New York court rejected an expert's valuation after he could not recall the prompts he had given a chatbot. And as noted above, a federal court has ordered an expert's prompts produced as discoverable methodology, though that order is stayed pending review. The safe assumption is that prompts are bench notes: part of the record of how you worked. If they are gone, counsel will argue your methodology is gone with them.

A STRONG ANSWER SOUNDS LIKE

“Yes. Prompts and outputs are retained as part of the case file under our retention standard, the same way I keep my examination notes. They are available.”

A WEAK ANSWER SOUNDS LIKE

“Those chats are long gone. I never thought of them as records.”

QUESTION 6

“Does that tool learn from what you type into it? Where does your data go?”

This question tests whether you understood the agreement you were working under. Data handling differs sharply by vendor and by account tier, and the terms change over time. If you cannot answer, counsel will invite the jury to imagine your case material sitting in someone's training data, and you will have no document to answer with. Guessing confidently is worse than not knowing, because the terms are checkable.

A STRONG ANSWER SOUNDS LIKE

“For the account tier we use, the vendor's terms in force at the time of use did not permit training on our inputs, and we retained a copy of those terms with the case file. I can state where the data was processed and how long the vendor retained it.”

A WEAK ANSWER SOUNDS LIKE

“I assume they don't keep it.”



QUESTION 7

“How did you verify what the tool told you?”

Fabricated output built the court record on AI: the roughly 1,500 tracked decisions worldwide all involve AI-fabricated citations, and Kohls shows how it ends: exclusion, with no second chance. The examiners who get through this question are the ones who treated every AI output as unverified until checked against its original source. The ones who do not are the reason the tracker keeps growing.

A STRONG ANSWER SOUNDS LIKE

“Nothing from the tool entered my work product until I verified it against the original source: the statute, the paper, the log, the underlying data. Each verification is logged, and I can show you the log for this case.”

A WEAK ANSWER SOUNDS LIKE

“It's usually accurate, and it cited its sources.”

QUESTION 8

“Did the AI write any part of the report you signed?”

In the Watson Grinding matter, opposing counsel obtained the expert's ChatGPT conversations and showed that most of the filed report tracked the chatbot's output. Compare *Ferlito v. Harbor Freight*, where testimony was admitted: the expert wrote his report first and used AI only to check it afterward. The order of operations is the whole question. Courts have accepted the second pattern. The first one hands counsel your signature next to a machine's sentences.

A STRONG ANSWER SOUNDS LIKE

“I formed the opinions and wrote the findings myself. AI was used afterward to check consistency, formatting, and completeness against my own draft. My drafts exist and show that sequence.”

A WEAK ANSWER SOUNDS LIKE

“It produced a first draft and I edited it. The conclusions are still mine.”

**QUESTION 9**

“Your forensic software uses AI too, doesn't it? Tell the court what its AI did to the evidence in this case.”

Most examiners prepare for the chatbot question and get blindsided by this one. In *State v. Puloka*, AI-enhanced video was excluded because the algorithm could not be explained or validated. The AI embedded in your forensic suite counts even if you never opened a chat window. And because SWGDE has published no best-practice standard on AI-assisted digital forensic examination, there is no standard to point at. There is only you, explaining what you relied on and what you confirmed yourself.

A STRONG ANSWER SOUNDS LIKE

“We maintain an inventory of the AI-assisted features in our tools. I can tell you which ones ran in this case, what they touched, and where I independently confirmed their output against the underlying data before relying on it.”

A WEAK ANSWER SOUNDS LIKE

“The tool is industry standard. I can't speak to what's under the hood.”

QUESTION 10

“Show me your lab's AI policy. What date is on it?”

In the 2025 Thomson Reuters survey, 52% of professionals said their organization has no generative AI policy at all. Counsel asks for the date because a policy written after the examination is nearly as damaging as no policy: it concedes the work was done in an unmanaged environment. The document itself matters less than what it proves about when your organization started paying attention.

A STRONG ANSWER SOUNDS LIKE

“Here it is. It predates this examination, it names approved tools, prohibited data, disclosure requirements, and retention rules, and this is the version that was in force when I did the work.”

A WEAK ANSWER SOUNDS LIKE

“We have general guidance about being careful with new technology.”



QUESTION 11

“What training have you had on the AI tools you used?”

The same Thomson Reuters survey found 64% of professionals received no AI training. Counsel uses this question to convert every earlier answer into a competence problem: you used a tool you were never trained on, under a policy you may not have read, and you want the court to rely on the result. Documented training closes that door.

A STRONG ANSWER SOUNDS LIKE

“I completed documented training on both the tools and our policy before using them in casework, and it is refreshed on a set schedule. The records are available.”

A WEAK ANSWER SOUNDS LIKE

“I picked it up on my own. It isn't complicated.”

QUESTION 12

“The tool you used has been updated since your examination. Is the version you relied on still available, and would it produce the same output today?”

Models change underneath their users, and vendors retire versions on their own schedules. Counsel asks this to suggest your work is unreproducible: whatever you validated no longer exists. You cannot freeze time, but you can keep a record that stands on its own, backed by a policy that treats every model change as a trigger to revalidate before the tool touches casework again.

A STRONG ANSWER SOUNDS LIKE

“We record the tool and version at the time of use, and a model or version change triggers revalidation under our policy before further use. Because my prompts and outputs are preserved along with the verification log, my record does not depend on what the vendor ships next.”

A WEAK ANSWER SOUNDS LIKE

“It's the same tool. The updates just make it better.”



Score yourself

Count how many of the 12 you could answer today, on the stand, with a document to back the answer. The document is the standard. An answer you cannot support in writing counts as a miss.

10 to 12: Ready. You have the makings of a disclosure-ready AI record. Ready means prepared to answer, not guaranteed to prevail; no honest document promises a court outcome. Your remaining work is keeping the record current as tools and policies change.

6 to 9: Exposed. You have discipline in places and gaps in others, and you do not get to pick which question counsel opens with. The gaps are usually the quiet ones: prompt retention, vendor-embedded AI, the policy date.

0 to 5: At risk. If AI is anywhere in your workflow right now, including inside your forensic tools, the record being built about your work is being built by accident. Fixing this before your next report costs far less than explaining it on cross.

What to do next

The fastest way to find your actual gaps is the AI Exposure Review: a 90-minute structured session with me, followed by a written AI Exposure Snapshot within 3 business days. It costs \$1,500, prepaid, and 100% of it credits toward the Court-Ready AI Assessment if you move forward within 90 days. No case material is needed or accepted; we work from how your team operates, not from your evidence.

Book an AI Exposure Review at thewaldrepcompany.com/court-ready-ai/.

The Waldrep Company. Documented, auditable AI use with human validation.
