

Inside the First Stipulated AI Review Protocol

Meet and Confer Podcast · Kelly Twigger ·  [Listen](#) · Case in Minerva26: *James v. Cerebras Systems, Inc.*, No. 4:25-cv-09361 · **Orders:** [ECF No. 74](#) (ESI Protocol) and [ECF No. 75](#) (Stipulated Protective Order) · **Entered:** July 7, 2026 · **Judge:** Magistrate Judge Robert M. Illman · Northern District of California (Oakland Division) · Magistrate Judge Robert M. Illman

Introduction

Hi, and thanks for joining me this week. Last week I was on vacation, enjoying the incredible beauty of Yellowstone and Grand Teton National Parks near Jackson, Wyoming. It was the first time in probably years that I have not touched my laptop for a week, and it was blissful. So if you need a push to take a real, unplugged vacation, I'm giving it to you right now. If it helps, you can do what I did and go somewhere that has no internet coverage, so you don't have a choice. That works.

All right — back to our topic du jour. For the last year on this show, I've been tracking how courts are working through artificial intelligence and discovery: the privilege fights over a litigant's prompts, the judges starting to write AI terms into their protective orders, the slow realization that the rules we have are not built for any of this. Today we hit a genuine milestone in that arc, and I want to walk you through all the pieces carefully, because it is going to shape how these protocols get written for the next several years.

This case is called *James v. Cerebras Systems*. And what the parties did here is, as far as I have seen, a first: a stipulated ESI protocol that governs the use of generative AI to review and produce documents as its own category — with its own disclosure rules, its own validation math, and a requirement that a party using AI to decide what gets produced disclose the prompts it used to do it. And the parties agreed to all of it. Nobody fought. That, as much as any single provision, is really the story here. And I'll tell you up front that when I read it, I thought I would never agree to any of this — so I'm not entirely sure why they did.

Welcome back to the Meet and Confer podcast. This is the Case of the Week series, where I choose a recent decision or order on the discovery of ESI and turn it into practical advice for you to use. My name is Kelly Twigger. I'm the CEO and founder at Minerva26 and the Principal at ESI Attorneys, and I've been a discovery strategist and practicing attorney for just about 30 years.

Let's talk about why orders like this one belong on your radar. There are more than 5,000 discovery decisions issued at the federal level every year, and that has held for the last four to five years. About 80% of those decisions on ESI come from magistrate judges. And less than 1% of civil cases go to trial, which means these cases are won and lost on discovery. If you are not watching how sophisticated parties are handling AI in that process, you are already behind. And you may want to make the argument that your client doesn't have as much information, they're not as sophisticated, they don't have the budget. Here's my response to that: the less money you have, the more diligent you need to be about addressing these ESI issues. So if you have clients dealing with smaller matters on lower budgets, you actually need to have more knowledge than your clients who have wiggle room for you to figure it out.

Now, in this case there are two orders, both entered by United States Magistrate Judge Robert M. Illman on July 7, 2026. And you need to keep them straight, because they do two different jobs. Document 74 on the docket is the ESI Protocol — the Stipulation and Order Regarding the Production of Electronically Stored Information. Document 75 is a companion order — the Stipulated Protective Order. I'll come back to how the two work hand in hand, because it turns out to matter a great deal how the parties set them up here.

Before we get into the provisions, I want you to remember something. Nothing in the Federal Rules required any of this. Rule 26 obligates a party to provide relevant, proportional discovery. Rule 34 governs the form in which documents are produced. Neither rule says a word about disclosing your prompts, your validation methodology, or your elusion rate. Those obligations exist in this case for one reason and one reason only: the parties agreed to them. This was a stipulated order — negotiated and consented to — not an order a court imposed after a fight, not the kind of scorched-earth discovery battle we watched in the OpenAI cases. So as I walk through how demanding this protocol is, keep thinking about this question: why would sophisticated parties agree to expand their own obligations this far beyond what the rules would require?

We've had this debate on Case of the Week about TAR and search terms — how much are parties really obligated to share in meeting their obligations under the rules? It's been left up to what the parties agree to, or what individual courts order on motion. We've not seen this before, where parties have just flat-out agreed to it. Let's dive into the facts.

The Facts

The case underneath these orders is a copyright class action. The plaintiffs are authors. The defendant, Cerebras Systems, released a family of open-source large language models called Cerebras-GPT, and the plaintiffs allege that these models were trained on a dataset called The Pile, which contains a subset known as Books3 — described in the complaint as derived from pirated books.

Let me be precise about exactly what this case is, because the "AI chip company" label for Cerebras can mislead you. Cerebras is best known as a hardware company that builds chips for AI — think of it like an NVIDIA that you might be familiar with. But this claim is not about the chips. It is about Cerebras-GPT, the models, and the allegation that they were trained on pirated

books. Cerebras used its own chips to build a model to show off, allegedly, how much faster its chips are than others on the market. That makes this case materially identical to the author class actions against OpenAI, Anthropic, and Meta — same theory, same kind of dataset, just a different defendant. The only real wrinkle is that a company better known for selling the hardware is now on the receiving end of the same copyright claim as the model builders. Now let's look at who is at counsel table, because it really matters here. For the plaintiffs: Lieff Cabraser Heimann & Bernstein and McGuire Law — serious class action firms who know and understand discovery. For Cerebras: Kecker, Van Nest & Peters. Kecker may not be a firm you've heard of, but they are a very well-respected litigation boutique with very high-profile clients. And Kecker is the same firm, with some of the same lawyers — including Robert Van Nest — that serves as OpenAI's lead trial counsel in the AI copyright cases. So the lawyers who negotiated this meticulous, AI-specific protocol are the same defense lawyers living through the discovery wars in the OpenAI litigation. That is not a coincidence, and I think it's a large part of why this order looks the way it does.

One more piece of context, specific to this courthouse. The Northern District of California has had a model stipulated ESI order on its books for litigants to adopt since 2015. It is, quite frankly — and I've discussed this before — pretty hopelessly out of date, because it predates the discovery problems we've dealt with since 2015, let alone generative AI. The ESI Protocol says on its face, in the opening recital, that the parties started from the court's model order, and the court's standing order required them to attach two things: a declaration explaining every modification they made, and a redline comparing their version to the model. They had to show their work. The fact that this judge let them modify a decade-old form this heavily — building entire TAR and AI appendices into it — is a good sign. It is also different from how a lot of judges treat their standing protocols, where deviation is discouraged and you're nudged back toward the form. Magistrate Judge Illman let the parties tailor the model to the technology. That's a great instinct, and I hope more courts follow that path.

The Analysis

All right, let's start with the shape of the entire protocol, because the AI piece only makes sense once you see where it sits.

The heart of this ESI protocol is Section VII, titled Search and Review. And it does something I want you to take in as a whole before we get to any one part. It expressly authorizes and governs three different ways to cull and review documents, and it addresses combining them — so-called layering. First, search terms — capped at 20 per custodian per party, with required hit reports and a null-set test on the documents that do not hit. Second, technology assisted review, or TAR — the predictive-coding process where you train a classifier to rank documents by likely responsiveness — which gets its own rulebook in Appendix 3. Third, artificial intelligence — which gets its own separate rulebook in Appendix 4. And Section VII.6 governs what I refer to as layering, which is using more than one of those methods on the same set of documents. So the architecture is simple to state: you choose your method or methods, you disclose that choice, and if you're using TAR or AI, you hand over the matching appendix — three or four — pursuant to the protocol. Everything else flows from that.

Now, most of what this protocol requires goes well beyond Rule 26. And a fair amount of it

brushes right up against material that is arguably privileged or work product — the prompts you wrote, how you trained your model, your validation results, your error rates. My instinct as counsel is generally to protect all of that and not hand it over. So why consent here? What is the thought process behind the parties? The most plausible read I can see is that counsel who lived through the OpenAI discovery fights decided that a detailed, agreed-upon protocol up front is going to be cheaper and more predictable than litigating every AI question by motion, one expensive fight at a time. It also draws the litigation out. And OpenAI lost most of those battles and had to disclose the information anyway. So this level of detail buys peace — but you buy that peace by expanding your own obligations and trading away protections you would otherwise have had. That's a real strategic decision, and not an obvious one. But it goes to one of the themes we're always talking about here on Case of the Week: you need to sit down early in your case, understand what the issues are going to be, understand who counsel is on the other side, and make strategic decisions appropriate for that case. Protraction and hostility ultimately cost money. You have to make good decisions for your clients — and perhaps that's what drove the protocol here.

Now, with that frame, let's talk about the actual protocol itself. Section VII.8 creates a distinct regime the order calls "AI Responsiveness Review." And it defines it broadly, as workflows using — quote — "large-language models, deep-learning classifiers, embedding-based similarity analysis, semantic clustering, predictive redaction systems, or generative summarization models" to decide whether a document gets produced, withheld, or redacted. Six different technologies that make up AI, not one. That's a far wider net than TAR. If you use a large language model, or any of the tools in that list, to make responsiveness or privilege calls, you are in the regime of this protocol, and you must disclose that election to the requesting party. And under Section VII.9, you must provide an AI Review Protocol containing everything in Appendix 4.

So let me tell you what Appendix 4 actually requires, because the scope of it is what will blow your mind. From a workflow standpoint: if you use AI to decide what gets produced, then how the AI reached those decisions — every single one of them — must be documented, your privilege calls must be human-confirmed, and the entire process is disclosed to your opponent. The AI's decision-making actually becomes part of the discovery. Specifically, Appendix 4 requires you to disclose, at a minimum, these things:

- **The system itself** — the identity, version, and hosting environment of each AI model, whether it runs on-premise, in a private cloud, or in a vendor-secured environment, plus a statement that the system is "generally accepted within the eDiscovery industry as a reasonably reliable tool." We'll come back to that. But if you use these AI tools on a regular basis, you'll notice that in the lower left-hand corner you often see a "relaunch" button almost every day you log in — meaning that every day you use the system, you may be using a different model. So one question I have in reading this protocol is: do I have to document every single time the model changes if I'm updating it? It sounds like the answer is yes.
- **The universe** — the criteria used to select the documents the AI reviewed, including any culling done before the AI ran (search terms, date ranges, custodians, deduplication, threading, near-duplicate suppression), and which document types were excluded from AI analysis and how those were reviewed instead. If that review is manual, who reviewed them? It even asks for the credentials of the reviewers doing this work.

- **The training and prompts** — the methodology used to train or instruct the system; the sources of training and validation data; the identity and qualifications of the people who designed, trained, and validated it; any documents excluded from training; and disclosure of all prompts, templates, instruction sets, and parameter configurations, with any change to a prompt served in redline within three business days. Now think about how often you iterate on your prompts. That is going to be quite a feat to provide within three business days. There is going to have to be a very well-articulated process up front, or this becomes a nightmare.
- **The scoring** — how the system generates scores, classifications, or decision boundaries, and the thresholds used; the distribution of scores from the initial run; and how many documents were placed in each category, including responsive, non-responsive, privileged, uncertain, or requiring further review.
- **The oversight** — who conducted oversight, privilege review, and validation; the review workflow used; the number of documents manually reviewed; and the quality-control procedures, which must include human confirmation of privilege determinations, checks for hallucination, over-summarization, and misclassification by the AI, and testing of prompt performance.
- **The iteration** — how many times the model or the prompts were adjusted during the review, and how that changed the results. Again, those redlined prompts have to be provided within three business days.
- **The validation** — the validation method used; a sample of the documents the AI excluded as non-responsive, sized at a 95% confidence level with a plus-or-minus 2% margin of error; the resulting richness, recall, precision, and elusion rates; and a duty to meet and confer and take corrective measures if elusion exceeds 3%.
- **The audit trail and security** — all AI processing must occur inside the secure review environment holding the documents, with no content exported to an outside model except by written agreement; the responding party must retain audit materials, including prompt-iteration logs and configuration records; and any genuinely proprietary prompt may be withheld but must be logged and made available for attorneys'-eyes-only inspection under the Protective Order.

Now that is a remarkable amount of transparency to sign up for. Let me pull out the pieces that matter most for us to consider.

- **Prompts are the provision people are going to focus on most.** You are disclosing all of them, and every change to them, in redline on a three-day clock. We've spent time this year on whether a litigant's own AI prompts are protected. [Warner v. Gilbarco](#) said yes — they are work product — and gave us the line that an AI platform is a “tool, not a person.” [Morgan v. V2X](#) agreed for a pro se litigant. So there is law that prompts can be protected work product. And here, the parties just agreed to hand them over. That is not a court overriding a work-product objection; it is a sophisticated set of parties on both sides of the case trading away that protection by stipulation — which is exactly the kind of protection I think would be hard to give up. It would be very difficult for many lawyers to sit down and think about doing this. I'm not saying it's a bad idea here — in the scope of what this litigation is, and the potential to move it forward faster and get to the merits, it probably has an incredible strategic angle. But it's going to be hard for a lot of folks to accept. The one escape hatch is the first place the two orders touch — the protocol and the protective order. A genuinely proprietary prompt can be withheld, but it has to be

logged and shown to opposing counsel's eyes only, under the protective order. So it's not a clean "this is privileged, you don't get to see it." It's "we're saying it's privileged, but we'll show it to the lawyers, and the lawyers can determine whether they get to argue it's privileged." So they're even giving up something on that level. I'm interested to see how this plays out.

- **The validation numbers also strike me as high, and that's worth pointing out.** A 95% confidence level is standard for this kind of sampling. But a plus-or-minus 2% margin of error is tight — a lot of TAR validation runs at a 5% margin, and tightening from 5% to 2% multiplies the number of documents you have to review in your validation sample by roughly six times, because the sample size grows with the square of the precision you demand. And an elusion expectation of 3% or lower is also pretty aggressive. The case law that blessed TAR — going back to *Da Silva Moore* and *Rio Tinto* — focused on whether the process was reasonable and cooperative, not on hitting a hard numeric target, and recall in the range of 70 to 80% has generally been treated as defensible. So what these parties agreed to sits at the very demanding end of anything the TAR world has required, and it's now being applied to AI. That's a lot to take on voluntarily — which is, again, my whole question about this order. I'm going to guess, given the sophistication of these parties, that they've used the systems they're engaging with and understand what they're taking on. But it's going to create a huge burden.
- **The quality-control requirement writes the failure modes of generative AI into a discovery protocol.** Appendix 4 requires checks for hallucination, over-summarization, or misclassification. A federal court order — and remember Rule 37(b), sanctions for violation of a court order — now expects your document review to include a documented check against the AI making things up. And it comes with a backstop I want every person on your team to hear: the order states that "legal counsel remains responsible for ensuring the accuracy and completeness of all productions." The AI does not absorb the lawyer's duty under Rule 26(g). You cannot point at the model when a production is wrong. These tools do not reduce the duty of care — they raise it.
- **On security**, all AI processing has to happen inside the secure review environment, with no document content exported to an outside model absent agreement. In plain terms, you cannot pipe the collection out to a public model. That one sentence tells you how firms and vendors can and cannot deploy these tools for this case, under this protocol. And if that's not how you want your case to work, don't copy the language of this protocol — adjust it for the way the technology works in your case. That is critical.
- **A quieter point on that "generally accepted" standard in Appendix 4.** The appendix makes you certify the system is generally accepted in the eDiscovery industry as reasonably reliable. It will not be a fight in this case, because the parties agreed their systems met that standard. But notice what the phrase is asking, because it will matter elsewhere: who decides what is generally accepted for a tool that may be six months old? We have decades of literature and case law validating TAR. We have nothing like that consensus for large language models used to make responsiveness calls, and the tools are changing faster than anyone can evaluate them. We don't know yet what "generally accepted" means for this technology, and that undefined standard is going to get tested somewhere.
- **Finally, layering — Section VII.6 — is the quiet provision that really matters here.** Layering is using two culling methods on the same set of documents — say, search terms and then TAR, or search terms and then AI. The document has to clear both filters. The problem is that each method independently misses some responsive

documents, and stacking them compounds the loss: the keyword pass strips out responsive documents that do not contain your terms, and the TAR or AI system never sees them. So the protocol makes you disclose the intent to layer before you do it, meet and confer, compare the hit counts with and without layering, and let the requesting party sample the excluded set. The risk is not the technology; it is the undisclosed, unvalidated combination of the two technologies. Transparency about process is what makes a review defensible, and this provision writes that down. We've not seen anything this comprehensive in addressing the layering concern in a protocol before. I think it's kind of genius.

How the Two Orders Work Hand in Hand

Now let's talk about how the two orders work hand in hand — the ESI Protocol and the Stipulated Protective Order. They are designed to interlock, and if you read one without the other, frankly, you're going to misunderstand both.

The difference you can't lose sight of is scope. The ESI Protocol — Document 74 — governs your entire review and production. Everything. Every document, whether or not it is ever designated under the protective order. The Protective Order — Document 75 — reaches only the material a party actually designates as protected, using the categories in the protective order. The protocol applies to everything; the protective order applies to a subset. That distinction drives the next point.

Section 13.4 of the protective order governs running Protected Material through generative AI tools. It says a party may use AI tools on that material only where the tool does not train on it and meets industry-standard security. And then it names names, stating that — quote — “the parties’ enterprise ChatGPT and Harvey accounts satisfy the foregoing requirements.” Read this carefully, because it is easy to get wrong. It does not mean enterprise ChatGPT and Harvey are the only two tools you can use, and it does not bless them for everyone. It means these parties reviewed the data processing agreements — the DPAs — behind their own enterprise ChatGPT and Harvey accounts, confirmed that those contracts prohibit training on the data and provide the required security, and put that finding in the order. This is a contract point, not a technology point. Plenty of other tools meet that same bar. The difference between a consumer account and an enterprise account is entirely a difference in the data terms — the enterprise DPA is what promises no training and delivers the security controls, and a consumer account does not. So if you adopt this language, the takeaway is not “use enterprise ChatGPT or Harvey.” It is: review your own tool's DPA, confirm it meets these requirements, and get that verification into the record — in your protocol if necessary.

This did not come from nowhere. It is the direct descendant of *Morgan v. V2X*, where Magistrate Judge Braswell wrote new AI-specific language into a protective order requiring that any AI platform carry a contractual prohibition on training, flow that prohibition down to third parties, provide a right to delete the data, and be documented in writing — which are, functionally, the terms of a DPA imposed from the bench. These parties took that same logic and stipulated to it in their protective order. They went one step further than Magistrate Judge Braswell required, by verifying and naming their tools. And note the contrast: in *Warner*, the Court said the parties did not have to disclose their tools. The question of tool disclosure has come up across these cases

and been handled differently — here, the parties chose to name them. Now, here's the over-designation problem I raised for you when we first looked at *Morgan*. Because the Section 13.4 restriction attaches only to Protected Material, its practical weight depends entirely on how much gets designated. And that creates a brand-new incentive to over-designate. If designating material controls what tools the other side can run it through, a producing party now has a reason to designate aggressively, to keep its data out of the opponent's unvetted tools. That is a different axis from the over-breadth problem in *Morgan* — there, the breadth was in the definition of the tool, where "any modern artificial intelligence platform" was broad enough to sweep in Westlaw, Relativity, and Microsoft Copilot. Here, the breadth is in the scope of the designation. But it is the same issue, and the guardrails are already in the order: Section 5.1 prohibits mass, indiscriminate, or routinized designations and threatens sanctions, and Section 6 gives the other side a challenge procedure. Those guardrails only work if the court enforces them and counsel actually challenges over-designation. Designation here — based on the issue we raised with *Morgan* — now carries AI-usage consequences it did not carry a few years ago.

Two more provisions in this protective order show how carefully it was tailored for an AI case. Section 2.19 makes Training Data attorneys' eyes only, so the datasets at the heart of the merits — Books3 and The Pile — are viewable by opposing counsel and experts under an attorneys'-eyes-only designation rather than walled off in the most restrictive tier. And Section 2.18 defines Source Code for this dispute specifically, expressly including — quote — "training pipeline code, data ingestion or filtering scripts, tokenization or preprocessing scripts, and model architecture files," while carving Training Data out of that definition. The ESI Protocol itself says source code is not covered and will get its own separate protocol — which tells you the real technical fight, over exactly how Cerebras-GPT was built, is still coming. And we'll be ready for it.

This whole section raises another theme we regularly discuss on Case of the Week. The level of complexity and thought put into this protocol means both sides sat down early and understood what they were going to want for this case. They had the OpenAI cases as a guide — multiple cases across the country dealing with exactly this issue of using copyrighted materials to develop an LLM. So there's a playbook for them to follow. But the point is: they followed it. They sat down and said, here are all the issues we're going to come across, let's put them all down and deal with them. And over-designation is one of the ones they included. Kudos to them. There is still a pretty big equity problem sitting underneath all of this. A protocol like this — disclosures, validation at a 95% confidence level, audit logs, prompt redlines on a three-day clock — is only achievable by a well-resourced party with sophisticated vendors. At the other end of the spectrum are the pro se litigants we saw in *Warner* and *Morgan*; we're seeing a wave of them using consumer AI tools to run their own cases, and that is already creating real problems for courts and for opposing counsel who have to sift through AI-generated filings and productions that nobody validated or governed at all. The same technology is producing two opposite problems at once: at the top, sophisticated parties voluntarily building elaborate governance like this protocol; at the bottom, parties using AI with none. This order is a template for the first group. It does nothing for the second. And the gap between them — what we refer to as access to justice — is widening. We flagged that gap when we talked about *Morgan*, because Magistrate Judge Braswell specifically raised it, and it is only getting sharper. Put these two orders on the map we've been building all year. In the [Warner](#), [Heppner](#), and [Morgan](#) episodes, the AI user was a litigant using a chatbot to think through a case, and the

question was whether the prompts were privileged. This order is a different region of the map entirely. Here, the AI user is counsel and their vendors, using AI as the engine that culls and produces documents, and the question is governance. Same technology, completely different discovery problem. What is striking about Cerebras is that the parties did not wait for anyone to rethink the rules. They wrote their own, by agreement, and the magistrate judge signed off on them.

Takeaways

All right, let's talk about our takeaways. Here's what to do with this.

Start with the thing that reaches every one of your cases, not just this one: your clients are already using AI, whether their businesses sanction it or not. That means litigators now have an affirmative obligation to understand how AI artifacts and evidence are created, and to plan for preserving them, producing them, and asking for them in discovery. This is where key evidence is going to live. And as the OpenAI cases have already shown us, prompts can be evidence of intent — what a person asked the model, and how they asked it, can go directly to state of mind. If you are not thinking about AI-generated material as a discoverable source of ESI in every matter, you need to start now. If you are in-house, you need to understand how AI is being used to create information within each of the business units in your company. If you are outside counsel, you need to sit down with your clients now and start understanding it — so that you're not spending two months, after your duty to preserve has already arisen, trying to figure it out, because then you'll have preservation problems.

Before you adopt a protocol like this one, decide whether you actually want its obligations. Nothing in Rules 26 or 34 requires prompt disclosure, validation reporting, or elusion testing. This is a stipulation, and stipulations bind you. The detail may buy predictability and fewer fights in this particular case — but you are expanding your own obligations and trading away protections. So make that trade on purpose, not by copying a form. If you are the responding party and you will use AI to review, build your AI Review Protocol before you deploy the tool, not after. Settle your hosting environment, your thresholds, your validation plan, and your prompt-disclosure position in advance, because you may be handing all of it to the other side.

If you are the requesting party, this order legitimizes real leverage. Demand the AI election, the validation metrics, the elusion testing, and the prompts, and ask whether they checked for hallucination or over-summarization.

On prompts, decide your position before you agree to anything. They are protectable work product under *Warner* and *Morgan*, and here the parties gave that up by stipulation. Now, the difference: in *Warner* and *Morgan*, the prompts were used to determine case strategy — they were the lawyers' perceptions. Here, the prompts are being used to determine the set of information to be reviewed, and then the subset to be produced. That process — how we determine the set to be reviewed, and then how we get to the subset to be produced — has always been privileged. This is changing that dynamic, in this case, based on the parties' stipulation.

On tools, know your data terms. The order verified the DPAs behind these parties' ChatGPT and Harvey accounts. Other tools will qualify to meet the obligations set out in this protocol, but you have to review your own tool's DPA and confirm it meets the no-training and security requirements. A consumer account will not clear the bar — just like Magistrate Judge Braswell said in *Morgan*. Get your verification into the record if necessary, so it's there and you don't have to prove it to the court in a later hearing.

Watch your designations. Because the AI-tool restriction attaches only to Protected Material, designate with discipline if you are producing, and challenge over-designation under Section 6 if you are receiving.

And do not overlook the unglamorous detail, because it is where cases get fought. The protocol handles short message data from Teams and Slack in units of no less than a 24-hour period, preserving emojis and threading. It lets a requesting party demand current versions of up to 100 hyperlinked documents within 14 days. It requires production of version history from collaboration tools. It sets the privilege log in Excel within 45 days and excludes party-to-outside-counsel communications after the complaint was filed on October 30, 2025. That's pretty standard. Sophisticated parties sweat these details for a reason — don't leave them out of your protocols.

There are so many takeaways from just reviewing this one protocol. And we include all of the ESI protocols and decisions in Minerva26 for exactly this reason — to let you see how other parties have handled the scope of drafting a protocol based on the complexity of their case. This case is highly complex, with a lot of information likely to be dealt with, and they're going to use sophisticated ways to determine the scope of that information and to review it — hence the detailed protocol we have here. Do you need that for every case? No. But you need to think about what you do need. So learn these things. Make educated decisions. Don't just pull a form off a shelf and decide it's going to work for your case because it worked for the last one. Understand your clients' sources of ESI. Understand how they're using AI — or, if they aren't, exclude it. Get in the game with these things. Don't just take what your client is telling you verbatim; test it. Understand what's out there, because you're the one signing the discovery responses. You're the one with the Rule 26(g) obligations, and AI tools don't change that. All right — that's our Case of the Week for this week. If you take nothing else, take this away: generative AI is no longer riding in under the TAR umbrella, and the parties writing these protocols right now are setting the template you will be handed next year. But a stipulation is a choice. Before you sign one that looks like this, be sure you want everything in it — because nothing in the rules made you agree to it, and everything in it will bind you.

That's also why we built Minerva26 — the discovery strategy platform that connects case law, rules, and real-world workflows — so teams can turn developments like this into a defensible plan instead of reinventing the protocol in every meet-and-confer. If it's about the discovery of ESI, it's covered in Minerva26.

If this was helpful, please share it with a colleague who is about to negotiate an ESI protocol in a case where either side might use AI, and post it on LinkedIn with your biggest takeaway — because the fastest way to level up discovery strategy is to have these conversations in public. And if you haven't already, please subscribe to the Meet and Confer podcast so you don't miss the next episode that changes what "reasonable" looks like.

Thanks for listening. We'll be back next time with another Case of the Week.