

You Can't Compel What You Didn't Negotiate: *Schulte v. LinkedIn* and the Limits of AI Transparency in Discovery

Meet and Confer Podcast · Kelly Twigger ·  [Listen](#) · Case in Minerva26: [Schulte v. LinkedIn Corp.](#), No. 22-cv-00237 · July 7, 2026 · Magistrate Judge Laurel Beeler · Northern District of California

Introduction

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 electronically stored information and turn it into something that you can use in your practice right away. My name is Kelly Twigger. I am the CEO and founder at Minerva26 and the principal at ESI Attorneys, and I have been a discovery strategist and practicing attorney for almost 30 years.

In my last episode of Case of the Week, I walked you through the protocol in [James v. Cerebras Systems](#), which was the first stipulated ESI protocol that I have seen that governs generative AI as its own category of document review, where two sophisticated parties simply agreed to hand over the prompts, the validation math, and the elusion rates in that generative AI review. And my question for that entire episode was: why would anyone voluntarily take on obligations that the rules never put on them?

This week we get the other half of the picture, and it is every bit as important as what we talked about last time, because this week a party did not agree. One side asked a federal court to compel the same kind of AI transparency the *Cerebras* parties gave away by stipulation, and the court said no.

I want to be precise about why this one is a first, because it is easy to overstate. This is the first decision I have seen where a party's use of generative AI to review documents is actually contested in litigation. But be careful about what is contested here, because it is not what you might assume. Nobody argued that LinkedIn could not use generative AI to review its documents. That fight did not happen, and honestly, I think that fight is effectively over. What the

plaintiffs contested here is narrower and far more interesting — not whether LinkedIn could use generative AI, but how much it had to disclose about how it used it. The plaintiffs did not try to stop the AI. They tried to see inside of it. And that's the frontier that we're on now — not whether you can use these tools in review, but what you owe the other side about how you used them.

I don't think we're going to see the same kind of analysis, or a [Rio Tinto](#)-esque discussion, and a federal judge signing off on the use of generative AI the way Magistrate Judge Andrew Peck did on TAR back in 2015. I think these tools are here to stay, and the question is what parties have to disclose about them. That's the issue that I want you to focus on for today.

What you can negotiate for in an ESI protocol and what a court will order over an opponent's objection are two very different things. We talked about that last time, and the gap between them is enormous. If you want visibility into how the other side is using AI to review documents, you need to build that into your protocol up front. You will not get it later by motion. This is one court saying that very emphatically.

Let's talk for a minute about why decisions like the one we're looking at today belong on your radar. There are more than 5,000 — we're actually approaching 6,000 for 2026 — discovery decisions issued in federal and state courts every year, and that has held up for the last five years. About 80% of those decisions on electronically stored information come from United States magistrate judges in the federal district courts, and less than 1% of civil cases go to trial, which means that these cases are won and lost in discovery. If you're not watching how courts are ruling on AI in that process, and making sure that the forms you are using are up to date, you are already behind.

As always, we'll provide links to the decisions in Minerva26, and you can use the Generative AI issue tag in the platform to be immediately notified of any new decisions on this topic if you're a subscriber.

Today's case comes from *Schulte v. LinkedIn Corporation*, out of the Northern District of California. The discovery order that we're talking about was signed by United States Magistrate Judge Laurel Beeler on June 30, 2026, and filed on July 1st. And I want you to notice the judge here, because it really matters. This is a magistrate judge in the Northern District of California — the exact same district that gave us the *Cerebras* protocol in the last episode — writing on the same technology, Relativity aiR, but in a completely different posture. There, the parties agreed. Here, they did not.

Now, before we get into the facts of this particular case, I want to put both of these protocols in time. And by both of them, I mean the *Cerebras* protocol and the ESI order that was entered here in the *Schulte* case that impacts what the judge ruled. Because the contrast between the two different protocols — when and how they were written — is a big part of the lesson for today, and you need it in your head before you hear how the Court ruled.

Timing is one of our key themes on Case of the Week, and it has a huge impact here. So let's start with *Schulte*, the case we're talking about today. LinkedIn was sued in January of 2022, almost a year before ChatGPT was released to the public. The parties negotiated their interim ESI order in January of 2023, and that order is exactly what you would expect from early 2023. It has a clause requiring disclosure of Technology Assisted Review, and not one word about generative AI, because generative AI document review was not a practice anyone was litigating when they wrote it. And then they just left it there. That's the ESI order that governs in 2026. Nobody went back to amend it as the technology changed. So when LinkedIn ran a Relativity aiR review in 2026, the fight got decided by a protocol drafted three years earlier, in a world that no longer exists.

Now look at who was on the other side of the *Cerebras* protocol, because it is the exact opposite. The lawyers who negotiated the meticulous AI-specific protocol in *Cerebras* were the same firm that is OpenAI's lead trial counsel in the copyright cases. And that firm drafted it after the OpenAI discovery wars, where OpenAI lost the fights that mattered. It was ordered to hand over some 20 million ChatGPT logs, its privacy objections were rejected, its appeal of the preservation order was turned away, and it is now staring at a sanctions motion over logs it was supposed to preserve. Those are the lawyers who sat down and wrote the *Cerebras* AI protocol. They had just lived through what happens when you try to fight AI discovery one motion at a time and lose. So in *Cerebras*, they did the opposite — they got ahead of it and built the disclosure structure into the protocol from the start.

Now, there are differences between *Cerebras*, because it's an AI company, and LinkedIn — including the fact that the ESI created during the relevant time period in *Cerebras* was generated using AI, and that didn't happen at LinkedIn. So there are big differences here in terms of the timing and the types of data that were created, and what needed to be handled in the protocol. The key takeaway is you've got to understand that as this technology is changing, you've got to look at what you have on file and make sure the newer technology is covered if you're going to want that data.

Now keep that contrast in mind. One set of lawyers wrote their protocol looking forward, having learned the hard way where AI discovery goes. The party here is stuck with a protocol that's looking backwards, drafted before that technology existed, and it was never updated. That's the point that I want to make over and over again. You have to watch how that technology is changing, and then go look at what you have on file — your ESI protocols, your standing orders, your form language — and ask whether it still fits the technology your client is actually using. Because the order you signed in 2023 is going to decide your fight in 2026, whether or not it was written for it.

Now let's set up what the case is about and walk through the three rulings, because all three of them are useful and they're easy to run together.

The Facts

The underlying case here is an antitrust class action. The plaintiffs, who are LinkedIn Premium subscribers, allege that LinkedIn monopolized and attempted to monopolize in violation of Section 2 of the Sherman Act, and that the monopoly led it to overcharge Premium subscribers. The two categories of alleged conduct: first, that LinkedIn gave potential rivals access to its private user data through APIs on the condition that they not compete with LinkedIn; and second, that LinkedIn integrated its user data with parent company Microsoft's Azure cloud product, tying up scarce hardware and driving up prices. You don't need the antitrust theory to be able to use this order, but hold on to the fact that the defendant here is a Microsoft-owned data business with an enormous volume of ESI, because that volume is doing the real work in the Court's proportionality analysis.

There are three discovery letter briefs in front of the Court, and the Court denied all three of them — all three brought by plaintiffs. That is worth saying at the top, because the plaintiffs, the requesting party, went zero for three. The three disputes are over: first, LinkedIn's use of the generative AI tool Relativity aiR to review documents and its use of search terms to cull the population before that review; second, plaintiffs' motion to add an in-house counsel as a document custodian; and third, plaintiffs' motion to compel text messages from 19 custodians. I'm going to take them in order, because the AI ruling is the one that connects to everything that we've been building this year so far.

Before I do, let's talk about one piece of technology specificity, because it matters, and because the Court is kind of loose with it. Relativity aiR is Relativity's generative AI review product. A lot of you are very familiar with aiR. It uses large language models to make responsiveness and privilege calls on documents, and it produces a rationale for each call. That is not the same thing as traditional Technology Assisted Review. Traditional TAR — what we have well over a decade of case law on, going back to *Da Silva Moore* in 2012, and *Rio Tinto* in 2015 that I mentioned — trains a classifier on a human-coded seed set and then ranks documents by likely responsiveness. Generative AI review, like aiR, does not necessarily use a seed set at all. It applies an LLM against your review criteria. LinkedIn disclosed that it did not use a seed or training set, that aiR was making the final responsiveness calls, and that human review was limited to quality-control sampling of each responsiveness category. So this is a genuine AI review, making final calls, with humans checking samples rather than every document. Keep that in mind, because the Court is about to treat it as though it were ordinary TAR, because that's what the interim ESI order addresses — the parties had that in place. It does not — the ESI order does not address generative AI review at all.

The Analysis

1. The AI ruling — search-string pre-culling and the demand for aiR metrics (ECF No. 186)

All right, so what is the Court's ruling on these three motions? Let's start with the AI ruling and the plaintiffs' arguments about search-string pre-culling and the demand for aiR metrics. There are actually two fights bundled inside of it, so let's talk about each one separately.

This is the ruling that makes the decision a first, because it's the first time I have seen a court rule on a challenge to how a party uses generative AI in review, rather than whether or not it could use it at all. So that's not the fight, and I don't think that's ever going to be a fight. We're talking strictly here about how a party used generative AI in review. And keep this point in mind too: this fight is happening after LinkedIn already put more than 200,000 documents through the process that I'm about to describe. So we're not at the beginning of the process. We're halfway through it.

The first fight is about pre-culling. LinkedIn gave the plaintiffs 25 search strings and told them it would run those strings against the custodial documents first, and then feed the results — the surviving population — into Relativity aiR to filter for responsiveness. So: the whole mass of custodial data, search terms applied against it, the results of that search-term culling then fed into aiR. The plaintiffs said, no, no, no, no — you cannot use search strings to pre-cull before the AI runs. Plaintiffs' theory was that search strings artificially reduce the population that aiR ever sees, so responsive documents that do not happen to contain a search term might get stripped out, and the AI would never see them. They asked the Court to prohibit the search-string pre-cull and compel LinkedIn to run aiR across all of the custodial information. That would be a substantial amount of data and cost added to review.

Now, if that argument sounds familiar, it mirrors the layering problem that we spent time on in *Cerebras* last episode. Layering is stacking two culling methods on the same set of documents, so a document has to clear both filters to survive. And the risk is exactly what the plaintiffs described — each method independently misses some responsive material, and stacking them compounds the loss. In *Cerebras*, the parties handled that risk by agreeing to disclose the intent to layer, meet and confer, compare hit counts with and without layering from search terms, and let the requesting party sample the excluded set. That was the stipulated answer to the problem we have here.

Here's the contested answer. Magistrate Judge Beeler denied the plaintiffs' request to stop LinkedIn from layering. Her reasoning is the practical heart of this order — this decision we're talking about — so listen closely. The plaintiffs argued in the abstract that pre-culling is improper, but they never showed that LinkedIn's 25 search strings were actually deficient. This is one of our themes here on Case of the Week, and you've heard me say this over and over: you have to have factual support for your arguments. As the Court put it, if the plaintiffs had

shown the strings were too narrow, their concern about pre-culling might be warranted. But they did not make that argument, and they did not even raise concerns about the strings when LinkedIn disclosed them back on May 15th.

The Court then held, quite flatly, that using search terms to pre-cull documents before handing them to a technology review platform satisfies the reasonableness and proportionality requirements of Rule 26(b) and Rule 34(b)(2). And she cited [*Livingston v. City of Chicago*](#) and [*In re Biomet M2a Magnum Hip Implant Products Liability Litigation*](#). Both of those cases are in Minerva26. Pre-culling is not improper. It is normal. It is the process we have followed for two decades — search terms are what we have been using — and there's no reason for a court to otherwise consider a problem with them unless the parties can show substantial issues. Here, the plaintiffs agreed to the search strings, so there were no issues.

Then the Court dropped what I call the proportionality hammer. The files of just two LinkedIn custodians totaled roughly 800 gigabytes. Across 19 custodians, running aiR on everything without any pre-cull would mean feeding multiple terabytes into the tool, at significant cost in processing, hosting, and human review. So the court denied the request and instead ordered the parties to meet and confer within 21 days about the search strings themselves, with leave for the plaintiffs to come back if they think specific adjustments are warranted and cannot get agreement.

Now, that's a huge win for the plaintiffs. We've seen judges who've told parties "you snooze, you lose," and allowed the producing party to proceed with no ability to revisit what they've already agreed to on search terms. And that's what's key here. The Court did not say you can never challenge a pre-cull. It said: challenge the search strings, with a specific showing, at the time they are disclosed — not the concept of pre-culling in the abstract, months later. The objection has to be particularized, and it has to be timely. A generalized complaint that AI should see everything is a loser.

Now, practically speaking, this means that once documents are produced to the plaintiffs, they can go through them and identify things that are missing — concepts, or specific testimony they have evidence should exist — that aren't in those documents. And they can then take that factual evidence back to the Court and make an argument for something new. They'll have to meet and confer with LinkedIn before they can do that, but that is the process they can now follow. The judge gave them that out, which is, as I mentioned, more than a lot of judges will do.

The second fight was over the demand for metrics from aiR. This is the part that ties directly to *Cerebras*. The plaintiffs moved to compel LinkedIn to disclose a set of metrics about its use of Relativity aiR. Specifically — and in my view, ironically, given the parties' stipulation in *Cerebras* — the plaintiffs sought elusion estimates, the document error rate, and the number of human reviewers validating the AI's predictions. Look at that list, because it's nearly the exact set of disclosures the *Cerebras* parties agreed to hand over in Appendix 4. It's the same trio. The *Schulte* plaintiffs asked a court to order what the *Cerebras* parties volunteered. Now, plaintiffs

didn't go as far as what *Cerebras* included and ask for the identities and backgrounds of the reviewers, but it didn't matter. The Court said no.

The Court said that plaintiffs were seeking what we call discovery on discovery. Discovery on discovery is when a party wants discovery about how the other side is conducting its search and collection, as opposed to discovery of the underlying facts. And the rule, straight from [Taylor v. Google LLC](#), is that discovery on discovery is not something courts grant, because it is typically not relevant to the merits and rarely proportional. You can get it only if you demonstrate a specific deficiency in the other side's production — mere speculation is not going to cut it.

Magistrate Judge Beeler denied the metrics request for two reasons. First, LinkedIn had already satisfied its obligations under the parties' interim ESI order. As I mentioned, the parties had an ESI protocol or order, whichever phrase you prefer — the Court calls it an order here — and the obligation in that order is what decides this motion. Paragraph 5(a) of the interim ESI order required a producing party to disclose if it intends to use Technology Assisted Review to filter out non-responsive documents. That's it. Disclose that you are using it. LinkedIn disclosed that it was using aiR, which the court characterized as, quote, "a form of Technology Assisted Review" to filter non-responsive documents, and it went further and answered follow-up questions. The Court held that those disclosures more than satisfied the ESI order.

Second, the only deficiency that the plaintiffs could point to was that the target population came out to 204,444 documents, and a raw document count, without more, is not a specific showing of a deficient production. Now, it does seem like not very many documents, given the scope of what we're talking about from an antitrust perspective, but it's not sufficient just to make the argument that that's not very many documents. Because the argument was based on speculation and not deficiency, the court denied it.

Now focus on the move that the Court made, because it's the most important thing in this order, and it runs in the opposite direction from *Cerebras*. LinkedIn used a generative AI tool making final responsiveness calls with no seed set, and Magistrate Judge Beeler folded it under the TAR umbrella for purposes of the ESI order's TAR-disclosure clause. And here's a wrinkle worth pointing out: this decision actually came first, before the *Cerebras* protocol. Magistrate Judge Beeler signed this order on June 30th. The *Cerebras* protocol I walked you through last episode was entered on July 7th, a week later. I covered *Cerebras* first because the stipulated AI protocol was the headline, but chronologically, the contested fight came before the tidy agreement. So inside a single week, in the same district, on the same technology, two magistrate judges pointed in opposite directions. One treated generative AI review as just a form of Technology Assisted Review and applied a decade-old TAR standard to it. The other deliberately broke AI out of the TAR category and gave it its own separate rulebook.

Now, one qualification to that statement: the *Cerebras* protocol that contained specific AI-related provisions was stipulated to by the parties. The judge didn't write it for them, didn't require it. What he did do was give them permission to work outside of the Northern District of California's standard protocol, which has been on file since 2015 and is wildly out of date. So the Court

didn't order it there. But here, the Court denied it because of the existing ESI order that the parties had entered into.

Now, that tension — is generative AI review just TAR by another name, or is it a different animal that needs its own rules — is now live and unresolved, and you're going to see it fought about. I think when you really look at the timing of when this ESI order was entered, in 2023, and that the parties were thinking about the technologies that existed at that time, the Court was really just trying to fashion something that was going to work for everyone going forward. I don't think that particular language from the Court, about it being just a form of TAR, is going to really come back to bite anyone. I don't think that's going to be a reasonable argument that's going to be relied on going forward. So keep that in mind — you've got to keep in perspective the timing and the language of the ESI order that the judge was dealing with here.

And now the timing that I talked to you about at the very top of the episode pays off. Remember, the parties' ESI order was from January of 2023, and its Paragraph 5(a) required disclosure only of Technology Assisted Review. Because there was no generative AI provision, when aiR showed up, that TAR clause was the only hook that Magistrate Judge Beeler had. That's not any kind of criticism — it's just the point I was making: here's the specific language of the order, and here's why she relied on it.

Let's go back to the bottom line here on the ruling for this first motion. Under a standard ESI order's TAR clause, disclosing that you have used TAR or generative AI to filter documents is generally enough, if that's what your order calls for. You're not automatically entitled to the elusion rate, the error rate, or the reviewer headcount. The discovery-on-discovery doctrine squarely stands in front of that request, and you only get past it with a specific, non-speculative showing that the production itself is deficient. Which means — and this is the connection to the last episode — if those metrics matter to you, the place to get them is in the protocol, negotiated at the outset, the way the *Cerebras* parties did. It is not a motion to compel 18 months in. Given that the judge didn't order it here, it seems unlikely that she would have forced a protocol that did, unless the parties agreed. And since LinkedIn is fighting it, that seems doubtful.

2. Adding an in-house lawyer as a custodian (ECF No. 188)

All right, let's talk about the second motion. It's shorter, but genuinely useful, because the add-a-custodian fight comes up in almost every matter. The plaintiffs moved to add a LinkedIn in-house attorney named Jane Slater as a document custodian, and they wanted to collect, search, and produce her internal and external communications, subject to a privilege review.

Their theory was specific: that Slater was one of three lawyers who negotiated LinkedIn's private API agreements, and she participated in a set of negotiations where the only other LinkedIn person in the room was a business executive who left the company in 2021 and whose files were not preserved. So the plaintiffs said Slater is the only person who can fill the gap left by the deleted files.

The standard that the Court applied is the one to note here. To add a custodian, the requesting party has to show that the proposed custodian has, quote, "uniquely relevant information that is not available from the sources already designated," close quote, and it has to make a particularized showing of the specific gaps that only this custodian can fill. That comes out of the court's own prior order in the case and [*In re Facebook, Inc. Consumer Privacy User Profile Litigation*](#). Keep those standards in mind: uniquely relevant, particularized, specific gaps. That's a high bar, and it's deliberately high, because custodians are expensive and duplicative custodians are pure waste.

The plaintiffs didn't meet the bar here. LinkedIn was already producing from five current and former business development leaders responsible for those API relationships, two of whom were the direct supervisors of the custodian whose documents were not preserved during the relevant time period. So the gap was largely covered by the existing custodial framework. And critically, Slater is a lawyer. LinkedIn represented that her role in the negotiations was to give legal advice on the terms of the API agreements, which means that her internal communications are likely to be overwhelmingly privileged. So the Court found this supposed gap speculative, and it rested on two unproven assumptions — that Slater and the former executive had internal deliberations with no other custodian present, and that those deliberations were in a non-privileged, quote, "business negotiation capacity," close quote. No evidence in the record supported either one.

So the burden of collecting, reviewing, and logging a lawyer's files to find the handful of arguably non-privileged messages not already captured elsewhere outweighed the benefit. The court denied the motion, but it did note that LinkedIn had already agreed to produce Slater's external communications with API partners, where the non-privileged substance actually lives.

And I think that's the key point here. I think if LinkedIn hadn't agreed to produce the external communications with API partners from the lawyer, it likely might have had to produce more from that custodian, because it's possible plaintiffs may have had to come back on a subsequent motion and give additional facts to support it. But come on — we've produced a lot of lawyers' custodial files in litigation for various reasons, and when they're involved in business negotiations like this one, which are the crux of the case, their files are relevant. So I think that getting the external communications will make a big difference. If the external communications point the plaintiffs to things that would be in another set of documents for that lawyer that are preserved, the plaintiffs may have another argument to come back.

Now, here's the practice point on this issue. When you move to add a custodian, and especially when that custodian is a lawyer, a generalized "they were involved and might have unique material" is going to lose. You need a particularized showing of specific, non-privileged, non-duplicative material that only that person holds and that the existing custodians do not. Build that record before you file. Here, the plaintiffs can do that once they get the information, and it will really only depend on whether the court allows this particular issue to be reopened.

3. The text-message motion (ECF No. 190)

Now, the third motion is the most familiar fight in modern discovery, and the ruling is a clean illustration of how a well-drafted ESI order controls the outcome. The plaintiffs moved to compel LinkedIn to preserve, collect, search, and produce responsive text messages — iMessage, SMS, and app-based messages like WhatsApp and Signal — from all 19 custodians, on both personal and work-issued devices. Their argument was that LinkedIn executives routinely use text for business and to discuss topics relevant to the case.

But apparently, the counsel who negotiated the ESI order didn't realize that early enough, because the parties' interim ESI order, at Paragraph 4(b), carved text messages out. We've seen this before, where parties have agreed to carve text messages out and realize later that the evidence was in the text messages. That section provided that text messages and app-based messaging data need not be preserved, absent a showing of good cause by the requesting party.

So this motion was not decided on a blank slate. The plaintiffs had actually agreed to a provision where they bore the burden to show good cause, which at minimum meant showing the texts were relevant, proportional, and not cumulative. And they couldn't carry that burden. The plaintiffs' evidence was mostly emails referring to texts, and when the Court read those emails in context, the referenced texts looked logistical or cumulative, not substantive. One email had an employee asking a colleague to text if they needed anything urgent while traveling. Another had an employee giving a, "overview of our progress," and saying "text me for anything urgent," where the overview itself was already in the produced email. The plaintiffs' strongest example was a produced screenshot of texts about topics for an upcoming QBR — a quarterly business review — but LinkedIn had already produced the final prep documents for that QBR, so any related texts were likely duplicative. As the Court put it, the plaintiffs could not point to a single substantive business discussion occurring by text that was not already captured somewhere else. There was no good cause.

Now, I don't know where the parties are in discovery, or whether any depositions have been taken at this point. But asking specific custodians how they communicated, and whether they communicated substantively by text, will likely inform this. Again, if you go back to when this protocol was written three years ago, it wasn't unusual for the parties to negotiate away text, because we didn't rely on them the same way that we do here. Now, 2023 was post-COVID — use of text messaging did skyrocket during COVID — so maybe there's a little bit more to that. But again, it wasn't as much an unusual thing then as it is today, in 2026. And here, the Court found that there was no good cause to provide text messages for those 19 custodians.

The Court again looked at proportionality. Nineteen custodians, multiple messaging platforms, personal and work devices — all of that added up to being disproportionate to the needs of the case. The Court flagged the disruption to the private lives of each custodian that collecting

personal-device texts would require, and even if the plaintiffs had shown some marginal relevance, the burden overwhelmed it.

There's also a nice good-faith beat that the Court notes here that's worth stealing. LinkedIn had identified three repositories as most likely to hold responsive, non-duplicative material: M365 mailboxes, Google Drive, and Slack. And the Court noted that LinkedIn voluntarily included Slack, which is itself a category of text messaging excluded under the ESI order. That voluntary inclusion helped the Court credit LinkedIn's search as reasonable and made in good faith. And the court closed by invoking the *Sedona Principles*, Third Edition — that the responding party is ordinarily best positioned to decide what is relevant and proportional in its own data.

Now, we all know that the things people say in text messages are far and away the best evidence, even more so than Slack or any other kind of chat, because nobody has a filter when it comes to text messages. So you've got to pay attention to what's happening with mobile device data, and specifically text and instant messages, these days. That's where the good data is. Email is still crucial — we still use it — but it's not where all the great evidence is these days. In most types of cases, we're looking at those text messages, that kind of social data that really has the unfiltered, off-the-cuff things that tend to show the most intent and what actually happened in a given situation.

Now, the lesson here in the *Schulte* case is the same one that decided the ruling on the first motion, and it's the lesson of the whole episode. The ESI order did the work here. Paragraph 4(b) carved out texts and put the good-cause burden on the requester. Paragraph 5(a) required only disclosure of TAR use. Two provisions, negotiated at the front of the case, decided two of the three motions before Magistrate Judge Beeler ever even reached the merits of them. The protocol that you draft at the outset is the protocol that you litigate under 18 months later. Make sure it stays up to date.

Where this sits in the AI-in-discovery line

Now, where does this ruling sit in our AI Discovery Case Law series that we've been building on Case of the Week over the last few months? Let's put this on the map of what we've been building.

In [Warner v. Gilbarco](#) and [Morgan v. V2X](#), the AI user was a litigant, and the question was whether that litigant's own prompts were protected work product. *Warner* gave us the line that an AI platform is a "tool, not a person," and both of those cases protected the prompts. In [United States v. Heppner](#), that was the exception — use the tool outside the lawyer relationship, and you lose the protection. [Conservation Law Foundation v. Shell](#): expert prompts were discoverable. And in the last episode, [James v. Cerebras](#): counsel and their vendors using AI as the engine that culls and produces documents, with the parties agreeing to sweeping transparency in a stipulated protocol.

Schulte is the contested counterpart to *Cerebras*, and that's why it belongs right next to it. It's the same posture — AI is the review engine, run by counsel and vendors. But where the *Cerebras* parties handed over the elusion rate, the error rate, the reviewer counts, and prompts by stipulation, the *Schulte* plaintiffs asked a court to compel essentially the same list — and ran straight into the discovery-on-discovery wall. What one set of parties gave away voluntarily, another set of parties could not pry loose by motion.

So here is the one sharp question to carry into your next matter, and to analyze your existing matters on: if the other side is using generative AI to decide what you get, how will you know whether it worked, and where does that right come from? Because after *Schulte*, the answer is not "I'll move to compel the metrics." That decision is against you. The answer is: I negotiated the disclosure I need into the ESI protocol at the front of the case, when I still had the leverage to do it. Now, maybe the next decision we'll see is the parties fighting over whether those things should be in a protocol, and what the basis is for one party asking for them under the Federal Rules of Civil Procedure or the state equivalent.

Takeaways — what to do this week

What do these decisions mean for litigators, and how should you be thinking about your own matters this week?

- If you want visibility into an opponent's AI or TAR review — elusion rates, error rates, validation, reviewer counts, the prompts — you have to negotiate it into the ESI protocol at the outset. *Schulte* tells you that the discovery-on-discovery doctrine will not give it to you later. *Cerebras* tells you exactly what those provisions can look like. Use one to build the other.
- If you are challenging an opponent's search terms or a pre-cull before AI review, do not argue that pre-culling is improper in the abstract. That argument lost here. Make a particularized showing that the specific search strings are deficient or too narrow, and raise it when the strings are disclosed, not months later. Timeliness and specificity are what the court is looking for.
- If you are the responding party, know what your ESI order actually requires. Under a standard TAR-disclosure clause, disclosing that you use TAR or generative AI to filter non-responsive documents is generally enough. You are not automatically obligated to provide elusion estimates, error rates, or reviewer headcounts. But don't overread that — the protection comes from the discovery-on-discovery doctrine plus a clean disclosure, not from stonewalling.
- If you're moving to add an additional custodian, especially an in-house lawyer, make a particularized showing of specific, non-privileged, non-duplicative material that only that person holds. A speculative gap plus a likely-privileged file set loses every single time.

- On text messages: if your ESI order carves them out absent good cause, "executives text about business" is not good cause. You need a substantive business discussion that happened by text, or some testimony that they communicated by text, that is not captured anywhere else. Voluntarily including a source like Slack helped LinkedIn establish good faith — reasonableness is a record that you build.
- And underneath all of the above: draft the interim ESI order like it will decide your case, because it will. Two provisions in this order — the TAR-disclosure clause and the text-message carve-out — resolved two of the three motions on their own terms. The leverage is at the drafting table, not at the motion-to-compel stage.

And it's not only the orders you are drafting today. Go pull the protocols and standing orders you already have on active matters, and read them against the technology your client is actually using right now. LinkedIn is litigating under a January 2023 order that predates generative AI review. If you have cases running on protocols from 2022 or 2023, they almost certainly say nothing about AI, and the time to amend them is before the fight, not in the middle of it.

Conclusion

That's our Case of the Week. And if you take nothing else away, take this: everything the *Cerebras* parties gave away by agreement is exactly what the *Schulte* plaintiffs could not get by court order. The transparency you want from the other side's AI is a thing you negotiate for up front, not a thing you win on a motion after the review is done. The protocol is the leverage. Draft it that way.

You can follow Minerva26 for all the updates on the case law and rulings you need to know about using the Generative AI issue tag. The whole run — *Warner*, *Heppner*, *Morgan*, *Conservation Law Foundation*, *James v. Cerebras*, and now *Schulte* — is organized under that tag, so you can watch the contested rulings and the stipulated protocols line up against each other.

That's our Case of the Week for this week. I'm your host, Kelly Twigger, and I invite you to subscribe to our blog at Minerva26 to follow our Case of the Week updates, as well as to subscribe to the Meet and Confer podcast. If you value this content and want us to keep providing it, please share it with your colleagues on social media so that we can continue to reach as many folks as possible. If it's about the discovery of ESI, it's covered in Minerva26, your discovery intelligence platform. Thanks for joining me. See you next time.