

Verifiable Sampling for AI Prompts and Outputs

Meet and Confer Podcast · Kelly Twigger ·  [Listen](#) · Case in Minerva26: *Disney Enter., Inc. v. Midjourney, Inc.* · Decisions: [Order re Statistical Sampling Protocol for Production of Private Subscriber Prompts and Outputs](#), Dkt. No. 152, entered August 10, 2026; [Order re Production of Public Subscriber Prompts and Outputs](#), Dkt. No. 166, stipulated August 10, 2026, approved August 12, 2026, filed August 14, 2026 · Lead Case No. 2:25-cv-05275-JAK (AJRx), consolidated with No. 2:25-cv-08376-JAK (Ex) · Court: Central District of California · United States Magistrate Judge A. Joel Richlin

The Discovery Problem With No Stopping Point

Hi, and thanks for joining me this week. I want to start this week's episode with a problem that's going to start coming up as we represent generative AI companies and we're looking for generative AI outputs as evidence in cases. Suppose the evidence in your matter is not documents in a file share, or email, or text messages. Suppose it is everything a piece of software ever produced for its users. Every request somebody typed in, every result the system handed back, sitting in a database that was never built with discovery in mind, as happens with most all of the tools that we use. There is no natural stopping point in a data set like that. You cannot collect all of it, you cannot review all of it, and the other side is unlikely to accept whatever slice you decide to hand over unless you discuss it and agree on a process.

The answer when the volume gets to that scale is sampling. And we've talked about sampling multiple times on the Case of the Week. Today is going to be a chance to see where counsel really came together and came up with a sampling process based on science and data, and that is where we're going to have to be going when it comes to eDiscovery.

Now, sampling takes a piece of the population, looks at it closely, and reasons from the piece to the whole. Courts have been comfortable with the sampling process for years, but sampling does raise an obvious question that most protocols counsel put together never answer. How does the other side know that your sample is honest? You picked it, you ran the search, you decided what came out.

What the two orders we're going to discuss today do, and the reason that I picked them, is that they answer that question in writing. They agree on a sample size that comes out of statistics rather than negotiation, a method of drawing the sample that anybody can

rerun and check, a signed certification that the method was followed, and a list of every record in the population the sample came from so that the other side can verify the draw themselves. That last piece, providing a record of the entire responsive population, that's not something that I've seen before. This is the first time I've seen it.

Welcome to the Case of the Week

Welcome back to the Meet and Confer podcast. My name is Kelly Twigger. I'm the CEO and founder at Minerva26 and the Principal at ESI Attorneys. I've been a discovery strategist and practicing attorney for just about 30 years. This series on our podcast is called the Case of the Week, and this week's decisions come to us from the *Disney v. Midjourney* case pending in federal court in the Central District of California.

Now, this is the second time I'm covering a stipulated order from the parties on the Case of the Week, and here's why. The decisions and protocols coming down right now are deciding how you write your next protocol in your next case, and whether you need to go back and change the one that you're already operating under today. This episode and the ones before it are the proof that you need. Nearly every ruling we have covered this year on artificial intelligence has turned up, in one form or another, in the language sophisticated parties are now negotiating into their own agreements.

Now, don't get lost in the sophisticated parties angle if you're not representing Fortune 50 or Fortune 100 companies. The reality in ESI is that the more you negotiate for the smaller level cases you have, the better evidence you can get at a lower cost point. So don't get caught up in that. Watch these decisions and you get to draft from where the law is now, and it's changing every day. Do not watch them, and you're going to risk being handed someone else's protocol and find out what it costs you by not getting the evidence your client needs somewhere down the line.

Finding Stipulated Orders and Protocols in Minerva26

We recently made an update to the Minerva26 platform that allows you to search directly for the type of stipulated orders that we're going to talk about today. And you can also remove them from your search. Now, using that functionality plus issue tags like generative AI, prompts, ESI protocol, or protective order will drive you directly to those orders. If you'd like to see how that functionality works, or if you haven't seen it demonstrated for you or your team that are already on the platform, please reach out to us at support@minerva26.com and we'll get you set up.

The Two Orders in *Disney v. Midjourney* (Dkt. Nos. 152 and 166)

All right, our case today that we're going to talk about is *Disney Enterprises v. Midjourney*. That's the lead case. It's in the Central District of California, and it was consolidated with the Warner Bros. case as well.

The first of the two orders we're going to talk about is the Order re Statistical Sampling Protocol for Production of Private Subscriber Prompts and Outputs, which is docket number 152. That was entered on August 10, 2026. The second is the Order re Production of Public Subscriber Prompts and Outputs, docket number 166, which the parties stipulated to on August 10th. Now, United States Magistrate Judge A. Joel Richlin approved that order on August 12th, and it was filed on August 14th. Both links to those orders are going to be in the show notes. If you're a subscriber on Minerva26, you'll also see all of the other additional decisions in the *Disney v. Midjourney* case that are attached in the linked decisions on the left-hand side of the panel.

Now, let's notice the shape of these before we go any further, because it really matters for how you read them. As we saw in the [Cerebras](#) protocol case, nobody fought about either one of these orders. They are joint stipulations that the Court approved. That makes this the second agreed-upon AI protocol that we've covered this year, after [James v. Cerebras](#). And it's worth asking why the parties in the biggest generative AI copyright case in the country worked together to solve the problem of how to get electronic evidence, rather than asking the Court, who has little knowledge of the technology, to decide it.

Where This Sits in the AI-in-Discovery Line

Now let's talk for just a second about where this sits in the AI and discovery line of decisions. We've been building a line of cases on artificial intelligence and discovery for the last two years, and every one of them so far has asked the same question from a different angle. When someone uses an AI tool, what does the other side get to see about it? That's [Warner](#), [Morgan](#), [Heppner](#), [Tate](#), [Conservation Law Foundation v. Shell](#), the [Cerebras](#) protocol, and [Schulte](#). And if you want details on those, you can go back to those episodes, which will be linked.

Today is a little different. Here, the AI output is not a tool that anybody used to prepare the case. It is the evidence itself. The studios are suing over what Midjourney's model generated, which means the prompts subscribers typed and the images and videos that came back are the documents at the center of the case. Same tree, new branch, and it raises a question that none of the earlier cases had to answer. When the AI output is the evidence and there is an effectively unlimited amount of it, how does anybody produce it?

The Underlying Case, and the Two Features That Shape Production

All right, what is the underlying case about here? In June of 2025, Disney and Universal sued Midjourney in the Central District of California. Warner Bros. filed its own case later and the two were consolidated. The plaintiff group is essentially the entire American character library: Disney Enterprises, Marvel Characters, MVL Film Finance, Lucasfilm, Twentieth Century Fox, Universal City Studios Productions, DreamWorks Animation, Warner Bros. Entertainment, DC Comics, Turner Entertainment, Hanna-Barbera, and the Cartoon Network. The claim is about copyright infringement based on both how the model was trained and what the subscribers were able to generate with it.

Now, Midjourney is a text-to-image and text-to-video service. A subscriber types in a prompt and the model returns image or video files, and the whole exchange is stored as what the platform calls a job. If the studios want to prove that subscribers were generating Darth Vader, or my personal favorite Han Solo, and Scooby-Doo or the Minions, the proof is in those individual jobs.

Now, two features of the platform shape the entire production problem. The first one is volume. This is a consumer-based service with a very large subscriber base generating continuously, so the population of responsive jobs is enormous, sort of akin to the prompt issue that we saw in the OpenAI cases. The second is that Midjourney has a feature called stealth mode, which lets a subscriber generate privately rather than publicly. That single product feature splits the discovery population in half, and the parties dealt with those halves in two separate orders. Docket 152 handles the private, stealth mode material. Docket 166 handles the public material and borrows all of its machinery, basically, from 152. So same process, but applied to a different set.

Both orders respond to the plaintiffs' Requests for Production Nos. 43 through 45, served under Rule 34, and both trace back to a minute order that the Court entered on April 16 at Docket 61.

Character Buckets, and the Coordination Behind Them

Now, one more definition you need to understand the orders. The parties defined a Character Bucket as a group of keywords associated with a specific character or group of characters. Think about the keywords that you might use to identify Iron Man, or Captain Marvel, or Indiana Jones, or SpongeBob SquarePants. The buckets are listed in two exhibits, and everything in both orders is organized by bucket.

Now, that's worth pausing on, because it's a design choice that counsel made here, and it's a smart one. They did not sample the corpus as a whole. They sampled per character, because each character is a separate asserted work and a separate infringement question. Your buckets have to line up with the questions the case turns

on, or a sample cannot answer anything.

Now, think about for a minute that whole range of studios that I gave you and the sheer number of characters that we are talking about. The level of coordination here to create these buckets means that you had to have counsel representing each one of those individual studios and thinking about each one of those individual characters to create those buckets. That is not a simple process.

Now, let's walk through Docket 152, because this is the part you're actually going to want to leverage in your matters.

Why the Sample Size Is 385

This is the engine, how the sample actually gets drawn. Now, the sample size that the parties chose was 385 per Character Bucket. Now, that's not a number that two lawyers split the difference on. Section 8 states the parties' agreement that a properly drawn random sample of 385 from a population of at least 385 is statistically sufficient to be representative of the underlying universe, with a margin of error of approximately plus or minus five percentage points at the 95% confidence level. If a bucket holds fewer than 385 jobs following the search terms, Midjourney produces all of them instead of sampling.

Now, let me translate that, because that formula behind it is pretty ugly and you don't need it to use this. Here's what a sample of 385 buys you. Look at 385 jobs out of a bucket, and whatever share of those 385 has the quality you're measuring, the whole bucket is within about five points of that share. Say the experts score the 385 and 70% depict the character. That bucket is around 70%, and the honest range is 65 to 75. The 95% confidence level is how reliable that claim is. Do this over and over, and you land inside the range 19 times out of 20.

Now, here's the part that surprises everybody the first time in this process. The size of the bucket barely matters. 385 is not a percentage of the population. It is a fixed number, and past a few thousand records, it barely moves. The same 385 covers a bucket of 50,000 jobs and a bucket of 50 million. That is not the parties cutting a corner. That is how sampling works.

But there are two limits on what that number can do, and you need to understand both.

Where Sampling Breaks Down

The first is the difference between a percentage and a count. Five points of error on a percentage becomes 5% of the population once you multiply it out. 70% of 50,000 jobs is 35,000, give or take 2,500. 70% of 50 million jobs is 35 million, give or take 2.5 million. Identical sample, identical quality of estimate on the percentage, and a range on the actual number that is a thousand times wider. If the count of infringing outputs ends

up driving damages, I think that's where the fight will be.

The second limit has nothing to do with the population size, and I think that's where the real exposure is here. It is about how rare the thing you are counting is. 385 works well for estimating something common in the bucket. It doesn't work for something rare. Lilit Bat-Leah at Epiq wrote this up for ACEDS in the technology assisted review context, and her example is the one to hold on to. If a sample of 383 documents turns up only eight of the thing you are measuring, the range around your estimate runs from 35% to 97%. That's not an estimate, that's a shrug. Now, getting to a usable answer in that situation could take anywhere from 9,500 to 38,000 documents.

Bring that back to these buckets. A Character Bucket is defined by keyword hits, and keyword hits carry noise. If the jobs that actually depict the character are a thin slice of what those keywords pulled in, then 385 jobs might contain very few of them, and the estimate built on those will be loose. Then you multiply it by a population in the millions. Now, nothing in Section 8 says a word about how common or rare the target is inside a bucket. The parties agreed the sample is representative and essentially gave up the right to argue otherwise.

But there is one thing that this whole approach cannot work without. In order to turn a percentage into a count at all, you have to know the true size of the population. And that's the single number a producing party almost never has to tell you. We've talked about that numerous times. Now, keep track of that, because Section 7 of this order is where the parties solved that issue.

Reproducible Randomness: SHA-256 and the Deduplication Carve-Outs (Section 5)

Now, the randomization method is written out step by step in Section 5, and it's reproducible. Midjourney runs the keyword terms for the bucket against its database to identify all responsive jobs. It deduplicates by unique identifier so each job appears once. It then applies a SHA-256 cryptographic hash function to each job's unique identifier, using the same hash function across every bucket. It sorts the resulting hash values in ascending order, and it produces the jobs holding the 385 lowest hash values.

A hash function takes any input and turns it into a fixed string of characters that looks like noise. The same input always produces the same output, and there's no way to predict the output from the input or work backward. So sorting by hash value is a way of shuffling the deck that is random in effect but completely reproducible. We've been using hash values in discovery for a long time. The parties' own stated rationale is that this process produces a verifiably random sample, because hash values are uniformly distributed and are not predictable. The producing party cannot steer the draw, and the requesting party can run the same hash and confirm the answer. Random selection from a script nobody else can see is a promise. This is arithmetic.

Next, the deduplication rule that the parties put together has two carve-outs that are equally and very quietly important. Distinct jobs are not treated as duplicates on the ground that they contain identical or similar prompt text. And a job appearing in more than one Character Bucket is not deduplicated away, though the parties may note the overlap.

Now, if you are an AI user, you know that you can put in the same prompt twice and get two different outputs. And that's part of what this is protecting. 10,000 subscribers typing in the same prompt is 10,000 jobs, not one. And a single image that infringes three franchises counts in all three buckets. So if it's an image that has Captain America, Iron Man, and Thor in it, that's three separate infringements for the purposes of the law. In a case where volume is the damages theory, the definition of a duplicate really is the whole ballgame here. So it's critical that the parties really called out this duplicate issue.

The Field List: Metadata That Tells the Story (Section 6)

Now, Section 6 is the field list that is necessary to be provided for each one of the entries from Midjourney, each one of the jobs. And it's worth reading if you ever draft one. For each job, the production includes the prompt text, the image or video outputs, the URL, the user ID and job ID, the submission date, the Character Bucket, the hash value from Section 5, and an indicator of whether the job was created in stealth mode.

Then it names the machine-readable fields explicitly, and the list runs well past what you would think to ask for, including the job and user identifiers, the user's plan, the platform, the full prompt and the full command, several output URLs, the source image URL, the parent job ID, batch size, and a set of status flags. Those status flags include `is_moderated`, `is_published`, `soft_ban`, `hidden`, `user_hidden`, `mod_hidden`, `flagged`, and `matched_keywords`.

Now, if you think about all of those status flags. Was it moderated? Was it published? Was there a soft ban? Is it hidden? Is it flagged? Those fields tell the studios what Midjourney's own systems concluded about the content, which is a different and arguably better story than even just the images or the video. And the parent job ID lets you reconstruct the chain when a subscriber kept iterating towards a result. If the field exists in the system, name it in your protocol. Nobody produces a field that you do not ask for.

Now, Midjourney can withhold or redact for privacy, for things like names and addresses and phone numbers, and it has to identify the basis for each redaction.

Certification and the Verification List (Sections 7 and 8)

Now, Section 7 is the one that I mentioned earlier, and I want you to focus on for this episode. It does two things past the production itself.

First, it requires a written certification signed by responsible personnel confirming that the hash function and sampling method in Section 5 were actually followed, that the hash was applied to the real prompt identifiers in the production database, that no prompts were excluded except through the keyword criteria and the agreed-upon deduplication, and that the production contains the 385 lowest hash jobs.

Second, and this is the provision I have not seen before, it requires a Verification List. Contemporaneously with each sample, Midjourney produces a machine-readable CSV listing every prompt ID in the eligible universe, meaning every ID the keyword search returned before any hashing, one row per ID, and with that ID is its SHA-256 hash value. No prompt text, no outputs, just identifiers and hashes. And the order says that obligation is independent of the obligation to produce the sample.

So pay attention to what that gives the requesting party. It gives them the true size of the population, which is the denominator I told you they would need to turn a percentage into a count. It lets them sort the hashes themselves and confirm that the 385 they received are really the lowest ones. And if a job they know about is missing from the list, they can prove the search missed it. Now, those two provisions, the certification and the Verification List, are what turns sampling from something you have to trust into something you can audit. Everything is due three weeks after the entry of this order.

Now, Section 8 then says what the sample can be used for. Experts can extrapolate from the sample to estimate the total number of prompts generating outputs depicting a character, and for buckets that cover more than one character, to estimate the franchise-level total. The parties agreed not to challenge extrapolation on the ground that the sample is unrepresentative. They expressly preserved every other challenge, including a challenge to any conclusion that particular outputs depicting an at-issue character are infringing. So by providing this evidence, they're saying they're still disputing whether or not there's an actual infringement.

What the Public Order Adds (Dkt. No. 166)

Now, what the public order adds. This next second order, which is Docket 166, covers the other half, the material not created in stealth mode. And rather than rebuild any of it, it actually borrows directly from the private order. Sections 4 through 7 of the private order apply as if set out in full, reading Public Prompts for Private Prompts and Public Sample for Sample. You're just changing the words from public to private, from private to public. Section 8 on extrapolation applies equally. Same 385, same hash value, same certification, same Verification List.

Four things that are new in the public order. First, the report comes before the files. Section 3 gives Midjourney 30 days to produce a complete report of all Public Prompts organized by Character Bucket, carrying every Section 6 field except the image and the video outputs themselves. So the studios get the entire public universe, with prompt text

and metadata, before a single output file moves. Now, the fight in a case like this is almost always about how much gets produced, and both sides usually argue it blind. Phasing the report ahead of the files here is a smart move, and it means you negotiate phase two with actual numbers on the table.

The second thing that's new here is the URL clause. For each job in a Public Sample, Midjourney produces the image and video output files from its own systems, and I'm quoting, "without regard to whether such outputs remain accessible at any public URL." A producing party running a web platform can very easily take the position that if a public link is dead, the content is gone. This says that the obligation runs to what is in the systems, not to what is still visible on the internet. If you are drafting for any platform that serves content at links, that sentence belongs in your protocol.

Third, the backstop in Section 5. For any additional specific jobs the plaintiffs identify from the report whose outputs are not accessible at the URLs provided, Midjourney produces the corresponding files upon reasonable request. Between the report, the audited sample, and this provision, the studios really have a route to catch anything the sample did not catch. A sample without a named path back to what the sample missed is not a sampling protocol. It's a cap on production.

The fourth thing is contingent scope. Exhibit A buckets proceed immediately. Exhibit B holds additional characters caught up in the live scope of discovery dispute, and those buckets come in automatically only if the Court rules them within scope or the parties agree, with deadlines running from that date. So they planned for what was already happening in the case to make changes to that order going forward. Very well thought out. Both orders state that including Exhibit B is not an admission that those characters or works are properly at issue, meaning infringing, and that entry is without prejudice to any party's position on scope. Midjourney agreed to a mechanism without conceding the merits. If you've ever hesitated to stipulate to a protocol because you did not want to concede relevance, this is how you do it.

And both orders carry the same two escape valves, which most protocols forget. We've talked about this on Case of the Week too. The parties may modify by written stipulation filed with the Court, without further leave of the Court. And if unforeseen technical problems make the methodology impossible, or Midjourney discovers an undue burden in the course of the work, the parties resolve it in good faith first, and Midjourney goes to the Court only if that fails. Each order then reserves the parties' right to seek more discovery or to move to compel.

Now, we've seen before on the Case of the Week where parties put in their language that there's good cause. I'm trying to remember, it was a hyperlinked files case, so there was good cause to amend the protocol. And that even there requires you to show good cause. The parties instead here said, we don't have to show good cause. If we both agree, we can file a stipulation with the Court without any leave of Court. So the Court doesn't have to agree, or we can just file it with the Court. That's really a critical difference, right? You're not putting some sort of good cause standard in here. So think

about the best approach that you want to use and what makes the most sense for your case.

How Prompts Got Produced in the OpenAI Cases

Now, let's put what's happening here in Midjourney next to other courts that have made an AI company produce prompts and outputs at scale, because the contrast was really the lesson.

In [the consolidated OpenAI copyright litigation](#) that we covered on two or three episodes of Case of the Week, the news plaintiffs went after ChatGPT output logs. OpenAI had retained tens of billions of them in the ordinary course of business. And in July of 2025, the plaintiffs moved to compel 120 million logs. In August, they accepted OpenAI's own counterproposal of a 20 million log sample, de-identified using OpenAI's internal tool. Then in October, OpenAI changed course and said, rather than produce 20 million logs, how about we run search terms across the sample and produce only the hits, arguing that that was less burdensome and better protected user privacy. The plaintiffs moved to compel the full sample of 20 million.

Now, Magistrate Judge Wang granted that motion by plaintiffs on [November 7th](#), [denied reconsideration](#), and extended the order to the class plaintiffs on [December 5th](#). [The district court then affirmed](#), meaning that OpenAI was required to produce all 20 million logs. It could not further segregate them by search terms.

Three pieces of that reasoning matter going forward. On relevance, the Court held that the output logs which do not contain reproductions of the plaintiffs' works were still relevant, because they bore on OpenAI's fair use defense. Now, that's a broad relevance holding, and it is why the search term proposal failed. On privacy, the Court found the interests adequately protected by three things together: the reduction from tens of billions of logs to 20 million, the de-identification, and the existing protective order. And on proportionality, the Court said there's no case law requiring a court to order the least burdensome discovery available.

Now, parties are learning from what other courts have done. The OpenAI decisions made clear that parties will have to provide prompts as evidence where they are relevant, and that the details of how that is done should be worked out by the parties. There are two differences between how the parties worked through this in each case, and they are important to consider.

The first difference is when each party asked. OpenAI and Midjourney wanted substantially the same thing, which is a produced set narrowed by keyword rather than handing over the whole. OpenAI asked for that in October of 2025, after it had already agreed to the size of a sample, and Magistrate Judge Wang said no. Now, there's a difference here, and that is that OpenAI essentially opened the door because it had this fair use defense that Magistrate Judge Wang said you can't be limited to just use of the

plaintiffs' works. So it is a little bit different. But Midjourney has keyword-defined populations written into the definition of the sample itself, agreed in Exhibits A and B and the Character Buckets before anything was collected. Same instinct, opposite outcome. And the difference is that one of them was a stipulation and the other was a motion filed after the fact. Of course, you still have to remember that the OpenAI cases were really the first of their kind, and everybody is learning from what happened in OpenAI.

The second difference is what each sample can be checked against, and this is the one that I would talk to your client about. 20 million is a negotiating number. It comes out of a proposal, and nothing in that process tells the requesting party how the 20 million were chosen from the tens of billions, or lets them verify the selection afterward. 385 instead is a statistical number, drawn by a published method, certified by a signed statement, and accompanied by a list of the entire eligible universe so that the other side can rerun the draw. The OpenAI plaintiffs got vastly more documents. The Midjourney plaintiffs got something the OpenAI plaintiffs do not have, which is the ability to prove the sample is what it claims to be.

Now, that's the trade, and it's worth thinking about on both sides of the v. A small verified sample can be worth more than a large unverified one. And a producing party that offers verification is in a much stronger position to argue for a small sample. Now, keep in mind that the [Midjourney](#) case had something the OpenAI case did not: defined characters to search for. And that matters. It basically changes the ballgame. Comparing discovery decisions is almost like comparing apples to oranges, but we have to try and draw the parallels where we can.

Now, there is a direct link between these two cases. In this same Midjourney case, on June 15th, Magistrate Judge Richlin largely denied Midjourney's motion to compel discovery into the studios' own use and development of AI, and cited the OpenAI litigation in doing it. In the same order, he allowed discovery of the prompts and contemporaneous outputs the studios used to generate images in the operative complaints, while protecting as work product the prompts behind examples that were tried and not used. So it gave Midjourney something. That is the [Warner](#) and [Morgan](#) rule applied to a plaintiff's litigation prompts, in a case about AI outputs, by the same judge who granted both of today's orders. So he made the distinction. Kudos to Magistrate Judge Richlin.

And that same June order left open a question that the magistrate judge in [Concord Music Group v. Anthropic](#) found relevant, which is how many prompts it took to generate each sample. Now, the parties considered that here, and the parent job ID field in Section 6 is how they answered it.

Takeaways

All right, let's talk about takeaways from this case. And I know this is a little longer episode, but stick with me, because these are important.

Before the specific ones, here's the pattern that I promised you earlier, and I think it matters more than any single provision in either order. Parties are reading these decisions and rewriting their protocols in light of how courts are holding. That's not me speculating. In [Tate](#), the Court did not just rule on the work product question. It told the parties to go amend their protective order to spell out whether and how confidential material may be put into an AI tool at all. In [Conservation Law Foundation](#), the holding was that counsel's general language of notes in a Rule 29 stipulation did not reach the expert's AI prompts, which was a direct instruction to name those artifacts in the document. The [Cerebras](#) protocol was negotiated by OpenAI's lead trial counsel in the copyright cases, in the months after OpenAI lost the discovery fights that mattered in the Southern District of New York. And the two orders in front of us today answer, in advance and by agreement, nearly every objection that OpenAI raised and lost. The population is defined by keyword up front instead of narrowed by motion afterwards. The sample is small and the statistics justifying it are stated on the face of the order. The method is published rather than asserted. The privacy redactions are built into Section 6 rather than litigated.

That is case law turning into contract language in under a year. More appropriately, protocol language, but that essentially becomes an agreement between the parties once the court enters it. It means that the protocol that you are handed this year should not look like the one you were handed last year. And if it does, somebody is not staying on top of the case law.

So, several individual takeaways for you to focus on.

One. A sampling protocol should be verifiable, not just agreed. Put [Schulte v. LinkedIn](#) next to this. There, the plaintiffs went to court for the elusion estimates, the error rate, and the human reviewer counts behind LinkedIn's use of Relativity aiR. And Magistrate Judge Beeler said no, that's discovery on discovery, and it is disfavored. Here, the producing party handed over the equivalent machinery voluntarily: a published selection method, a signed certification that it was followed, and a Verification List covering the entire eligible universe. What [Schulte](#) tells you that you cannot win on a motion, [Midjourney](#) and [Cerebras](#) show you that you can get by negotiation. If you are the requesting party, ask for all three. If you are producing, offering them is how you justify a smaller sample. If you are the producing party, you need to be prepared to negotiate this up front. Make sure you know and understand your client's technology.

Second, use a defensible number and put the statistics in the order. 20 million in the OpenAI litigation was a negotiating position that came out of a counterproposal. 385 here is a statistical conclusion, with the confidence level and margin of error stated in Section 8 in the order. When the number is a conclusion rather than a compromise, it's

much harder to attack later. And these parties gave up the right to attack representativeness precisely because the method earned it.

Third, name the AI artifacts explicitly in every document you draft. This is the [Conservation Law Foundation](#) lesson arriving from the other direction. There, counsel used the general language of notes in a Rule 29 stipulation, and Magistrate Judge Farrish held that it did not cover the expert's prompts. Here, Section 6 does not say relevant metadata. It names more than 20 fields, including the moderation flags, the publication status, and the parent job ID. Generic categories will not get you AI artifacts, whether you are trying to protect them or trying to obtain them. Use the actual words: prompts, queries, outputs, job identifiers, moderation status. If you're not sure what they are, look them up.

Four, define what counts as a duplicate, in writing. Identical prompt text does not make two jobs one job, and a record falling in two categories stays in both. Remember the open question from the June order in this same case, the one that the magistrate judge in the [Concord Music Group](#) case found relevant. How many prompts did it take to generate each sample? You cannot answer that if your deduplication rules collapse the attempts. And when the count is the point, the definition of a duplicate decides the case before anybody even looks at a document.

Fifth, get an inventory of the population before you get the documents. Section 3 of the public order is a report, not a production. Every responsive job organized by bucket, with the prompt text and the metadata, and no output files at all. 30 days. What that buys the studios is the size and shape of the population before anybody argues about how much of it gets produced. Compare that to the OpenAI litigation, which opened with a motion to compel 120 million logs and took four orders and six months to work down from there, with both sides arguing volume blind the whole way. Ask for the inventory first. Then you are negotiating the production with numbers in front of you instead of adjectives.

Six. Never agree to sampling without a path to verify it. This is the provision that OpenAI wanted and couldn't get. Having agreed to a 20 million log sample, it tried to convert that sample to a search term subset, and Magistrate Judge Wang said no. Section 5 of the public order is the negotiated version of the same instinct, and it works because the parties wrote it in at the start. Without a written mechanism to request specific items outside the sample, you have agreed to a ceiling and called it a methodology.

Seven. Build the amendment clause in when you draft. These protocols can be modified by written stipulation without leave of court. Compare that to [Schulte](#) again, where a 2023 interim ESI order with a technology assisted review clause and no mention of generative AI decided a 2026 dispute over a genAI review tool, because nobody went back to update it. The technology in your case will change during the case. Give yourself an inexpensive and cooperative way to keep up with it.

Eight. Know the shape of the prompt rules, and go back to the episodes of the Case of the Week here for the details. The short version is this. Your client's own prompts used in connection with the litigation are generally work product, depending on the jurisdiction you're in. Remember, [Tate](#) was in Texas state court, where they had a specific rule protecting prompts from the party. Your expert's prompts are generally discoverable. And the other side's product prompts, the ones its model actually ran for its users, are producible on a protocol like this one here. That is enough to know where you stand. The details are where the cases live, and they matter. [Warner](#) and [Morgan](#) for the work product rule under Rule 26(b)(3), [Heppner](#) for the exception, [Tate](#) for what happens when a party puts discovery documents into a chatbot, and [Conservation Law Foundation](#) for experts. They're all in Minerva26 under the Generative AI tag. Go and pull the one that matches your facts.

Finally, these parties did not just make a plan. They planned for the plan. Look at everything that had to happen before anybody could draft Section 5. Somebody had to know and understand that Midjourney stores an exchange as a job, and that the job is the unit you sample. Somebody had to know stealth mode existed, and that it split the population in half and needed its own order. Somebody had to think about how people actually use the Midjourney product, that they run the same prompt over and over and iterate on a result, just like on any AI tool, and then realize that deduplicating on prompt text would erase the evidence of that repetition, which is the entire volume story. Somebody also had to know there was a parent job ID field worth asking for, as well as the other metadata fields. And somebody had to see that a scope fight was still live and build Exhibit B as a contingency rather than wait for it.

None of that is legal drafting. All of it is understanding the technology first and writing the protocol second. That is what we talk about on this show every single week. And this is the cleanest example of it that I have covered.

What I Am Watching: The Damages Question Nobody Has Answered

Now, those are our takeaways. There's something else that I'm watching in this case, and I want to be clear that this is my read and not something that the Court decided.

The most interesting line of either order is the one that the parties drew in Section 8. They agreed that experts may extrapolate from the sample to estimate how many prompts generated outputs depicting a character, and they gave up the right to say the sample is unrepresentative. Then they expressly kept the right to challenge any conclusion that particular outputs are infringing. So the count can be extrapolated, but the infringement cannot.

That boundary holds very cleanly in discovery. Now, I'm not at all sure that it holds at the damages stage. And remember what I told you about the breadth of these estimates. Statutory damages under the Copyright Act run per work infringed. An estimate of how many prompts produced a character, plus or minus a few points that may translate into

millions of jobs, is not the same thing as finding out how many works were infringed. Somebody's going to have to bridge that gap here, and I don't know if anybody's worked out how they're going to do that in this particular case.

So let me be clear about what is settled and what is not. The parties have settled the process. Nobody is going to fight about how the sample was drawn, because the method is published and anybody can check the draw. What they did not settle, and what no court has answered yet, is what a number produced that way is good enough to prove. Extrapolating a count in discovery is one thing. Using that extrapolated count to establish how many works were infringed, and therefore how much is owed, assuming liability is shown, is something else entirely, and nothing in these orders gets you there. My guess is that the first real fight over that will come at the expert stage, and we'll cover it when it does.

The other thing I'm watching is how far outside AI this travels, because none of this is really about artificial intelligence. The 385 number, the hash, the certification, the Verification List, all of it works for any data source too large to produce whole. Think about consumer platform content, chat and messaging across large custodian sets, call recordings, telematics, sensor data, clickstream and server logs, body cam video, transaction data. I would expect to see this protocol cited in cases that have nothing to do with generative models at all.

Now, the hits just keep on coming, and they're only going to get more complicated as the generative AI tools become the latest technology that eDiscovery has to focus on. If you want to follow this line of cases yourself, all of them are in Minerva26 under the Generative AI issue tag. That's the entire arc, from [Warner](#) and [Morgan](#) through [Heppner](#), [Tate](#), [Conservation Law Foundation](#), everything that we've talked about, as well as both [Midjourney](#) orders, with the orders themselves. So you can pull the language straight into your own protocol.

Know the Technology Before You Draft

And here's what I want to leave you with after this longer episode, because it is bigger than this case. To know what you want in discovery, you have to know the technology you are trying to get the data from. Not roughly. Specifically. How does the system store what happened? What does it call the thing you are asking for? What does it keep? What does it discard? What does it flag? How many different ways does it flag things? These parties could write a protocol this good because somebody on each side had answered those questions before the drafting started.

If you do not know the technology in your case, go find out. Read the product documentation. Ask your client's engineers. Bring in a consultant. Call somebody who does know and ask them to walk you through it. Because if you negotiate a protocol without understanding how the system actually holds the data, you are not going to ask

for the right thing. And the worst part is that you will probably never know what you missed. The critical information in your case may be sitting in a field nobody named.

That's our Case of the Week for this week. If this was helpful, please share it with a colleague who's negotiating a protocol right now, especially one involving a data source too big to produce whole. And if you haven't already, subscribe to the Meet and Confer podcast so you don't miss the next order that moves this line.

Thanks so much for joining me, and we'll see you next time.