

Google's AI Search Is Now A Liability Disaster
https://www.youtube.com/watch?v=ehsq_0Cw6e4
Transcript: <https://dontveter.com/ec/googleai.pdf>

Some time ago, I made a video about Google's AI search. And in that video, I cited a specific number. Google's AI overviews are accurate 91% of the time.

And I made a specific argument. 91% sounds reassuring until you apply it to 5 trillion annual searches because the remaining 9% translates to roughly 57 million wrong answers every hour.

I said the consequences of that 9% would be real if a commercial airline had a 91% success rate. We wouldn't call it aviation. We'd call it an extreme sport with a mandatory life insurance policy.

That video has over half a million views now. And look, I know what you're thinking already. El, at least 100,000 of those views were just Sundar Pichai rocking back and forth in a dark room maybe.

But I mentioned this not to boast, but because what happened this week suggests that half a million people were paying attention to the right thing.

Well, a German court just proved the whole point. The regional court of Munich has ruled that Google's AI overviews are not search results. They are not summaries of third-party content.

They are not links to other people's websites. They are Google's own words. and Google is directly liable for every false claim they generate.

In the same week, a federal judge in Mississippi discovered that lawyers on both sides of a court case had used AI to prepare their filings.

Both sides cited cases that did not exist. The judge canceled the whole trial and kicked everyone out. The judge basically looked at a courtroom full of fake laws and basically said, "All right, everyone out".

The humans, the laptops, the charging cables. Go sit in the hallway and think about your lives. It is the first time in legal history that a trial was postponed due to a double cocktail of digital delusion.

Two courts, two countries, same lesson. My name is El. I have a PhD in computer science and I analyze AI developments to understand what's actually happening beneath all of this hype.

In this video, I'm going to break down both of these rulings and explain what they actually mean for Google, for the legal profession, and for anyone who uses AI.

Then, I want to explain something technical about how AI search actually works under the hood because it makes the implications of this German ruling genuinely extraordinary.

And in the final section, I want to make an argument about how all of this is actually a significant silver lining, an amazing thing, one that will improve the field of computer science, make AI considerably more useful for all of us, and bring some order to what has until now been a spectacular mess.

If you find this analysis useful, please give this video a like and subscribe. And the best way to support this work and keep it free and without any gatekeeping is through Buy me a Coffee or a channel membership. The links are down below.

Now, let's start with Munich because this ruling is a landmark. Two Munich based publishers discovered that when people search for their company names on Google, the AI overview at the top of the results page was falsely linking them to scams, subscription traps, and dubious business practices.

The AI had confused them with other genuinely problematic companies, mixing up entities, drawing connections, and generating content statements like, "Yes, company X is known for dubious business practices."

It then built its own structure, a summary, a list of red flags, and even tips for users on how to avoid being scammed.

The AI essentially acted like an overzealous true crime podcaster who skipped the investigation entirely and went straight to the dramatic finger pointing.

None of these claims appeared in any of the linked sources. By the way, the AI made them all up. It fabricated an entire narrative about two real companies, presented it at the top of Google search results with a visual authority of settled fact and served it to every person who searched for these company names.

The publishers sent Google a cease and desist letter. Google did not respond appropriately, so the publishers took Google to court.

The Munich Regional Court's ruling did something that no court has done before. It classified Google's AI overviews as Google's own content, not a reorganization of third-party search results, not a summary, but and I'm quoting, independent new and substantive statements generated by Google's own AI.

The court's reasoning was very precise. The AI rewrites and judges results in its own words and according to its own structure.

It made claims that are not even made in the search results. The sources linked beneath the overview did not contain the statements the AI presented.

These were in the court's language, the defendant's own statements. The court then examined whether Google could fall back on the legal protections that have shielded search engines for decades.

In Germany, the Federal Court of Justice had previously ruled that search engine operations were only liable as indirect infringers. They made third-party content findable. They didn't create it themselves.

Imposing a duty to check every result would threaten how search engines work. The Munich court found that this reasoning does not apply to AI overviews.

And this is the biggest difference. A regular search engine points to outside websites. AI overviews generate new content.

That is a fundamentally different activity and it attracts fundamentally different liability.

Google also tried to claim protection under the EU digital services act as a host provider, the same safe harbor framework that protects social media platforms from liability for user generated content.

The court rejected this as well. You cannot be both the author and the neutral host of the same content.

That distinction, if it holds on appeal, potentially collapses the legal shield that has protected tech platforms for the better part of two decades now.

And it applies not just to Google, it applies to any AI provider whose system generates content from web sources. Open AAI, Anthropic, Perplexity.

Every company operating in this space would need to reckon with the exact same logic.

It turns out you cannot run a multi-billion dollar printing press of automated gossip and then claim you're just the guy delivering the Sunday paper.

Google's defense at the hearing was interesting, too. Check this out. The company argued that users could check the link sources themselves to verify whether the AI summary was correct.

Google said that users generally knew that information generated with AI should not be blindly trusted. But the court's response was devastating.

The possibility of disproving a statement through further research does not exempt the person who published it from liability.

And a Pew Research study found that only 1% of users ever click a source link in AI overviews. So, Google built a product designed to give people answers without clicking.

It then argued that people should click to verify the answers. The court was not persuaded.

It is a bold defense strategy. It's like a car manufacturer selling you a vehicle with no brakes and then telling the judge, "Well, the consumers should have had the foresight to open the door and drag their foot on the asphalt like Fred Flintstone. We merely provided the momentum."

Google was, of course, ordered to pay 80% of the legal costs.

There was one more detail that I think deserves attention here. The court addressed whether AI generated statements deserve free speech protection.

Its conclusion was as follows. An AI's opinion is not the expression of an acquired conviction, but the result of an algorithm. Offering AI powered research is above all an expression of Google's business activities.

When balancing privacy rights against Google's commercial interests, Google loss, an algorithm does not have convictions. It has outputs and the company that deployed the algorithm owns those outputs whether it likes the implications or not.

Google has since provided a statement on the ruling. The company says its AI overviews are designed to reflect information that already exists on the web.

It adds that AI overviews can occasionally miscontext or misinterpret web content just like traditional search results. And then remarkably, the statement repeats the very argument the courts rejected that people can dig deeper and verify.

The court explicitly ruled that the possibility of disproving a statement through further research does not exempt the publisher from liability for that statement.

Google's response to the ruling was to reiterate the defense the ruling rejected, which tells you something about how seriously the company is taking the court's reasoning and how far the gap between legal accountability and corporate culture still exists.

Google's legal team essentially looked at a binding court order and replied with, "I hear what you're saying, your honor, but have you considered that we don't really want to do that?"

It's the corporate equivalent of just covering your ears and shouting, "La, I really can't hear you over the sound of our stock price."

The ruling is not yet final. Google will almost certainly appeal. And the German legal system is different from the American one. This ruling does not automatically apply in the United States or anywhere else.

But the logic of the ruling is portable. The distinction between pointing to third-party content and generating new content from third-party sources is not just a little quirk of German law.

It is a structural observation about what AI overviews actually do. Any court in any jurisdiction could arrive at the same conclusion by just examining the same technology.

And the European Commission is already examining whether AI overviews violate the digital markets act and the EU copyright directive suggesting that this line of reasoning has legs far beyond Munich.

Spark Toro data from 2026 suggests that 68% of Google searches in the United States now end without a click. Across all searches, by the way, not just AI mode, the information is being consumed at the point of generation.

The publishers whose content feeds the AI are not being visited. And if the AI gets it wrong, the people harmed by the false information have until this ruling had very limited legal recourse.

But the Munich court identified this as a protection gap. If Google is only liable for obvious violations and the third party sources didn't even make the claims in question, victims have nowhere to turn.

The AI fabricated the claim. The sources didn't contain it. Under the old framework, nobody was responsible. The Munich court closed that gap.

Whether other courts follow remains to be seen, but the argument for doing so is now on the record with reasoning and precedent attached to it.

Now, while the Munich court was ruling on what happens when Google's AI gets it wrong, a federal court in Mississippi was dealing with what happens when everyone else's AI gets it wrong as well.

The case was straightforward, a contractual dispute between a lawyer and the city of Aberdeen, Mississippi, over unpaid legal fees.

What was not straightforward was how the lawyers on both sides prepare their arguments. They used AI, both of them, and both of them cited cases that do not exist, hallucinated presidents, fabricated legal authorities on both sides of the same case.

Senior United States District Judge Sherion Akott described the situation in a sanctions order as unusual, which in judicial language is roughly equivalent to setting the building on fire and calling the weather a little warm.

She wrote, "In an era of rampant unverified AI usage within the legal field, this case represents a prime example of the risk associated with serving as a rubber stamp."

And keep in mind, Judge Akott didn't just give them a stern talking to. She fined them up to \$3,500 and banned the lead attorneys from her courtroom for 2 years.

They didn't just lose the case. They got legally grounded. Lawyer Rob Floyd, who first identified the case, called it a comedy of AI errors and noted that the two clients were essentially paying their lawyers for ChatGPT to argue against itself.

Two artificial intelligences generating fake legal authorities, citing non-existing precedents, marshaling hallucinated arguments, while four human lawyers signed their names to the output without checking whether any of it was real.

It is a landmark moment in human laziness. We finally achieved the dream, my friends.

Two algorithms hallucinating a functional legal universe at 60 tokens per second while four human beings who went to law school, and passed the bar by the way, acted as glorified copy and paste interns for a piece of software.

The billable hours on this must have looked hilarious, like 3 hours watching ChatGPT make up fake law about solar panels.

That'll be \$1,200 one hour napping while the chatbot creates a fictional Supreme Court justice named Arthur Pendragon 400.

The legal profession has a word for the skill that distinguishes it from a large language model. It's judgment. Judgment is the word. What was missing from this courtroom was not technology. It was judgment.

And this is of course not an isolated incident where there's human there is some shenanigans going on. Judges across the United States have been increasingly frustrated with lawyers submitting AI generated filings.

Just last week, a New York judge publicly ripped into attorneys for citing hallucinated cases.

The Mississippi case is simply the most extreme example yet because this time it wasn't one side cutting corners. It was everyone.

Let me be very clear about where I stand on this because I think it's very important. I have nothing against lawyers using AI. I have nothing against anyone using AI.

Use it to do research. Use it to construct arguments. Use it to organize your thoughts, draft initial structures, explore angles you hadn't really thought about. This is what the tool is good at.

What you cannot do, what these lawyers actually did, what cost them their case, their client's interests, and their professional standing is treat AI as a replacement for the work of being a lawyer.

You still have to check, you still have to verify, you still have to read the cases you cite and confirm they actually exist. That is not optional. That is literally the job description.

Imagine being the client who hired a lawyer, paid for their expertise, trusted them to represent your interest in a federal court, and then discovered that the arguments presented on your behalf were generated by a chatbot and never verified by a human being.

That is not an AI failure. The AI did what the AI does. It generated plausible sounding text. The failure was entirely human.

Four lawyers treated AI output as finished work product, signed their names to it, and submitted it to a federal court. The tool worked exactly as it works.

The humans abdicated their responsibility. And this connects to something I think is generally important. AI right now is like a remarkably capable 10-year-old.

It can do wonderful things within reason and with appropriate supervision. But you still don't hand a 10-year-old the keys to the family car, give them a briefcase full of legal briefs, and say, "Good luck with a corporate litigation, honey. Dinner's at six."

It can research faster than you can. It can draft faster than you can. It can identify patterns and connections you might miss.

But it is not ready to be left alone with consequential decisions. It still needs a human in the room, not as a rubber stamp, not as a passive observer, but as an active, engaged, thinking participant who checks the output against reality.

Used that way, as augmentation, a partner, a tool that makes humans more productive, it is genuinely valuable.

But used as a replacement for thinking, it produces exactly what we saw in Mississippi. Two AIs arguing against each other with fake evidence while four humans watched.

I want to take a step back now and explain something technical because it changes how you understand the scale of the German rulings implication.

Google's AI overviews are not pre-written. They are not stored in a database waiting to be shown when somebody searches for something. They are generated in real time per query using a technique called retrieval augmented generation or rag.

Here's what this means in plain language. When you type a search query, Google system does two things in sequence.

First, it searches its existing index, so the same database that powers traditional search results, and pulls the top relevant web pages for your query.

Then, instead of showing you those pages as a list of links, it feeds the text from those pages into its Gemini AI model and asks the model to write a summary.

The AI reads the sources, synthesizes an answer in its own words, generates citation links, and serves the result to you.

The entire process takes between 1.4 and 5 seconds. For complex queries, it can pull from four to eight sources simultaneously and weave them into a single coherent response.

So, every single AI overview is a unique piece of content that did not exist until somebody searched for it. It is manufactured and consumed in the same moment.

If you search the same query tomorrow, the AI might generate a slightly different response because the underlying sources may have changed or because the model's generation process involves a degree of randomness.

There is no master copy. There is no editorial review. There is no human being who reads the output before it reaches you.

It is written, served and read in a single unbroken sequence. It is pure unadulterated automated jazz except instead of musical notes, it's improvising corporate liabilities at scale.

Now think about what the Munich court has said. These are Google's own words and Google is liable for them. At Google's scale, billions of queries per day. The court is effectively requiring quality assurance on content that doesn't exist until the instant is delivered.

Imagine running a restaurant where the chef doesn't decide the menu until the customer sits down, randomly throws ingredients into a blender based on a vibe, and you, the owner, are legally liable if it poisons them.

That is essentially Google search engine right now. It is a real time culinary roulette, except the soup is made of algorithmic hallucinations.

That is, with current architecture, a problem of genuinely extraordinary computational difficulty. You cannot prevent a statement that hasn't been written yet.

You cannot fact check content that is generated at the moment of serving for now.

For context, moderating user generated content on a social media platform, which is already an enormous challenge, by the way, involves reviewing content that exists before publication.

AI overviews do not exist before publication. They exist only at the moment of delivery.

The court has demanded accountability for a product whose output is by design unpredictable in advance.

And this is where my background in computer science becomes very relevant because the ruling effectively identifies a research problem. One that the field has been aware of but has not prioritized with the urgency it deserves.

The failure modes in AI generated search responses break down into three broad categories and understanding them matters because each one requires a different kind of solution.

The first one is retrieval failure. The system pulls the wrong documents or confuses entities with similar names.

This is exactly what happened in the Munich case. The AI retrieved information about genuinely problematic companies and attributed it to the plaintiffs who had no connection to them whatsoever.

This is fundamentally a named entity disambiguation problem which is a deeply boring academic way of basically saying the AI looked at two different companies and basically went, “ah, they both start with a capital letter”. It's basically the same thing.

The solutions exist, knowledge graphs, entity linking, verification against structured databases, but they have not been deployed at the speed and scale that Google's real time generation demands.

The second category of failure is synthesis failure. The model has the right sources in front of it, but draws connections or makes claims that are not supported by those sources.

This is the 56% ungrounded sources problem I cited in the Google video before.

More than half of AI overview answers that were technically correct cited sources that did not actually contain the information presented.

The AI is confabulating, generating plausible sounding content that has no basis in the material it was given.

Research on faithfulness verification, which is essentially a second pass that checks whether the generated text is actually entailed by the source documents, exists and is advancing.

But it does add computational cost and latency. That means it's going to make it slower. That the current product architecture is not designed to accommodate.

The third category of failure is confidence calibration failure.

The model presents uncertain information with the same visual authority as certain information. Giving the user no signal about when it might be wrong.

Everything looks equally confident, equally definitive, equally trustworthy. Research on uncertainty quantification flagging claims where the model's confidence is low or where sources conflict has been progressing in academia.

There is a lot of papers on this but has not been deployed at scale because frankly it makes the product look less impressive.

An AI overview that says I'm not really sure about this is a less compelling product than one that says yes this company is known for dubious business practices.

I'm 100% sure the incentive to appear confident has until now outweighed the incentive to be accurate.

The tech industry built these models to mimic the absolute worst traits of a tech bro.

Being completely wrong, but delivering it with the unshakable confidence of a condescending jerk who owns a crypto startup.

It is the digital embodiment of a man explaining your own PhD topic to you at a party.

These are serious structural problems that I just described.

Ultimately, they are for serious researchers to sit down with and work through methodically. It is out of scope of a YouTube video to solve them.

But naming them matters because it shows that the path to more reliable AI generated content is not mysterious.

The technical approaches exist. They are simply not being prioritized because the market has rewarded speed and capability over accuracy and reliability.

The Munich ruling just changed that calculation. And this is where I arrive at the argument I promised at the beginning.

Our beautiful silver lining. Everyone is going to cover this story as Google getting slammed and it did.

But I think something much more important is happening and I think it is generally a positive thing.

The field of large language modeling has been on a specific trajectory for the past 2 years. The overwhelming majority of research energy, competitive pressure, and investment has been directed towards enterprise utility, making models better at coding, better at structured tasks, better at the kind of verifiable output that enterprises are willing to pay for.

And that push has produced real improvements. Yes. But accuracy for general knowledge synthesis, which is the ability to retrieve information from multiple sources, combine it faithfully, present it without fabrication, and signal uncertainty when appropriate, has not received the same intensity of focus.

It has been comparatively kind of a secondary priority. And this has been the case mainly because the market wasn't punishing inaccuracy with the same force that it was rewarding capability.

But the Munich ruling changes the incentive structure. If you are liable for what your AI says, accuracy is no longer a nice to have feature. It is a legal requirement.

And this is not just one regional court in Germany.

The EU AI act becomes enforceable on August 2nd, 2026, less than two months from now.

Article 15 specifically requires appropriate levels of accuracy, robustness, and cyber security for high-risk AI systems.

The European Commission has already launched a probe into AI overviews to examine whether the feature violates the digital markets act and the EU copyright directive.

China has delivered its first court ruling on AI hallucinations. The legal systems of three major economic blocks are simultaneously drawing lines around AI generated content.

So this is not a local story. It is a global saddle point, an inflection point. And this matters enormously because humanity genuinely needs this technology to develop.

The volume of research papers published every year, the expanding body of human knowledge across every discipline, the sheer quantity of information that exists, it is growing beyond what any individual or team of humans can synthesize alone.

A single researcher in any field cannot read every paper published in their own native discipline. A doctor cannot stay current with every study relevant to their patients.

A policymaker cannot track every data point relevant to the decisions they make. AI powered knowledge synthesis is not a luxury. It is becoming a necessity.

The ability to retrieve, combine, verify and present information from millions of sources accurately and reliably.

That capability would be transformative for education, medicine, scientific research, policymaking, journalism, and every field where understanding the current state of knowledge matters, which is basically all of them.

I cover this in a lot more depth in my video about AI in education and academia.

I'm going to link it down below if you want to dive deeper on this topic. But it has to be accurate. It has to be trustworthy. It has to be reliable.

A medical summary that is 91% accurate is a medical summary that gives wrong information to 9% of the people who trust it.

A legal research tool that hallucinates the president, waste the court's time, and harms the clients who relied on it.

A search engine that falsely accuses publishers of fraud damages real businesses and real people.

The technologies potential is extraordinary, but potential without reliability is just a more sophisticated way to be wrong. Right?

And left to its own devices. The market was not prioritizing reliability because accuracy is expensive.

Speed to market is rewarded and nobody was being held legally responsible for getting it wrong.

The corporate strategy for the last 2 years has basically been ship it now, apologize later, and if anyone gets hurt, we'll just claim the AI is sentient and has its own constitutional rights.

The companies deploying these systems were in a race to ship, to scale, to capture market share.

The 91% accuracy rate was treated as a triumph rather than a problem because the 9% wasn't costing the companies much or anything. It was costing the people who trusted the output.

What the Munich court has done, what the EU AI act is about to do, what courts from Mississippi to China are beginning to establish, is create legal and regulatory incentive for accuracy that the market was not providing on its own.

And I know that there will be some voices in the comments who say, well, this is just regulation slowing AI down. But I think it's actually pointing AI in the right direction.

We are not banning the technology or restricting its development, just making the people who deploy it responsible for what it says, which really shouldn't be such a radical concept.

If my dog bites the postman, I can't really tell the police, "Well, his neural network is just a black box." And he was just predicting the next logical placement of his teeth. I would get a fine. The dog would get a leash.

Welcome to accountability, Silicon Valley. And honestly, even the dog has the decency to look slightly guilty afterwards.

Silicon Valley just updates its terms of service quietly and kind of hopes you scroll past it fast enough to experience thumb fatigue.

The mess is real. The lawsuits are real. The hallucinated legal citations are real.

The falsely accused Munich publishers are real. But out of this mess, the incentive to build AI that is genuinely accurate, not just impressively fluent, is finally becoming as strong as the incentive to build AI that is fast, cheap, and capable.

And that is the development that will make all of this technology actually useful in the ways that matter most.

I said at the beginning of this video that a German court proved a point I made a few weeks ago, but the real point is bigger than one ruling, one court or one company.

The real point is that AI generated content has for the past 2 years existed in a kind of legal vacuum. too new for existing frameworks, too fast for regulators, too profitable for the companies deploying it to voluntarily slow down.

That vacuum is closing. Not all at once, not perfectly of course, but it is measurable, verifiable in courtrooms and legislatures across the world.

And the technology that emerges on the other side of this, the AI that is built to be accurate because it has to be, not because somebody chose to make it so, that technology will be worth the mess it took to get here.

I genuinely believe that and I really look forward to covering it when it arrives.

If you want the full picture of what Google's AI search strategy looks like, why people are leaving for duck duck go and what the accuracy data actually says, I covered all of that in a lot of detail in my video on Google's AI search.

It has the context for everything we discussed today and is the video that I would watch next. Thanks so much for watching this one. Subscribe and I'll see you all on the next.