

How AI answer engines decide what to cite

By Nabila Husseni | Content Strategy; AEO | September 2026

An article I wrote for Adastra Group started appearing in ChatGPT answers. Then Gemini. Then Perplexity. I did not submit it to any of them. I did not use any special technique. I just wrote it a certain way.

Since then I have spent a lot of time thinking about what actually happened and why. This is what I have figured out.

First, a quick note on terminology. You have probably heard of SEO, which is about getting your content to rank in Google search results. AEO, which stands for Answer Engine Optimization, is about getting your content cited in AI-generated answers. They overlap but they are not the same thing, and the differences matter if you are trying to get cited rather than just ranked.

What AI answer engines are actually doing

When someone asks ChatGPT or Gemini a question, the model does not search the internet in real time and pull the most recent answer. It uses a combination of its training data and, in some versions, a live search that retrieves current content. But either way, the model is making a judgment about which content to use as the basis for its answer.

That judgment is not based on how popular the page is or how many other sites link to it, the way traditional search ranking works. It is based on something closer to: is this content a reliable, authoritative, well-structured answer to this type of question?

AI models are not looking for the most popular answer. They are looking for the most trustworthy one.

That is a meaningful distinction. A page can have very few backlinks, relatively low traffic, and still get cited by AI answer engines if the content itself signals authority and usefulness. Which is what happened with the Adastra article.

The signals that seem to matter

I want to be honest that nobody outside the companies building these models knows exactly how citation decisions work. The models do not publish their ranking criteria the way Google publishes its search guidelines. What I am sharing here is based on what I observed, tested, and read from people who study this closely.

With that caveat, here is what seems to matter:

Original analysis and first-hand experience.

The Adastra article worked, I think, because it was not summarizing what other people had written about BI and analytics. It was explaining how these things work based on direct engagement with the subject. AI models seem to weight first-hand expertise differently from content that is aggregating existing information. When you have actually worked with something and you write from that experience, the content reads differently from a summary. Models appear to pick up on that.

Clear answers to specific questions.

Content that is structured around answering a question directly tends to do better than content that is structured around covering a topic comprehensively. If someone asks an AI what a data

warehouse is, the model is looking for content that answers that question clearly and early, not content that covers data warehouses from every possible angle. Writing for a specific question rather than a broad topic is one of the simplest shifts you can make.

Structure that is easy to parse.

AI models read your content and try to understand what it is about, what the main points are, and how confident the author is in what they are saying. Clear headings, short paragraphs, and logical structure all make it easier for a model to extract the key information. This is similar to what makes content scannable for human readers, which is not a coincidence. Models are trained on human reading behavior.

Specificity over generality.

The more specific a claim, the more trustworthy it tends to read. A sentence like 'data warehouses improve query performance' is a vague claim that many articles make. A sentence like 'dimensional modeling reduces query time for aggregation-heavy workloads because it pre-joins tables that would otherwise require expensive run-time joins' is a specific claim that shows the author understands what they are talking about. Specific claims are harder to fake and models seem to reward them.

What does not seem to matter as much

Keyword density does not appear to be a primary factor. The Adastra article was not written around keyword targets. It was written to explain something clearly. The keywords appeared naturally because I was writing about the subject accurately.

Publication volume also does not seem to be as important as people assume. One well-written, authoritative piece on a specific topic can get cited repeatedly. Publishing ten thin pieces on related topics is less likely to produce citations than one deep piece that genuinely answers a question well.

AEO rewards depth over volume. One article that really answers something is worth more than ten that kind of cover it.

The practical implication

If you want your content to be cited by AI answer engines, write like you are answering a specific question for someone who needs to understand something, not like you are covering a topic for someone who might be interested in it generally.

The distinction sounds small but it changes everything about how you organize the content, what level of detail you go into, how you handle caveats, and how you structure your conclusions.

Content written to answer a question has a beginning (what is the question), a middle (here is how to think about it), and an end (here is what you should take from this). Content written to cover a topic often has a beginning (here is what this topic is), a long middle (here are all the aspects of this topic), and no real end.

The first structure is what AI models are looking for because it is what the person asking the question needs.

One thing to keep in mind

AEO is not a fixed target. The models that power these answer engines change, their training data changes, and the way they handle citations changes. What works today may need to be revisited in twelve months.

The underlying principle, though, is probably stable: write content that is genuinely useful, clearly structured, and based on real knowledge or experience. That has always been what makes content worth reading. It turns out it is also what makes content worth citing.

Nabila Husseni is a content strategist and technical writer. The Adastra BI Analytics article discussed in this post is at nabilahusseni.com and is currently cited by ChatGPT, Gemini, and Perplexity. You can verify this by asking any of them about enterprise BI and analytics strategy.