

Citation Monitor Lite

Find out which domains actually win your category in ChatGPT. Not a guess, a named list. About an hour to set up, then twenty minutes a week.

What is in this kit: this guide, the Citation Monitor Lite spreadsheet, and a copy-paste Claude Code prompt that automates the whole thing if you would rather not do it by hand.

Why this beats checking whether you show up

Checking whether your name appears in an AI answer tells you your score. It does not tell you the game. Citation tracking asks a different question: **when you are invisible, who is filling the space, and through which sources?**

That distinction matters because the answer is frequently uncomfortable. One citation study published in July 2026 ran 15 buyer queries against ChatGPT's search-enabled model, three samples each, and logged roughly 300 citations. The target brand's own website was cited exactly once across all 300. When that brand appeared in ChatGPT's answers at all, it arrived through third-party sources, almost never through anything the company controlled.

That is the whole argument for measuring before optimizing. Publishing another comparison page on a blog nobody is citing changes nothing, and you would never learn that from a mention check alone.

1 Build your query set

15 to 20 questions your buyers actually ask, with no brand names in them. Ask what someone asks before they know you exist.

Split them three ways: about eight **evaluation** questions ("what's the best X for Y"), about six **comparison** questions, and about six **definitional** questions ("what is X"). Evaluation is where buying decisions form. Definitional is where people who have never heard of you form their first impression.

Put these in the Queries tab of the spreadsheet, replacing the example rows.

2 Ask with web search on, three times per query

This is the technical piece that matters most. A plain chat answers from training data and cites nothing. Only a search-enabled model returns real sources.

Three samples per query, because a single run is a coin flip rather than a measurement. Fifteen queries at three samples is 45 answers.

3 Log every cited URL

Query, sample number, domain, full URL, page title. If a sample returns no citations at all, log that too with the domain "NONE", otherwise your hit-rate math quietly inflates itself.

Citation Log tab. Room for 300 rows.

4 Classify each domain

Six buckets. This classification is the whole point, because each type has a completely different playbook.

Type	What it means
Your site	A page you own and control
Competitor	A direct competitor's own domain
Review platform	G2, Capterra, Trustpilot, industry-specific review sites
Community	Reddit, forums, Q&A sites, Stack Overflow
Publisher	Media, blogs, roundup posts, newsletters
Docs	Documentation, help centres, wikis

The Domain Rollup tab counts everything for you. Type each unique domain into column A and it fills in the rest.

Sort by "Queries it appears on", not by total citations. A domain cited twenty times on one query is a niche win. A domain appearing across four or more different queries is doing structural work in your category.

Three ways to read the same data

- **Down a single query.** Which domains own this specific answer? If a query that matters to your revenue returns two review platforms and one forum thread, you now know the three surfaces to work on, and that none of them is your blog.
- **Across the source types.** What share of your category's citations go to each type? That distribution is your channel strategy, derived from data rather than instinct.
- **The trend line.** Weekly runs show whether the work is moving anything, which is the only way to tell "this is not working" apart from "this is working and I cannot see it yet."

What each pattern tells you to do

If the citations show	The play is
Review platforms dominate	Review volume, recency, and a complete profile on the specific platforms cited. Not all of them, the ones actually showing up.
Communities dominate	Genuine participation where the cited threads live. Sometimes one specific thread is doing all the work. Read it first.
Publishers dominate	Earned media. Get into the specific roundups being cited. Your citation data just wrote your target list.
Competitor sites dominate	Read what they published that you did not. Frequently a comparison page framing their own alternatives.
Docs dominate	Problem-first queries are being answered by documentation. Publish the real how-to your buyers search for.
Your site barely appears	Check whether the page even exists for that query intent. A missing page is a content spec. An ignored page is a structure problem.

The honest limit

The monitor tells you where the doors are. Opening them is a different kind of work, and frequently not content work at all. Review programs, community participation, and earned media each sit with a different function inside a company, which is why AI visibility behaves more like a go-to-market program than an SEO task.

For a solo operator that is good news: you own all of those functions already. There is nobody to coordinate with. The constraint is your time, not organizational alignment.

Run it weekly

Citation sources shift as models update their search behavior. A domain that dominates this month can fade next month. One run is a snapshot. Weekly runs are a monitoring system.

Want it run for you?

The AI Visibility Audit runs this across every buying query in your category, across ChatGPT, Perplexity, and Google AI, and hands back a ranked fix plan. Free scan first, no email required.

closermethod.com/seo

Methodology adapted from Sara Soleymani's published citation-monitoring build, Medium, July 2026, which included both an n8n workflow and the original Claude Code prompt. Worth reading directly for the full technical reasoning behind the sampling and classification choices. Supporting data: Forrester 2026 Buyers' Journey Survey; Muck Rack / Generative Pulse, 2026.