

How to Track Your Brand Visibility in AI Search With Peec AI

Metehan Yesilyurt

AI Search & Information Retrieval Researcher · metehan.ai

Published: September 06, 2026 · Canonical: <https://metehan.ai/articles/how-to-track-brand-visibility-in-ai-search/>

Tracking brand visibility in AI search means measuring how often ChatGPT, Perplexity, Google AI Overviews, AI Mode, and Gemini mention or cite your brand across a fixed set of prompts, then monitoring mentions, position, sentiment, citations, and revenue over time. The best way to track AI brand visibility is a repeatable loop, not a one-off audit, because AI answers change between runs and a single snapshot misleads. This guide walks through the full process step by step, using Peec AI as the tracking layer.

The five-step summary

- Define a prompt set that mirrors real buyer questions across your funnel and markets.
- Set a baseline in the first 7 days before you judge any number.
- Track each AI channel separately, because ChatGPT, Perplexity, AI Overviews, AI Mode, and Gemini behave differently.
- Benchmark competitors with share of voice, not just your own mention counts.
- Connect visibility to traffic and revenue so the program survives budget review.

What AI Visibility Tracking Actually Measures: Metrics and Formulas

AI visibility tracking measures five core metrics, and each one answers a different question about how answer engines treat your brand.

Visibility rate is the percentage of tracked prompt runs in which your brand is mentioned in the answer. **Citation rate** is the percentage of runs in which one of your URLs is linked as a source. **Average position** is where your brand appears in the answer relative to other brands, counted from the top. **Share of voice** is your share of all brand mentions in your competitive set on the same prompts. **Sentiment score** is how positively or negatively the answer describes your brand when it does appear.

Metric	Formula	Healthy range	Warning range
Visibility rate	Mentioned runs / total runs x 100	30%+ on branded and solution prompts	Under 10% on prompts you should own
Citation rate	Runs citing your URL / total runs x 100	15%+ on content you built for AI	Under 5% despite mentions
Average position	Sum of positions / mentioned runs	Top 3 in your category prompts	Consistently 5th or lower
Share of voice	Your mentions / all competitor set mentions x 100	Above your market share	Falling while a rival rises
Sentiment score	Positive minus negative descriptions, normalized	Stable positive	Any negative trend on one channel

Two rules for reading these numbers. Mentions and citations move independently, so track both: a model can name you without linking you, and link you without naming you. And position matters more than raw mentions for high-intent prompts, because answer engines front-load their recommendations the same way users read them.

Step 1: Build the Prompt Set That Mirrors Real Buyer Questions

Your tracking is only as good as your prompt set. The goal is not to guess every phrasing your buyers type, because LLMs understand intent and most real prompts are unique anyway. The goal is coverage of the questions that decide deals.

Selection criteria that work in practice:

- Cover all funnel stages: category discovery ("best AI visibility tools"), comparison ("Peec AI vs alternatives"), and validation ("is X worth it") prompts behave differently and all matter.
- Include your top 3 competitors' branded prompts, because that is where switching decisions happen.
- Write prompts in the languages of your priority markets. AI engines answer differently per language, not just per country.
- Start with 50 to 100 prompts per brand. Enough for stable trends, small enough to review weekly. Scale to 300 to 500 once the routine works.

The Peec difference here is prompt discovery from fanout data. When an AI engine answers a prompt, it silently runs its own sub-queries behind the scenes, and Peec surfaces those real fanout queries so you can track the questions the engines actually ask, not just the ones you imagined. In my own dataset, a single buyer-intent prompt typically fans out into 3 to 8 sub-queries, and those sub-queries are where citation decisions get made.

[GÖRSEL PLACEHOLDER 1: Peec AI prompt setup / prompt suggestions view]

Step 2: Establish Your Baseline in the First 7 Days

Do not judge any number on day one. AI answers are non-deterministic: in my audits, 15 to 30 percent of prompts change their brand lineup between two runs on the same engine on the same day. A one-day snapshot is noise wearing a suit.

What to do in the first week instead:

- Let the full prompt set run daily for 7 days before drawing any conclusion.
- Record the weekly average per metric per channel. That average is your baseline, not the best or worst day.
- Note your variance. If a prompt's visibility swings between 20% and 80% across the week, it is a volatile prompt and needs trend reading, not daily panic.

From week two onward, read every number against the previous period. "Visibility rate 42%" means nothing alone. "Visibility rate 42%, up from 35% last week, driven by AI Overviews" is a finding. Peec's dashboard shows previous-period deltas by default, which quietly enforces this discipline.

[GÖRSEL PLACEHOLDER 2: Peec AI dashboard with previous-period comparison]

Step 3: Track Each AI Channel Separately

Averaging your visibility across engines hides exactly the differences you need to act on, because each channel retrieves, cites, and describes brands differently. This is the heart of the process.

How to Track Brand Mentions in ChatGPT

What is different on this channel: ChatGPT separates what the model says (mentions) from what its search layer reads (sources), and the two often disagree. Its source selection runs through internal search queries that differ from the user's prompt, so a page can be read without being cited and cited without your brand being named.

What to measure: mention rate and citation rate separately, plus which of your URLs appear as sources. If ChatGPT mentions you but never cites you, your brand authority is carrying the answer while your content is invisible to its retrieval, and those need different fixes.

How it looks in Peec: ChatGPT tracking shows mentions, position, and the exact source URLs per prompt, with daily refresh so you catch answer changes the day they happen.

[GÖRSEL PLACEHOLDER 3: Peec AI ChatGPT mentions and sources view]

How to Track Brand Mentions in Perplexity

What is different on this channel: Perplexity is citation-first. Nearly every claim in an answer links to a numbered source, which makes source tracking the main event rather than a bonus. Continuous monitoring matters more here because Perplexity re-retrieves on nearly every run, so its source list changes faster than ChatGPT's.

What to measure: which domains and URLs appear in the citations for your prompts, how often yours are among them, and which third-party pages keep winning citations you want. That earned-media list is your PR target list.

How it looks in Peec: Perplexity tracking lists the cited sources per prompt over time, so you see both your citation rate and the exact competing URLs. For a deeper Perplexity-specific setup, see my guide on [tracking brand mentions in Perplexity AI^{\[1\]}](#).

[GÖRSEL PLACEHOLDER 4: Peec AI Perplexity source tracking view]

How to Track Your Brand in Google AI Overviews

What is different on this channel: AI Overviews sit on top of Google's index and inherit its ranking systems, so your classic SEO strength influences your AIO presence more than on any other channel. Mentions, cited sources, and your underlying organic rankings move together here, which means you track three layers at once.

What to measure: mention rate in the AI Overview itself, which URLs the overview links, and how both trend over time against your organic rankings for the same queries. A page that ranks well but never gets pulled into the overview is a formatting and extractability problem, not an authority problem.

How it looks in Peec: AI Overviews tracking shows mention and source data per prompt with weekly and monthly rollups, so you can separate a real trend from Google's constant experimentation.

[GÖRSEL PLACEHOLDER 5: Peec AI Google AI Overviews tracking view]

How to Track Google AI Mode Rankings Over Time

AI Mode is not AI Overviews: it is a fully conversational search surface where the entire results page is a generated answer, it runs deeper query fanout, and there is no numbered ranking to fall back on. Tracking AI Mode means tracking inclusion, position within the answer, and cited sources over time, because "ranking" here is really answer share. Peec tracks AI Mode as its own channel with the same mention, position, and source breakdown, and daily runs matter because AI Mode answers shift with model updates rather than index updates. I compared the dedicated options in my [AI Mode rank tracker guide](#)^[2].

[GÖRSEL PLACEHOLDER 6: Peec AI AI Mode tracking view]

How to Track Gemini Rankings Over Time

What is different on this channel: Gemini blends model knowledge with Google Search grounding, and its answers lean harder on the model's internal brand associations than AI Overviews does. That makes it the channel where brand perception work shows up first.

What to measure: mention rate and position over time, plus sentiment, because Gemini tends to editorialize more than the other engines. Watch for divergence: if Gemini describes you differently than AI Overviews cites you, the model's stored view of your brand lags your current content.

How it looks in Peec: Gemini is tracked alongside the other channels with the same prompt panel, so cross-channel divergence is visible on one screen instead of five tabs.

[GÖRSEL PLACEHOLDER 7: Peec AI Gemini tracking view]

Step 4: Monitor Rankings and Visibility Over Time, Not in Snapshots

Every number in AI search is a sample from a distribution, so the trend line is the product, not the daily value. Three practices make time-series tracking honest.

Aggregate weekly, decide monthly. Daily data catches incidents, weekly averages reveal direction, and monthly comparisons justify decisions. Peec runs daily and rolls up to weekly and monthly views, which matches how the noise actually behaves.

Measure volatility per prompt. A prompt whose answer changes every run is telling you the engines have not settled on an answer for that question yet. Those are your best opportunities, because the citation slot is still contested.

Set alerts for steps, not wiggles. Configure alerting for sustained changes, like a visibility drop that holds for 3+ days or a competitor entering your top prompt, rather than every daily fluctuation. Alert fatigue kills tracking programs faster than bad data does.

Step 5: Benchmark Against Competitors With Share of Voice

Your own visibility rate can rise while your position in the market falls, and share of voice is the metric that catches it. The calculation: your brand's mentions divided by all mentions of your tracked competitor set on the same prompts, in the same period.

How to work with it:

- Build a topic-level leaderboard, not just a global number. You can lead in "AI visibility tools" and be invisible in "enterprise AI analytics" at the same time, and the global average hides that.
- Check competitor prompt overlap. The prompts where a rival appears and you do not are your gap list, ordered by prompt volume.
- Read share of voice against your actual market share. Punching below your market weight in AI answers is the clearest early warning this channel produces.

In Peec, competitor benchmarking is built into the same prompt panel, so the leaderboard updates daily with no separate setup per rival.

[GÖRSEL PLACEHOLDER 8: Peec AI competitor share of voice leaderboard]

Step 6: Track Citations and Sources: Which Pages AI Engines Actually Read

Citation tracking answers the question mention tracking cannot: which exact pages are shaping the answers. In the last 30 days, my fanout dataset logged over 1,100 citation-related sub-queries from AI engines, which tells you how central source selection is to how these systems compose answers.

Work the data in two layers:

Domain level. Which domains keep earning citations on your prompts? Sort them into three buckets: your own pages, earned media you influence (reviews, comparisons, communities), and competitor-controlled pages. The second bucket is usually the fastest lever, because a single well-placed third-party mention can win citations across dozens of prompts.

URL level. Which of your specific pages get cited, and for which prompts? This is where the on-page versus off-page decision rule lives: if your page is cited but your brand mention rate is low, work on brand authority off-page. If you are mentioned but never cited, fix extractability on-page: summaries, direct answers, tables, and clean structure.

Peec reports citations at both levels and splits owned versus earned sources, which turns the report directly into a task list. If data quality in citation reporting is your buying criterion, I compared vendors on exactly that in [which AEO insights company delivers the best data](#)^[3].

[GÖRSEL PLACEHOLDER 9: Peec AI citation and source report]

Step 7: Measure Sentiment Per Channel

Sentiment in AI search is channel-specific, and averaging it wastes the signal. The same brand routinely reads positive in AI Overviews, neutral in ChatGPT, and cautious in Gemini, because each system draws on different sources and different stored associations.

The workflow that finds the cause, not just the score:

- Track sentiment per channel weekly, and only investigate when one channel diverges from the others.
- Use theme analysis to see what the engines praise or criticize: pricing, support, reliability, and ease of use are the usual recurring themes.
- When a channel turns negative, pull the cited sources for the affected prompts. The negative language almost always traces to specific third-party pages, and that gives you a concrete PR or content target instead of a vague reputation problem.

Peec scores sentiment per prompt per channel and groups the drivers by theme, so the path from "Gemini sentiment dropped" to "this review page is the cause" is a two-click investigation.

[GÖRSEL PLACEHOLDER 10: Peec AI sentiment analysis view]

Step 8: Cover the Query Fanouts Behind Every Prompt

Query fanout is the hidden layer of AI search: when a user asks one question, the engine silently runs several of its own sub-queries to gather material, and your content competes in those sub-queries, not in the visible prompt. A prompt like "best way to track AI brand visibility" fans out into sub-queries about specific channels, metrics, tools, and pricing, and each fanout is a separate chance to get retrieved.

How to build a coverage matrix:

- Pull the real fanout queries for your top prompts. Peec surfaces these from live engine behavior, which is the feature I lean on most in my own research.
- Map each fanout query to the page that should answer it. Empty cells in that matrix are your content gap list, pre-validated by the engines themselves.
- Turn gaps into actions: a missing comparison section, a missing pricing answer, a missing definition. Fanout-level gaps are usually section-sized, not article-sized, so the fixes ship fast.

Most tools show you the scoreboard. This fanout-to-gap-to-action loop is where tracking stops being reporting and starts being optimization. The fanout layer shows you the game.

Step 9: Connect AI Visibility to Traffic and Revenue

AI visibility programs die in budget reviews when they cannot show money, and the measurement problem is real: standard analytics under-attributes AI traffic badly. LLM referrals often arrive stripped of referrer data, users copy answers without clicking, and agents fetch pages without firing JavaScript.

Three measurements that survive scrutiny:

- **Self-reported attribution.** Add "How did you hear about us?" to signup, and count the ChatGPT and AI answers. In my client work, self-reported AI attribution runs 2 to 5 times higher than what GA4 credits to AI referrals, and that gap is the under-attribution made visible.
- **LLM referral analysis.** Segment the referral traffic that does arrive from chatgpt.com, perplexity.ai, gemini.google.com, and copilot.microsoft.com, and watch its trend against your visibility rate. Direction matters more than magnitude here.
- **Prompt-to-revenue correlation.** Tag your high-intent prompts and compare visibility movements on them against pipeline in the following weeks. It is correlation, not attribution, and stated honestly it is still the most convincing chart in the quarterly review.

Peec pairs visibility data with referral and attribution views, which keeps the revenue conversation inside the same tool that tracks the mentions.

How to Run This at Scale: Teams, Cadence and Automation

The process above is one brand's loop. Running it across teams, clients, or business units changes the setup more than the method.

For In-House SEO Teams

Run a weekly rhythm: Monday metric review against the previous week, one experiment shipped per week, monthly deep-dive on citations and sentiment. Keep a simple RACI: SEO owns prompts and analysis, content owns extractability fixes, PR owns the earned-media citation targets. The most common failure is treating AI visibility as a side dashboard nobody owns.

For Agencies

Multi-client setup lives or dies on templates: a standard prompt-panel structure per client vertical, a standard weekly report, and white-label reporting so deliverables carry your brand. Peec's unlimited seats matter here because every client stakeholder can get direct dashboard access instead of waiting for the monthly PDF. I reviewed the agency-specific options in [the best AI visibility tools for marketing agencies](#)^[4].

For Enterprise Brands

Enterprise scale adds API access, BI integration, and governance: push visibility data into BigQuery so Looker dashboards sit next to revenue data, define per-brand and per-market workspaces, and set access by role. Seat metering is the silent rollout killer at this level, which is one reason unlimited-seat pricing changes enterprise adoption. I covered the full procurement checklist in my [enterprise AI search analytics tools comparison](#)^[5].

For Marketing Teams

Report three numbers to the C-level, always with previous-period deltas: share of voice against named competitors, visibility trend on the top 20 revenue prompts, and AI-attributed pipeline. Skip the per-channel detail in the executive view and keep it one screen. The moment the dashboard needs a walkthrough, it stops being read.

Methodology: How This Data Is Collected

Transparent methodology is what separates measurement from marketing, so here is how the numbers behind this guide work.

Tracking-tool data comes from scheduled prompt runs: each prompt in the panel is executed daily on each tracked channel, answers are parsed for brand mentions, positions, cited URLs, and sentiment, and results are aggregated into weekly and monthly views. Sampling matters because answers vary run to run, which is why single-run numbers are never quoted as findings in this guide, only period averages.

The fanout statistics referenced throughout come from my own research dataset of real sub-queries generated by ChatGPT and Perplexity over 30-day windows, collected from live engine behavior rather than vendor estimates. Channel coverage in the walkthrough is ChatGPT, Perplexity, Google AI Overviews, Google AI Mode, and Gemini; Claude and Copilot follow the same method where tracked. Where a number comes from client work, it is labeled as such and rounded, because point precision would overstate what sampled, non-deterministic systems can support.

Frequently Asked Questions

How do I check my brand's AI visibility for free?

Run a manual audit: pick 20 high-intent prompts, ask them on ChatGPT, Perplexity, and Google AI Overviews on two different days, and record mentions, positions, and cited URLs in a spreadsheet. It is enough to find your baseline and biggest gaps, and free trials of tracking tools like Peec AI extend the same check across more prompts automatically.

How often should I track AI search rankings?

Track daily, read weekly, decide monthly. AI answers change between runs, so daily tracking is needed to build reliable averages, but daily numbers alone are noise. Weekly aggregation reveals real direction, and monthly comparisons are stable enough to justify content and PR decisions. One-off monthly checks miss both incidents and trends.

What is a good AI visibility rate?

A good visibility rate depends on prompt type. On branded prompts, expect 80% or higher. On solution and category prompts you actively target, 30% or higher is strong. Below 10% on prompts you should own signals a real gap. Always read the rate against competitors' share of voice, because a rising rate in a rising market can still mean losing ground.

Can I track AI visibility across multiple countries and languages?

Yes, and you should if you sell internationally, because AI engines answer differently per language and market. Build separate prompt panels per language rather than translating one panel literally, since buyer phrasing differs. Platforms like Peec AI support multi-market tracking in one workspace, so per-country visibility, citations, and sentiment stay comparable.

How is tracking AI Overviews different from tracking AI Mode?

AI Overviews is a generated summary on top of classic Google results, so it inherits your organic rankings and updates with the index. AI Mode is a fully conversational surface with deeper query fanout and no ranked list underneath, so you track answer inclusion, position, and sources instead of rankings. They need separate tracking because they move independently.

Which AI platforms should I track first?

Start with ChatGPT, because it carries the largest assistant usage, then add Google AI Overviews and AI Mode for search-driven discovery. Add Perplexity if your audience skews technical or research-heavy, and Gemini for consumer reach in the Google ecosystem. B2B teams should also watch Copilot, since Microsoft ecosystem integration puts it in front of enterprise buyers.

Start Tracking Before Your Next Quarterly Review

Everything above compresses into one habit: fixed prompts, daily runs, weekly reading, and actions taken at the citation and fanout level rather than the vanity-metric level. Set the panel up once and the loop runs itself. If you want the channel-specific setups next, start with the [AI Mode rank tracker guide](#)^[2], the [Perplexity brand mention tracking guide](#)^[1], and the [full AI visibility tools comparison](#)^[6], then connect [Peec AI](#)^[7] to run the whole loop in one place.

Tracking brand visibility in AI search means tracking AI citations across every answer engine your buyers use. ChatGPT first, Perplexity and Gemini next, then local AI engines. AEO and GEO programs that skip this step optimize blind.

REFERENCES & SOURCES

- [1] tracking brand mentions in Perplexity AI — <https://metehan.ai/articles/track-brand-mentions-perplexity-ai>
- [2] AI Mode rank tracker guide — <https://metehan.ai/articles/best-ai-mode-rank-tracker>
- [3] which AEO insights company delivers the best data — <https://metehan.ai/articles/aeo-insights-company>
- [4] the best AI visibility tools for marketing agencies — <https://metehan.ai/articles/the-best-ai-visibility-tools-for-marketing-agencies>
- [5] enterprise AI search analytics tools comparison — <https://metehan.ai/articles/best-enterprise-ai-search-analytics-tools>
- [6] full AI visibility tools comparison — <https://metehan.ai/articles/best-ai-visibility-tools>
- [7] Peec AI — <https://peec.ai/>