

# Best AEO Tools in 2026: 22 Answer Engine Optimization Platforms Tested

Metehan Yesilyurt

*AI Search & Information Retrieval Researcher · metehan.ai*

Published: September 25, 2026 · Canonical: <https://metehan.ai/articles/best-aeo-tools/>

Ask ChatGPT who leads your category and it will name a short list of brands. Your name is either on it or it is not, and no rank tracker will tell you which. That gap is why the best AEO tools in 2026 exist. I tested 22 answer engine optimization platforms and ranked them on one question: do they show whether AI names you as the answer, and do they tell you how to get named when it does not.

PeeC AI takes the top spot for most teams, followed by Profound, AEO Vision, Scrunch AI, Conductor, AthenaHQ, and Rankscale. Every claim below is checked against peer-reviewed studies, vendor pricing pages, and certifications as of September 2026, not marketing decks. The stakes are measurable: a 2026 study of 11,500 real queries found [AI Overviews appear above the standard results 51.5% of the time<sup>\[1\]</sup>](#), and a separate study of 55,393 queries found they [trigger on 13.7% of searches and 64.7% of question-style searches<sup>\[2\]</sup>](#). When a question-shaped query is the one your buyer types, a ranking is not the prize. Being the answer is.

**One distinction runs through this whole guide.** Citation share tells you your page was used as a source. Answer share tells you your brand was named as the recommendation. The two move independently, and a tool that only reports the first is selling you half the picture. This is the split I keep coming back to in [100 questions that show AEO and GEO are different from SEO<sup>\[3\]</sup>](#).

## What AEO actually measures

AEO, or Answer Engine Optimization, is the practice of being the answer an AI assistant gives, not a source it quietly cites. The mechanics differ from SEO in a way that decides which tools are useful. A classic search engine matches a query to a page and returns a list. An answer engine assembles a response from several sources, decides which brands to name, and picks which pages to attribute. Answer engine optimization tools instrument that assembly process: they run prompts, record whether your brand appears, read how it is framed, capture the URLs behind the response, and compare all of it to competitors. If you want the fastest way to see where you currently stand, start with whether [ChatGPT recognizes your brand at all<sup>\[4\]</sup>](#), because a model that has no entity for you cannot recommend you.

## How I evaluated these 22 platforms

I weighted the criteria for AEO rather than for tracking in general, which changes the order. **Answer detection** comes first: does the tool distinguish being named from being cited? **Prompt discipline** second: can it build a prompt set from real questions instead of a keyword list, so you are not [guessing which prompts to track<sup>\[5\]</sup>](#)?

**Framing and sentiment** third, at dimension level rather than as one blurry score, because a wrong description of your brand is worse than silence. **Engine breadth** fourth, weighted toward the assistants your buyers actually use. **Actionability** fifth, so the tool produces a brief or a fix rather than a chart. And **measurement honesty** sixth, meaning it publishes sampling frequency, collection method, and a way to open the raw answer behind a metric.

## The best AEO tools in 2026

Here are the 22 platforms, ordered by how well they serve answer engine optimization specifically. Pricing and engine counts reflect what each vendor published as of September 2026, and where a vendor gates an engine or a feature to a higher tier, I say so.

| Tool                                 | One-line verdict                                                                               |
|--------------------------------------|------------------------------------------------------------------------------------------------|
| <b>Peec AI</b>                       | Named-versus-cited tracking across six answer engines, unlimited seats, \$95 entry.            |
| <b>Profound</b>                      | Category-defining AEO, Agents that act on gaps, nine engines and SOC 2 on Enterprise.          |
| <b>AEO Vision</b>                    | \$9 entry, citation-level page attribution, Reddit Insights and cited-domain mapping.          |
| <b>Scrunch AI</b>                    | SOC 2 Type II, real-time bot feed, \$250 Core, nine answer engines on Enterprise.              |
| <b>Otterly AI</b>                    | \$29 entry with AEO audit, four base engines, extras cost more.                                |
| <b>Conductor</b>                     | AgentStack inside ChatGPT and Claude, MCP and API, SOC 2 plus ISO 42001 and 27001.             |
| <b>AthenaHQ</b>                      | Free Essential tier, \$295 Starter with nine models, GA4 and GSC, Citation Engine.             |
| <b>Writesonic</b>                    | \$79 entry, Action Center that fixes pages in place, all 10 engines on Enterprise.             |
| <b>Goodie AI</b>                     | \$399 Explorer, up to 11 models on Enterprise, MCP access and revenue attribution.             |
| <b>Ahrefs Brand Radar</b>            | 464M+ search-backed prompts, ties citations to backlinks, \$199 per platform index.            |
| <b>Semrush AI Visibility Toolkit</b> | \$99 per domain, 289M+ prompts, Google-first coverage, one suite for SEO and AEO.              |
| <b>SE Ranking (SE Visible)</b>       | Real browser-collected answers, Basic \$99 and Core \$189, five engines, no Claude yet.        |
| <b>Rankscale</b>                     | From \$20 Essentials, 17+ engines on every plan, credits roll over, white label from Growth.   |
| <b>Gauge</b>                         | \$599 Growth, 108K answers a month, six engines, 18 articles, Claude via your own key.         |
| <b>Evertune</b>                      | \$800 Pro, 100 samples per prompt, EverPanel consumer panel, 100,000 prompts.                  |
| <b>Hall AI</b>                       | Acquired by Tracksuit in July 2026; Slack-first alerts, nine engines, sales-led pricing.       |
| <b>Gumshoe AI</b>                    | \$99 Starter, free sample audit, persona-level share of LLM, no annual contract.               |
| <b>Trakkr</b>                        | Prism serves AI-ready JSON-LD, Growth \$100, Scale \$500 with white-label portals.             |
| <b>Nightwatch</b>                    | Euro 79 entry, six engines included at no extra cost, geo data to city level, unlimited seats. |
| <b>SISTRIX</b>                       | Free during beta, euro 119 Start with 100 prompts, 25M prompt dataset, European depth.         |
| <b>Botify</b>                        | Server-log crawler analysis, GPTBot and PerplexityBot visibility, sales-quoted.                |
| <b>Omnia</b>                         | Euro 79 Growth, daily refresh, sentiment by dimension, Omnio agent ships the fixes.            |

## Peec AI: The Cleanest Answer-Share Tracker for Most Teams

*Peec AI dashboard*

[Peec AI](#)<sup>[6]</sup> answers the AEO question that matters first: is your brand being named in the answer, or only cited under it? It tracks that distinction across ChatGPT, Google AI Overviews, Google AI Mode, Perplexity, Gemini, and Microsoft Copilot, then shows how your framing shifts over time. Setup takes minutes and the dashboard is readable on day one, which matters because AEO reporting usually lands on a non-technical stakeholder's desk. The reason it holds the top spot here is the combination of breadth and usability at \$95 per month. That tier covers 50 prompts across three selectable engines per project, and unlimited seats are included on every plan, so a whole marketing team or a client roster fits on one license. Higher tiers add brand sentiment, competitor benchmarking, LLM prompt suggestions, Looker Studio reporting, and an Actions view that sorts your next moves into owned and earned media with a priority score. Enterprise reaches up to 11 models, adding Claude, GPT-5 Search, DeepSeek, Qwen, and Mistral. For AEO specifically, the sentiment and framing data is the part that earns its keep, because being named in a negative or wrong context is worse than not appearing at all. Annual billing saves about 15 percent.

## Profound: The AEO Platform That Named the Category

*Profound dashboard*

[Profound](#)<sup>[7]</sup> coined the term Answer Engine Optimization, and it is still the name most teams encounter first. It reads brand mentions, citations, answer share, sentiment, and ranking across AI answer engines, and it monitors crawler activity so you can see exactly which bots are pulling which pages. The entry tier is where the caveats start. At \$99 per month on Starter you get ChatGPT only, 50 tracked prompts, 100 Agent credits, and no data export, which is thin if your buyers also work in Perplexity or Google. The \$399 Growth tier is the first one worth calling complete for a small team: it adds Perplexity and Google AI Overviews, raises the limit to 100 prompts, analyzes roughly 9,000 responses a month, allows three seats, and unlocks CSV and JSON export. Enterprise is quoted rather than listed, and it is where the platform opens up: nine engines including Google AI Mode, Gemini, Microsoft Copilot, Grok, DeepSeek, and Claude, plus API access, SAML single sign-on, multi-company tracking, and SOC 2 Type II. Its credit-based Agents turn a gap into an action, and Prompt Volumes mines real user conversations so your prompt set mirrors buyer language. The company raised a \$96 million Series C at a reported \$1 billion valuation in February 2026, publishes the Profound Index weekly, and runs the Zero Click conference series. Agency Growth starts at \$99, and annual billing is advertised as two months free.

## AEO Vision: Citation-Level AEO From \$9 Per Month

*AEO Vision dashboard*

[AEO Vision](#)<sup>[8]</sup> is built around a question most trackers skip: which specific pages of yours get pulled into an answer, and how does the model frame your brand when it does? It reads ChatGPT, Perplexity, Gemini, Claude, Google AI Mode, and Google AI Overviews at that citation level rather than stopping at a mention count. Beyond

tracking, it adds competitor discovery, Reddit Insights that surface the exact threads AI engines cite, Top Cited Domains, AI task management, and workflow automation, so a content team can move from a finding to a brief without leaving the tool. The research-first framing is deliberate: it pairs measurement with prioritized next steps drawn from the citation patterns documented in this guide. Disclosure: I work on AEO Vision, so treat this as an informed but biased pick. Lite is \$9 per month for ChatGPT-only tracking, Solo is \$99 and covers five platforms, and Growth is \$299 with multi-brand workspaces. For a small team testing whether AEO is worth the investment, the \$9 door is the cheapest credible entry in this list.

### Scrunch AI: The Compliance-Ready AEO Option

*Scrunch AI dashboard*

[Scrunch AI](#)<sup>[9]</sup> is the pick when AEO has to survive a security review. It carries SOC 2 Type II certification, SAML and OIDC single sign-on, role-based access control, and audit logs, and it tracks answer visibility at the prompt level with persona, topic, funnel stage, and location labels. The Core plan is \$250 per month and covers 125 prompts, five site audits, five licenses, and four platforms: ChatGPT, Perplexity, Google AI Overviews, and Microsoft Copilot. Enterprise adds nine platforms including Gemini, Claude, Google AI Mode, Grok, and Meta AI, plus query fan-out analysis, multi-brand workspaces, Looker Studio, an MCP surface, and an Agent Experience Platform that can serve an AI-ready version of your site without a redesign. Its live bot feed is the feature AEO teams use most in practice: it shows GPTBot, PerplexityBot, and ClaudeBot hitting your pages in real time, which turns crawler access from a theory into an observable fact. Core has a 7-day trial and Enterprise is quoted by sales.

### Otterly AI: Cheap Prompt Tracking With an AEO Audit Attached

*Otterly AI dashboard*

[Otterly AI](#)<sup>[10]</sup> covers the baseline AEO workflow at the lowest price that still includes an audit. Every plan ships with ChatGPT, Google AI Overviews, Perplexity, and Microsoft Copilot, daily tracking, unlimited team members, multi-country support across 65 plus markets, and an AI prompt research tool. The catch is engine depth: Google AI Mode, Gemini, and Claude are paid add-ons, so the \$29 sticker is a starting point rather than the real spend. What you get for it is prompt generation from keywords, daily monitoring, instant alerts when visibility drops, and an audit that flags the technical blockers that stop engines from pulling you in, like missing schema and crawl problems. For a small PR or brand team that mainly needs early warning when the narrative shifts, that is enough. Plans are \$29 per month for 15 prompts, \$189 for Standard, and \$489 for Premium.

### Conductor: Enterprise AEO With Agents Inside the Assistants

*Conductor dashboard*

[Conductor](#)<sup>[11]</sup> treats AEO as an enterprise content operation rather than a monitoring dashboard. It joins content creation, AI visibility insight, site monitoring, and performance tracking, then puts the work surface inside the tools your team already uses. AgentStack adds native LLM apps inside ChatGPT, Claude, and Copilot, a Model

Context Protocol server, a Data API, and turnkey agents that carry a content team from insight to published, optimized content without switching context. Coverage spans nine AI engines across more than 160 countries. For procurement, it holds SOC 2, ISO 42001, and ISO 27001, which is the certification set enterprise and regulated buyers tend to require before an AI tool can touch their data. It is one of the most capable options in this list and also one of the most expensive. Pricing is usage-based rather than per seat, with Essentials, Growth, and Enterprise tiers quoted by sales.

### AthenaHQ: Nine Answer Engines With a Free Door and Real Actions

*AthenaHQ dashboard*

[AthenaHQ](#)<sup>[12]</sup> pairs answer-engine tracking with expert-led recommendations, so a team can see which content gaps to fill, which prompts to pursue, and which technical issues are blocking AI crawlers, then get alerted when the picture changes. Starter is \$295 per month, includes 3,600 credits where one credit equals one AI response, and covers nine models: ChatGPT, Perplexity, Google AI Overviews, Google AI Mode, Gemini, Claude, Copilot, and Grok, with more on request. A free Essential tier gives 300 credits across five models, which is one of the few genuinely usable free tiers in this category and a good way to test AEO before committing. A unified visibility score keeps reporting simple, and native GA4, Google Search Console, Shopify, Webflow, and Framer integrations connect answer visibility to traffic and revenue rather than leaving it as a vanity chart. Enterprise adds multi-region and multi-language tracking, the Athena Citation Engine, SSO, audit logs, persona targeting, and BI tool support.

### Writesonic: AEO Monitoring With an Action Center That Edits Pages

*Writesonic dashboard*

[Writesonic](#)<sup>[13]</sup> closes the gap that most AEO tools leave open. Instead of reporting that you are missing from an answer, its Action Center ranks the gaps by impact and can apply the fix directly to a WordPress, Webflow, Framer, or Cloudflare site. Self-serve annual plans start at \$79 per month for ChatGPT tracking, \$199 for ChatGPT plus Gemini and Google AI Overviews, and \$399 for 200 prompts and 600 answers a day. Enterprise covers all 10 engines including Perplexity, Claude, Copilot, Grok, DeepSeek, Meta AI, and Google AI Mode. It simulates AI queries by country, plans content against competitors, and an AI Visitors dashboard shows which pages AI bots crawl most, which pairs well with the AEO habit of checking crawler access. It also tracks ChatGPT Shopping and agentic commerce, which matters for product brands. The full Action Center and agentic workflows sit on Enterprise, with a limited Action Center trial on Growth.

### Goodie AI: The Broadcast-Reach AEO Option With 11 Models

*Goodie AI dashboard*

[Goodie AI](#)<sup>[14]</sup> is the AEO platform for teams that want the widest engine spread and are willing to pay for it. It pairs visibility monitoring with an Optimization Hub and an AEO content writer, so measurement and action live in the same place. Pricing is where it separates from cheaper rivals. Explorer is \$399 per month for ChatGPT,

Google AI Overviews, and Perplexity, capped at 100 prompts, 3,000 AI responses, 10 optimization actions, and three seats. Pro, quoted separately, widens the engine set to Gemini, Copilot, and Amazon Rufus and raises the prompt ceiling to 250. The top Enterprise tier reaches up to 11 models including Claude, Google AI Mode, Meta AI, DeepSeek, and Grok, the broadest spread among the mid-market tools reviewed here. Work from the tool arrives as ranked actions across owned content, earned media outreach, technical fixes, and social entity signals, each carrying a difficulty and an impact estimate, and the plan includes MCP server access plus revenue attribution through Google Analytics. Its case studies are unusually concrete for this category, though the entry price sits well above most competitors.

### Ahrefs Brand Radar: Answer Visibility Built on a Real Search Index

*Ahrefs Brand Radar dashboard*

[Ahrefs Brand Radar](#)<sup>[15]</sup> has an advantage most AEO tools cannot match: its prompt set comes from an actual search index rather than a synthetic list. It runs 464 million search-backed prompts through ChatGPT, Perplexity, Gemini, Copilot, Google AI Overviews, and Google AI Mode, plus Grok and Claude for custom prompts, and stores every response. Because those prompts expand from real keywords with People Also Ask and semantic fan-out, the coverage reflects how people actually phrase questions instead of how a vendor imagines they do. You get brand share of voice, citation data, estimated impressions, and topic gaps with almost no setup. Access costs \$199 per month per platform index, or \$699 per month for all platforms, which includes 2,500 custom prompt checks. The structural benefit is that it lives inside Ahrefs, so an AI citation can be traced back to the backlink and ranking profile of the page behind it. The tradeoff is depth of remediation: it is strong on data and light on telling you what to change.

### Semrush AI Visibility Toolkit: AEO Inside the SEO Suite You Already Run

*Semrush AI Visibility Toolkit dashboard*

[Semrush AI Visibility Toolkit](#)<sup>[16]</sup> Semrush added answer-engine tracking to the SEO platform most teams already pay for. The toolkit reports brand mentions, a zero-to-100 AI Visibility Score, sentiment, citations, share of voice, prompt research with volume and intent, competitor gaps, and an AI-readiness site audit. It draws on one of the largest prompt datasets in the category at 289 million plus LLM prompts. The base tier costs \$99 per month per domain, tracks 25 custom prompts daily, and covers ChatGPT, Google AI Overviews, Google AI Mode, Gemini, and Perplexity. The limitation to know before you buy is that Claude, Microsoft Copilot, DeepSeek, and Grok are reserved for the separate Enterprise AIO product, so teams that need them should factor in the upgrade. For an agency already running client SEO in Semrush, the appeal is one dashboard and one invoice rather than a new vendor relationship. Semrush One bundles SEO and AI visibility from \$199 per month.

### SE Visible: Answer Data Collected From Real Browsers

*SE Ranking (SE Visible) dashboard*

SE Ranking<sup>[17]</sup> SE Visible collects answers from real browser sessions rather than API simulations, which is the methodological difference that makes its numbers worth a second look. API-based collection can miss system prompts and user memory, so a browser read is closer to what a buyer actually sees. It covers ChatGPT, Gemini, Google AI Mode, Google AI Overviews, and Perplexity, and reports visibility score, share of voice, average position, competitor benchmarking, sentiment, prompt insights, and source analysis in one place. The interface is pre-built and readable, so you do not need a technical analyst to interpret it. The three tiers are Basic at \$99 per month for 200 prompts and three projects, Core at \$189 for 450 prompts and five projects, and Plus at \$355 for 1,000 prompts and 10 projects, with a 10-day trial and roughly 20 percent off annually. Claude is not covered yet, which is the main gap for teams standardizing on that engine.

### **Rankscale: Seventeen Engines With No Tier Gating**

*Rankscale dashboard*

Rankscale<sup>[18]</sup> makes a simple AEO promise: one plan, 17 plus engines, no per-engine or per-seat surcharges. Coverage spans ChatGPT, Gemini, Perplexity, Claude, DeepSeek, Mistral, Grok, Copilot, Google AI Overviews, and Google AI Mode across 240 plus countries and regions. Pricing is credit-based, which is where the arithmetic matters. Essentials starts at \$20 per month, Pro is \$99 for 1,200 credits, Growth is \$385 for 5,500, and Enterprise is \$780 for 12,000, with credits that roll over and can be topped up. A typical engine query costs 0.25 credits, while DeepSeek costs one and Claude costs two or more, so model the full prompt, engine, region, and schedule mix before buying rather than dividing the credit total by the prompt count. It reports a visibility score, average position, competitor benchmarking, page-level citation tracking, and page audits, plus white-label reports and a REST API from the Growth tier. It is the strongest value in this list for teams that want engine breadth without tier games.

### **Gauge: A Statistically Serious AEO Sample Plus Content**

*Gauge dashboard*

Gauge<sup>[19]</sup> is built around sample size, which is the quiet problem in AEO measurement. Its self-serve Growth plan runs 600 prompts daily, producing roughly 108,000 answers analyzed per month, enough that a share-of-voice movement reflects a real shift rather than a single lucky draw. At \$599 per month it also includes 18 publish-ready articles a month, an Ask Gauge agent and Action Center recommendations, exports, and five seats. It monitors ChatGPT, Google AI Overviews, Google AI Mode, Gemini, Perplexity, and Microsoft Copilot. Claude and Grok are available by connecting your own Anthropic or xAI API key, or through the custom Enterprise tier. The tier structure is the clearest example of gating in this category, so read the plan sheet against the engines you actually need rather than the headline. For a well-funded team that wants both measurement rigor and content output in one contract, it is a defensible pick.

### **Evertune: Repeated Sampling to Kill AEO Noise**

*Evertune dashboard*

[Evertune](#)<sup>[20]</sup> attacks AEO measurement error head-on by sampling every prompt 100 times per model across 11 plus AI models. That turns a single check into a stable score and tightens the margin of error enough that topic rankings become comparable. Its prompts come from EverPanel, a proprietary panel of more than 150 million real user conversations weighted to reflect the composition of the internet, so you are measured on the language buyers actually use rather than a guessed keyword list. The Pro plan is \$800 per month and includes 100,000 prompts, daily, weekly, or monthly tracking, global language and country coverage, and 25 AI-optimized articles a month. Enterprise adds SSO, AI website optimization, and bot analytics. It is the most expensive option here and the least suitable for a small team, but for a brand that needs to defend a share-of-voice number in a board meeting, the statistical footing is the point.

### **Hall AI: Slack-First Alerts, Now Inside Tracksuit**

*Hall AI dashboard*

[Hall AI](#)<sup>[21]</sup> took the position that AEO monitoring belongs where the team already talks. It built a Slack-first dashboard with citation heatmaps and alerts, and it added agent analytics showing how AI assistants and crawlers move through your site. It covered nine engines and analyzed how products surface inside ChatGPT conversations, which made it popular with content and mid-market teams that wanted fast feedback rather than a heavy platform. The structural change to know about is ownership: Tracksuit acquired Hall in July 2026, and AI visibility is now rolling into Tracksuit's brand tracking product rather than being sold as a standalone tool. Historically Hall offered a free Lite tier with one project and 25 tracked questions, with paid tiers reported around \$199, \$599, and \$1,499 plus per month. Current pricing on its site is sales-led, so if you were evaluating Hall specifically, the honest advice now is to evaluate Tracksuit.

### **Gumshoe AI: Persona Audits for a Pay-As-You-Go Price**

*Gumshoe AI dashboard*

[Gumshoe AI](#)<sup>[22]</sup> runs conversations as real buyer personas rather than as keywords, then reports a Share of LLM metric showing how often AI names your brand against competitors, along with citation sources, buyer criteria, and sentiment. That persona framing is useful in AEO because answer engines respond to how a question is phrased, so a keyword list can miss the way a buyer actually asks. It tracks 11 model families and runs each audit across six, with a monthly Visibility Audit plus weekly or daily Snapshot Audits in between. Published plans are Starter at \$99 per month and Pro at \$299, both with no annual contract, and there is a free sample audit covering four personas and four models before you pay. A pay-as-you-go option has been listed around \$0.10 per conversation, which is friendly for testing but can scale unpredictably at high volume, so set a ceiling if you go that route.

### **Trakkr: Making Your Pages Machine-Readable for Answer Engines**

*Trakkr dashboard*

[Trakkr](#)<sup>[23]</sup> starts from a different premise than most trackers. Before measuring, it makes your content easier for answer engines to consume: its Prism technology serves AI crawlers a machine-readable version of the page with JSON-LD schema, key facts, an auto-generated FAQ, an AI summary block, and entity recognition. That version is delivered through a Cloudflare Worker, so human visitors and traditional crawlers still see the normal page, which avoids the usual cloaking tradeoff. On the measurement side, every plan tracks eight models: ChatGPT, Claude, Gemini, Perplexity, DeepSeek, Grok, Meta AI, and Google AI Overviews, with 50 prompts per brand, daily refresh, competitor tracking, citation sources, and perception analysis. Growth costs \$100 per month for one brand and three seats. Scale costs \$500 for 10 brands and adds unlimited seats, the REST API, and white-label client portals. Enterprise starts at \$1,000. For an agency that wants the technical AEO work and the reporting in one place, the white-label portals at Scale are the differentiator.

### Nightwatch: Answer Visibility Beside Your Rank Tracking

*Nightwatch dashboard*

[Nightwatch](#)<sup>[24]</sup> combines classic rank tracking with answer-engine visibility in a single dashboard, and unlike most competitors it includes AI tracking in the plan rather than selling it as an add-on. Starter is euro 79 per month for 500 keywords, 50 AI prompts, and 1,500 AI responses, Professional is euro 159 for 150 prompts, and Agency is euro 399 for 500 prompts, all across ChatGPT, Claude, Gemini, Perplexity, Google AI Mode, and Google AI Overviews, with unlimited seats and no per-seat fees. Its distinguishing strength is geographic precision down to the ZIP code or city, which suits multi-location and local businesses where an answer can differ street by street. Its Citation Intelligence feature maps AI mentions back to the Google rankings that drove them, which is genuinely useful for AEO because it lets you fix the underlying ranking rather than only the symptom. For an SEO team that wants one view of both surfaces, this is the most economical route in this list.

### SISTRIX: European Answer Tracking, Free While in Beta

*SISTRIX dashboard*

[SISTRIX](#)<sup>[25]</sup> brought its European search dataset into AEO, and its AI Visibility features are free for all customers during the beta, which makes it the cheapest way to start on this list if you are already a customer. It tracks Google AI Overviews, AI Mode, ChatGPT, and Perplexity, with a daily AI Visibility Index for your brand and automatic competitor comparison. Its expanded SERP archive stores full AI Overview content, recognized entities, and historical changes, which is useful for the AEO habit of auditing how an answer evolved rather than only how it reads today. Prompt allowances scale with the package: 100 prompts a month on Start at euro 119, 300 on Plus at euro 239, 600 on Professional, and 1,500 on Premium, all with unlimited projects. Coverage leans on its own dataset of roughly 25 million prompts drawn from real user questions, so you do not have to build a prompt list from scratch, though its prompt tracker beyond Google is still maturing.

### Botify: Answer Visibility Grounded in Server Logs

*Botify dashboard*

[Botify](#)<sup>[26]</sup> connects AEO to the technical reality of your site. Its AI Visibility Dashboard tracks presence rate, citation share, sentiment, product citations, and prompt intent across ChatGPT, Perplexity, and Google AI Mode, using vetted browser scraping for ChatGPT and AI Mode because it captures system prompts and user memory that API-only tools miss. Its LogAnalyzer is the feature that separates it: it ingests your server logs daily and shows how often GPTBot, OAI-SearchBot, ChatGPT-User, PerplexityBot, ClaudeBot, and Meta-ExternalAgent crawl your site, then flags technical blockers automatically. In AEO practice this is the diagnostic that explains most invisible-brand problems, because if a bot cannot reach a page, no content strategy will surface it in an answer. The platform rewards teams who can act on log-file insight and first-party data. AI Visibility is an optional module on all Botify plans, LogAnalyzer is included, and both are sales-quoted, with third-party estimates around \$500 plus per month.

## Omnia: Daily AEO Tracking With an Agent That Ships Work

*Omnia dashboard*

[Omnia](#)<sup>[27]</sup> is aimed at small teams at scaleups who need fixes rather than a sixth dashboard. It tracks ChatGPT, Perplexity, Google AI Overviews, and Google AI Mode with daily refresh and country and language targeting, and can unlock Claude, Microsoft Copilot, and Gemini on higher tiers. The reporting goes further than a single brand score: it shows the exact sources behind each answer, stores full response snapshots, and breaks sentiment down by dimension such as pricing, product, and support, which is the level of detail an AEO team needs to know what to change. Its action layer studies winning citations and produces a brief covering what to build, which prompts to target, and which sites to pitch. The Omnia agent then drafts and ships the work with your approval. Growth is euro 79 per month for 25 prompts, Pro is euro 279 for 100 prompts, and Enterprise starts at euro 499 for 200 plus prompts, with a 14-day trial and no credit card required.

## Answer share versus citation share: the metric that decides your AEO plan

Most AEO dashboards lead with a single visibility number, and most of them blend two things that should stay separate. **Citation share** is the proportion of answers that pulled from your pages. **Answer share** is the proportion that named your brand. A 2026 measurement framework draws exactly this line when it separates [citation selection from citation absorption](#)<sup>[28]</sup>, and finds that the pages a model actually leans on tend to be longer, better structured, and dense with extractable evidence like definitions, numbers, comparisons, and steps.

The practical consequence is that the two metrics call for different work. High citation share with low answer share means engines trust your pages as sources but do not consider your brand the recommendation. The fix is owned and earned positioning, not more content: clearer entity definition, comparison pages that name you against competitors, and third-party mentions that establish you as the obvious answer. Low citation share with low answer share means the retrieval problem comes first, and the work is structural. Reading the two numbers separately tells you which job you are actually doing this quarter, which is why I rank answer detection ahead of every other feature in this list.

| Tool               | Answer detection | Framing depth   | Why it stands out               | Entry price |
|--------------------|------------------|-----------------|---------------------------------|-------------|
| Peec AI            | Named and cited  | Yes, per prompt | Unlimited seats, Actions view   | \$95/mo     |
| Profound           | Named and cited  | Enterprise only | Agents act on gaps              | \$99/mo     |
| AEO Vision         | Named and cited  | Growth tier     | Page-level citation attribution | \$9/mo      |
| Scrunch AI         | Named and cited  | Enterprise      | Real-time bot feed              | \$250/mo    |
| Otterly AI         | Named and cited  | Premium         | Audit plus alerts               | \$29/mo     |
| Conductor          | Named and cited  | Enterprise      | AgentStack and MCP              | Custom      |
| AthenaHQ           | Named and cited  | Enterprise      | Citation Engine, GA4 and GSC    | \$295/mo    |
| Writesonic         | Named and cited  | Enterprise      | Action Center edits pages       | \$79/mo     |
| Goodie AI          | Named and cited  | Enterprise      | Optimization Hub, 11 models     | \$399/mo    |
| Ahrefs Brand Radar | Cited first      | Not offered     | Backlink traceability           | \$199/mo    |
| Semrush Toolkit    | Cited first      | Enterprise AIO  | One suite with SEO              | \$99/mo     |
| SE Visible         | Named and cited  | Not in scope    | Real browser answers            | \$99/mo     |
| Rankscale          | Named and cited  | Growth tier     | No engine gating, 17+ models    | \$20/mo     |
| Gauge              | Named and cited  | Enterprise      | 108K answers per month          | \$599/mo    |
| Evertune           | Named and cited  | Enterprise      | 100 samples per prompt          | \$800/mo    |
| Hall AI            | Named and cited  | Via Tracksuit   | Slack-first alerts              | Sales-led   |
| Gumshoe AI         | Named and cited  | Not offered     | Persona-level share of LLM      | \$99/mo     |
| Trakkr             | Named and cited  | Scale tier      | Prism JSON-LD for crawlers      | \$100/mo    |
| Nightwatch         | Named and cited  | Not offered     | City-level geo accuracy         | Euro 79/mo  |
| SISTRIX            | Named and cited  | Not offered     | 25M prompt dataset              | Euro 119/mo |
| Botify             | Cited first      | Not offered     | Server-log bot analysis         | Custom      |
| Omnia              | Named and cited  | Enterprise      | Sentiment by dimension          | Euro 79/mo  |

## Building a prompt set that measures answers, not keywords

AEO lives or dies on the prompts you track, and a keyword list is the wrong starting point. Answer engines respond to phrasing, context, and implied intent, so the same underlying need can produce a different answer depending on how it is asked. The workable method is to stratify prompts into four bands. **Category prompts** ask who leads a space. **Comparison prompts** ask which option is better for a situation. **Problem prompts** describe a symptom and ask what to do. And **Brand prompts** ask directly about you, which tests framing rather than discovery. Then seed each band from real data: Search Console queries you already serve, sales-call language, and community threads, which is the approach behind [using real query data to pick prompts](#) <sup>[5]</sup>.

Two refinements matter in practice. First, run each prompt across every engine you track rather than assuming ChatGPT stands in for the rest, because the sources behind ChatGPT, Perplexity, and Google [overlap very little](#) <sup>[1]</sup>. Second, keep the set small enough to review weekly. A 40-prompt set you actually read beats a 400-prompt set

you never open. If you want a deeper method, I have written up [how I reached number one in SearchGPT<sup>\[29\]</sup>](#) and the [RRF Top-n approach for earning citations in ChatGPT's web mode<sup>\[30\]</sup>](#).

## Schema and entity readiness for answer engines

Answer engines do not read your page the way a person does, and structured data is how you tell them what the page is about without ambiguity. It is not a ranking trick and it will not rescue thin content, but it removes the interpretive guesswork that stops a model from using an otherwise good page. The table below maps the schema types that correspond to the answer shapes engines produce most often.

| Schema type                     | Who needs it              | What it feeds                                       |
|---------------------------------|---------------------------|-----------------------------------------------------|
| Organization and sameAs         | Every brand               | Establishes the entity answer engines match against |
| Product and SoftwareApplication | SaaS, ecommerce           | Feeds comparison and list answers                   |
| FAQPage                         | Content and support pages | Maps directly to question-shaped prompts            |
| Article and author              | Publishers, experts       | Supports the expertise signals engines weigh        |
| HowTo                           | Product docs, tutorials   | Feeds step-by-step answers                          |
| Review and AggregateRating      | Ecommerce, local          | Influences recommendation answers                   |
| BreadcrumbList                  | All sites                 | Clarifies site structure for retrieval              |
| Person and sameAs               | Consultants, authors      | Ties a name to an entity across sources             |

Two things are worth doing before you invest in schema tooling. First, confirm the crawlers can reach your pages at all, because a blocked bot makes every other AEO effort irrelevant, and the [SIGIR study is explicit that blocking AI crawlers reduces retrieval<sup>\[1\]</sup>](#). Second, make sure the entity you are marking up is consistent everywhere: same name, same description, same links across your site, your profiles, and third-party mentions. Answer engines resolve entities across sources, so inconsistency reads as uncertainty. Tools like Trakkr automate the delivery layer by serving an AI-ready version of the page, while others expect you to handle markup yourself.

## Which answer engines each tool covers

Coverage is where plans quietly differ, so read this table against the assistants your buyers open. ChatGPT is universal and therefore rarely decisive. Perplexity and Gemini are where the mid-market field separates. Claude, Grok, Meta AI, and DeepSeek are the engines most often gated to a top tier, and Amazon Rufus appears only where ecommerce answers matter.

| Answer engine       | Coverage     | Tools                                                                                                                                                                 |
|---------------------|--------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| ChatGPT             | All 22 tools | The baseline. Absence here is disqualifying.                                                                                                                          |
| Google AI Overviews | 18 tools     | Peec, Profound, AEO Vision, Scrunch, Otterly, Ahrefs, Semrush, SE Visible, Rankscale, Gauge, Nightwatch, SISTRIX, Botify, Trakkr, Omnia, AthenaHQ, Writesonic, Goodie |
| Perplexity          | 17 tools     | Rewards long, well-sourced pages that read as citable.                                                                                                                |
| Google AI Mode      | 12 tools     | Frequently gated to a higher plan.                                                                                                                                    |
| Gemini              | 12 tools     | Add-on at Otterly AI, Enterprise at Profound.                                                                                                                         |
| Claude              | 10 tools     | Usually the first engine cut from a cheaper tier.                                                                                                                     |
| Microsoft Copilot   | 9 tools      | Often bundled into base plans.                                                                                                                                        |
| Grok                | 9 tools      | Newer, and commonly on the top plan.                                                                                                                                  |
| Meta AI             | 4 tools      | Rare outside enterprise tiers.                                                                                                                                        |
| DeepSeek            | 6 tools      | Costs more per query on credit plans.                                                                                                                                 |
| Amazon Rufus        | 2 tools      | Only relevant for product and ecommerce answers.                                                                                                                      |

If your buyers concentrate in one assistant, buy the cheapest plan that covers it well rather than paying for breadth. If they are spread across several, engine count becomes the deciding column, and a plan that gates Claude or Grok is a plan you will outgrow. Ask each vendor for the full engine list in writing, because plenty of tools cover fewer engines and never publish which.

## What an AEO tool costs, and why the sticker price misleads

Four billing models are in play, and the mismatch between them is why two tools with the same headline price can cost very different amounts at the same volume. Prompt-based plans charge per question tracked. Credit-based plans charge per AI response, so the engine mix moves the bill: on Rankscale a typical query costs 0.25 credits, DeepSeek costs one, and Claude costs two or more, while Profound Agents burn a credit each time they run. Per-domain plans like Semrush charge \$99 per domain, which multiplies across a client roster. Per-seat plans punish large teams, which is why unlimited seats at Peec AI, Nightwatch, Trakkr, and Temso are worth real money, and why Profound routes teams above three people to a custom quote. Add overage rates, extra domains, and the seats you will need in six months before comparing anything.

The counterintuitive advice is that the cheapest credible entry is usually not the cheapest plan. AEO Vision at \$9 and Rankscale at \$20 are excellent ways to test whether AEO matters for your category, and AthenaHQ's free Essential tier is genuinely usable. But most teams outgrow a single-engine tier within a quarter, and the upgrade usually costs more than starting one band higher. Where AEO is already a line item, Peec AI at \$95 is the most capability per dollar in this list because six engines, unlimited seats, sentiment, and an action layer arrive together. Where the buyer is an agency, the decisive feature is not price at all: it is unlimited seats and white-label reporting, which is where Peec AI, Trakkr Scale, Rankscale Growth, and Scrunch AI separate from the field.

| Budget band                 | Tools                                                                                                            | What to watch                                          |
|-----------------------------|------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------|
| Test the water              | AEO Vision Lite \$9, Rankscale Essentials \$20, Otterly AI \$29, AthenaHQ Essential free, SISTRIX beta free      | Confirm which engines the cheap tier actually reaches  |
| One person, one brand       | Peec AI \$95, Profound Starter \$99, SE Visible Basic \$99, Trakkr Growth \$100, Nightwatch euro 79              | Check seat limits before you invite the team           |
| Small team, several engines | Profound Growth \$399, Writesonic Growth \$399, Goodie Explorer \$399, Scrunch Core \$250, SE Visible Core \$189 | Multi-engine coverage starts in this band              |
| Agency, many clients        | Trakkr Scale \$500, Peec AI agency \$245, Scrunch agency \$500, Rankscale Growth \$385, SE Visible Plus \$355    | Prioritize unlimited seats and white label             |
| High-volume monitoring      | Gauge \$599, Evertune Pro \$800, Rankscale Enterprise \$780                                                      | You are buying statistical confidence, not features    |
| Enterprise procurement      | Conductor, Profound Enterprise, Botify, Writesonic Enterprise, AthenaHQ Enterprise                               | SSO, SOC 2, ISO, audit logs, and data retention decide |

## How to tell whether an AEO tool is measuring honestly

Two tools can report the same score and mean different things. Four questions expose the difference. How many times is each prompt sampled? Evertune samples 100 times per model, which tightens the margin of error to roughly a point, while most tools sample once a day. Where does the data come from? SE Visible, Peec AI, Promptwatch, and Gauge collect from real browser sessions, arguing that API results miss system prompts and user memory, while others rely on API calls. Are the prompts real questions or synthetic? Ahrefs expands its set from a live search index, while many vendors generate prompts internally. And is variability acknowledged? The research makes the case that you should not measure once<sup>[31]</sup> and that visibility numbers carry statistical uncertainty that demands repeated sampling<sup>[32]</sup>.

A practical test takes ten minutes: run one prompt inside the tool and by hand on the same day, then compare both the naming and the cited sources. If the tool cannot show you the raw answer behind a metric, or cannot state its sampling frequency, treat its score as directional rather than comparable. This matters most when you are reporting to someone who will make a budget decision on the number.

## Which AEO tool fits your team

The right pick depends on who is doing the work and who reads the output. A solo marketer optimizing a single brand needs a cheap prompt set and plain language. An agency needs unlimited seats and a client-ready report. An enterprise needs certifications and an audit log. The table below maps the common cases.

| Team                     | Shortlist                                           | What decides it                                  |
|--------------------------|-----------------------------------------------------|--------------------------------------------------|
| Solo marketer or founder | AEO Vision Lite, Rankscale Essentials, Otterly AI   | Start cheap, upgrade when a prompt starts paying |
| In-house SEO team        | Peec AI, Nightwatch, SE Visible, Semrush Toolkit    | One dashboard beats two vendors                  |
| Content team             | Writesonic, Goodie AI, Trakkr, AthenaHQ             | The action layer matters more than the chart     |
| PR and communications    | Otterly AI, AEO Vision, AthenaHQ                    | Sentiment and narrative framing                  |
| Agency with clients      | Trakkr Scale, Peec AI, Rankscale Growth, Scrunch AI | Unlimited seats and white-label reporting        |
| Enterprise and regulated | Conductor, Profound, Scrunch AI, Botify             | Certifications, SSO, audit logs                  |
| Ecommerce and product    | Goodie AI, Writesonic, Trakkr                       | Product citations and Rufus coverage             |
| Local and multi-location | Nightwatch, Semrush Toolkit, SE Visible             | City-level answer differences                    |

Two rules cover most situations. If you report to clients, unlimited seats and white label outrank engine count. If you report to a board or a procurement team, certifications and SSO outrank everything, which is the ground Conductor, Profound, Scrunch AI, and Botify compete on. Everyone else should buy the smallest plan that reaches their engines and includes an action layer, because a dashboard that tells you what is wrong without telling you what to do rarely survives its second renewal.

## Integrations, data access, and security for AEO stacks

An AEO tool eventually has to feed something. Check for API access, an MCP server so assistants can query your own data, Looker Studio connectors, Google Search Console and GA4 integrations to tie answer visibility to traffic, Slack alerts, webhooks, CSV export, and white-label reporting. Conductor, Profound, Scrunch AI, Peec AI, Trakkr, ZipTie, Temso, Promptwatch, Knowatoa, and Goodie publish an API or MCP surface, and AthenaHQ and ZipTie connect Search Console natively. On the security side, the questions that appear in a real procurement review are SOC 2 Type II, ISO 27001 and ISO 42001, SSO with SAML or OIDC, role-based access control, audit logs, GDPR and CCPA coverage, data retention, and whether the vendor sells personal data. Scrunch AI publishes SOC 2 Type II with SSO, RBAC, and audit logs, Conductor holds SOC 2, ISO 42001, and ISO 27001, Profound is SOC 2 Type II certified with SSO and SAML on Enterprise, and Writesonic lists SSO, SOC 2 Type II, HIPAA, and GDPR on Enterprise.

## What the research says about getting named in an answer

The evidence points in a consistent direction, and most of it comes from outside the vendor world.

Earned sources carry the answer. A peer-reviewed 2025 study found a [systematic bias toward earned, third-party, authoritative sources over brand-owned pages and social posts](#)<sup>[33]</sup>. Vendors see the same pattern from the other side: Profound's analysis of 27 million prompts reported that 97.4% of AI citations trace to non Tier-1 earned media like Reddit threads and niche videos. If you want to be named, the mentions that count are often not on your domain.

Citation is not the same as naming. The [citation selection and absorption framework](#)<sup>[28]</sup> shows that pages which actually shape an answer are longer, better structured, and full of extractable facts, which is why comparison and definition pages over-perform. [Kevin Indig's correlation work](#)<sup>[34]</sup> adds that the engines reward different things:

Perplexity leans longer, while ChatGPT leans on domain trust and readability. Treating them as one channel loses the signal.

Repetition and recency shape who gets named. AI answers are not stable, which is why a single check is a poor basis for any decision, and it is the reason [token economics and log probabilities explain so much of what you observe<sup>\[35\]</sup>](#). I confirmed this from both directions, which is the point of [finding it in the code and watching science prove it in the lab<sup>\[36\]</sup>](#).

## How to get named, not just tracked

---

Start with the prompts where you are cited but not named, because that is the cheapest ground to take. Those answers already trust your pages, so the work is positioning: sharpen your entity definition, publish comparison pages that name you against the alternatives, and earn third-party mentions that establish you as the recommendation rather than one of several sources. Then move to prompts where you are absent entirely, and diagnose before you write. If your pages are not crawled, fix access. If they are crawled but not retrieved, fix structure and evidence density. If they are retrieved but not named, fix positioning. Each failure has a different remedy, and the [Google AI Overviews case study<sup>\[37\]</sup>](#) walks through the diagnostic order I use.

## AEO questions buyers actually ask

---

### What is the difference between AEO and GEO?

AEO, or Answer Engine Optimization, targets being named as the answer an AI assistant gives. GEO, or Generative Engine Optimization, targets how generative engines describe and cite your brand across synthesized responses. They overlap heavily in practice, and most tools serve both, but the metrics differ: AEO leans on answer share and framing, while GEO leans on citation behavior and source selection across generative surfaces.

### Do I need an AEO tool if I already track rankings?

Yes, because a ranking and an answer are different outcomes. Rank tracking tells you where your page sits in a list. An AEO tool tells you whether an assistant names your brand when someone asks a question. The 2026 data shows AI Overviews appearing above standard results on 51.5% of tested queries, so the answer is frequently the first thing a buyer sees, and no rank report captures it.

### What is answer share and how is it different from citation share?

Answer share is the proportion of AI answers that name your brand. Citation share is the proportion that used your pages as a source. A page can be cited without your brand being named, which is why a single blended visibility score hides the most useful signal. Read the two separately and they tell you whether your problem is retrieval or positioning.

### How many prompts should I track?

Fewer than you think, tracked consistently. A 40-prompt set stratified across category, comparison, problem, and brand prompts, reviewed weekly, produces better decisions than a 400-prompt set nobody opens. Seed the prompts

from Search Console queries, sales-call language, and community threads rather than from a keyword tool, because answer engines respond to phrasing and implied intent.

### **Which AEO tool tracks the most answer engines?**

Rankscale covers 17 plus engines without gating them by tier. Goodie AI reaches up to 11 on Enterprise including Amazon Rufus, Meta AI, and DeepSeek, and Profound reaches nine including Google AI Mode, Gemini, Copilot, Grok, DeepSeek, and Claude. Many tools cover fewer engines and do not publish the list, so ask for it in writing before you commit.

### **Is there a free AEO tool worth using?**

AthenaHQ's Essential tier gives 300 credits across five models and is genuinely usable, and SISTRIX AI Visibility is free for existing customers while in beta. AEO Vision starts at \$9 and Rankscale Essentials at \$20, which is low enough to test fit properly. Free tiers are for validating that AEO matters in your category, not for ongoing competitive tracking.

### **Do AEO tools work for ecommerce and product brands?**

They can, with two additions. First, confirm coverage of Amazon Rufus, which only Goodie AI and Writesonic track, since product answers often route through it. Second, check for product citation tracking, which Botify and Goodie expose and most brand trackers do not. For most product brands the useful prompts are comparison and recommendation shaped, not category-definition shaped.

### **How do I know an AEO tool is measuring honestly?**

Four checks. Does it state how many times each prompt is sampled? Does it say whether data comes from a live interface or an API? Are the prompts drawn from real questions or generated internally? And can you open the raw answer behind a metric? Tools that answer all four are auditable. Tools that report a single composite score without methodology are not comparable to anything.

### **What schema matters most for AEO?**

Organization with sameAs links, because it establishes the entity answer engines match against, and FAQPage on question-shaped content, because it maps directly to how prompts are phrased. Product and SoftwareApplication markup feeds comparison answers, HowTo feeds procedural answers, and Person with sameAs ties an expert name to an entity. Schema removes ambiguity, but it cannot compensate for thin content or a blocked crawler.

### **Which AEO tool is best for an agency?**

The deciding features are unlimited seats and white-label reporting. Peec AI includes unlimited seats on every plan, Trakkr Scale adds white-label client portals from \$500 per month, Rankscale Growth adds white-label reporting and a REST API, and Scrunch AI has agency tiers from \$500 per month with compliance built in. Engine count matters less than whether you can put 30 client brands and five team members on one license.

### **Do AEO tools meet enterprise security requirements?**

Procurement usually narrows the field to four names. Scrunch AI publishes SOC 2 Type II alongside SAML and OIDC single sign-on, role-based access control, and audit logs. Conductor holds SOC 2 plus ISO 42001 and ISO 27001, which covers AI governance as well as general security. Profound is SOC 2 Type II certified and adds

SAML SSO on Enterprise. Writesonic lists SSO, SOC 2 Type II, HIPAA, and GDPR on its Enterprise tier. Before signing, request the trust center, the data retention policy, and the subprocessor list, since those three documents decide most real objections from legal and security review.

## Sources and further reading

---

- Aggarwal et al., [GEO: Generative Engine Optimization](#)<sup>[38]</sup>, arXiv:2311.09735 (KDD 2024).
- Chen, Koudas et al., [Generative Engine Optimization: How to Dominate AI Search](#)<sup>[33]</sup>, arXiv:2509.08919.
- Xu, Iqbal, Montgomery, [Measuring Google AI Overviews: Activation, Source Quality, Claim Fidelity, and Publisher Impact](#)<sup>[2]</sup>, arXiv:2605.14021.
- Grossman et al., [How Generative AI Disrupts Search: An Empirical Study of Google Search, Gemini, and AI Overviews](#)<sup>[1]</sup>, arXiv:2604.27790 (SIGIR 2026).
- Zhang, He, Yao, [From Citation Selection to Citation Absorption: A Measurement Framework for GEO Across AI Search Platforms](#)<sup>[28]</sup>, arXiv:2604.25707.
- [Don't Measure Once: Measuring Visibility in AI Search \(GEO\)](#)<sup>[31]</sup>, arXiv:2604.07585.
- [Quantifying Uncertainty in AI Visibility: A Statistical Framework for Generative Search Measurement](#)<sup>[32]</sup>, arXiv:2603.08924.
- Kevin Indig, [What content works well in LLMs](#)<sup>[34]</sup>, Growth Memo.
- Microsoft Advertising, [From Discovery to Influence: A Guide to AEO and GEO](#)<sup>[39]</sup>, January 2026.

The best AEO tools share one trait: they report answer share and citation share as separate numbers, then point at the work that closes the gap. Everything else, from engine count to dashboard polish, is secondary to that.

---

## REFERENCES & SOURCES

- [1] AI Overviews appear above the standard results 51.5% of the time — <https://arxiv.org/abs/2604.27790>
- [2] trigger on 13.7% of searches and 64.7% of question-style searches — <https://arxiv.org/abs/2605.14021>
- [3] 100 questions that show AEO and GEO are different from SEO — <https://metehan.ai/blog/100-questions-that-show-aeo-geo-is-different-than-seo/>
- [4] ChatGPT recognizes your brand at all — <https://metehan.ai/blog/does-chatgpt-recognize-your-brand/>
- [5] guessing which prompts to track — <https://metehan.ai/blog/stop-guessing-which-prompts-to-track/>
- [6] Peec AI — <https://peec.ai/>
- [7] Profound — <https://www.tryprofound.com/>
- [8] AEO Vision — <https://aeovision.ai/>
- [9] Scrunch AI — <https://scrunch.com/>
- [10] Otterly AI — <https://otterly.ai/>
- [11] Conductor — <https://www.conductor.com/>
- [12] AthenaHQ — <https://athenahq.ai/>
- [13] Writesonic — <https://writesonic.com/generative-engine-optimization-geo>
- [14] Goodie AI — <https://higoodie.com/>
- [15] Ahrefs Brand Radar — <https://ahrefs.com/brand-radar>
- [16] Semrush AI Visibility Toolkit — <https://www.semrush.com/ai-seo/overview/>
- [17] SE Ranking — <https://visible.seranking.com/>
- [18] Rankscale — <https://rankscale.ai/>
- [19] Gauge — <https://www.withgauge.com/>
- [20] Evertune — <https://www.evertune.ai/>
- [21] Hall AI — <https://usehall.com/>
- [22] Gumshoe AI — <https://gumshoe.ai/>
- [23] Trakkr — <https://trakkr.ai/>
- [24] Nightwatch — <https://nightwatch.io/>
- [25] SISTRIX — <https://www.sistrix.com/ai/>
- [26] Botify — <https://www.botify.com/>
- [27] Omnia — <https://www.useomnia.com/>
- [28] citation selection from citation absorption — <https://arxiv.org/abs/2604.25707>
- [29] how I reached number one in SearchGPT — <https://metehan.ai/blog/how-to-rank-1-in-searchgpt-a-step-by-step-guide/>
- [30] RRF Top-n approach for earning citations in ChatGPT's web mode — <https://metehan.ai/blog/the-rrf-top-n-playbook-how-to-get-cited-by-chatgpts-web-mode/>
- [31] not measure once — <https://arxiv.org/abs/2604.07585>
- [32] statistical uncertainty that demands repeated sampling — <https://arxiv.org/abs/2603.08924>
- [33] systematic bias toward earned, third-party, authoritative sources over brand-owned page... — <https://arxiv.org/abs/2509.08919>
- [34] Kevin Indig's correlation work — <https://www.growth-memo.com/p/what-content-works-well-in-llms>
- [35] token economics and log probabilities explain so much of what you observe — <https://metehan.ai/blog/token-economics-log-probabilities-ai-search-visibility/>
- [36] finding it in the code and watching science prove it in the lab — <https://metehan.ai/blog/i-found-it-in-the-code-science-proved-it-in-the-lab-the-recency-bias-thats-reshaping-ai-search/>
- [37] Google AI Overviews case study — <https://metehan.ai/blog/how-to-optimize-your-website-for-ai-overviews-google-cloud-case-study/>

[38] **GEO: Generative Engine Optimization** — <https://arxiv.org/abs/2311.09735>

[39] **From Discovery to Influence: A Guide to AEO and GEO** — <https://about.ads.microsoft.com/content/dam/sites/msa-about/global/common/content-lib/pdf/from-discovery-to-influence-a-guide-to-aeo-and-geo.pdf>