

FREE PLAYBOOK

The AI Visibility Audit

Are you cited in **ChatGPT, Perplexity, Gemini,**
and the dozen engines routing traffic away from Google?

INSIDE THIS PLAYBOOK

ChatGPT

Perplexity

Gemini

AI Overviews

Copilot

Claude

Grok

Brave AI

AI BrandFactory

Table of Contents

Who this is for

What is in this playbook

Part 1, Why AI Visibility Needs a Formal Audit in 2026

Part 2, The AI Engines List: 15 Surfaces and the Manual Search Patterns for Each

Part 3, Automated Monitoring Tools Ranked by Use Case

Part 4, The Metrics and the Dashboard

Part 5, Diagnosis: When You Are Not Cited, Why Not and How to Fix

Part 6, The Weekly Audit Workflow and 8 Common Mistakes

Built by AI BrandFactory

Are you cited in ChatGPT, Perplexity, Gemini, and the dozen engines routing traffic away from Google? Most professional services do not know. This is the audit.

Who this is for

Three vertical groups, one shared blind spot.

- **Medical practices, clinics, dental, behavioral health, aesthetics, DPC.** Patients increasingly start health questions in ChatGPT or Perplexity before ever opening Google. If the AI engines are not citing you when they answer "best [your specialty] in [your city]", your funnel is leaking before the first click.
- **Law firms, solo attorneys, multi-practice firms.** Prospective clients now ask ChatGPT or Claude for help understanding their legal situation before contacting an attorney. The firm cited in those conversations gets the call.
- **Home service businesses, contractors, trades, multi-location service brands.** Local search has fragmented across Google AI Overviews, Perplexity, Gemini, and Bing Copilot. If one engine cites your competitor and not you, you lose that traffic permanently because customers do not switch engines once they have an answer.

The fragmentation problem is the same across all three. The audit is the same. The fixes (covered in Part 5 with cross-references) differ by what is broken.

What is in this playbook

Six parts.

- **Part 1, Why AI Visibility Needs a Formal Audit in 2026.** What changed, why per-engine measurement is now necessary, the cost of staying invisible.
- **Part 2, The AI Engines List: 15 Surfaces and the Manual Search Patterns for Each.** The complete inventory with URL, query patterns, what visibility looks like on each engine, known limitations.
- **Part 3, Automated Monitoring Tools Ranked by Use Case.** Eight tools that scale the manual audit, with pricing and use-case ranking.

- **Part 4, The Metrics and the Dashboard.** Share of AI Voice, citation count, sentiment, competitor benchmark, trend over time.
- **Part 5, Diagnosis: When You Are Not Cited, Why Not and How to Fix.** The diagnosis flowchart, the signal map, and the fix playbook cross-referenced to The AEO and GEO Playbook, The Local Schema Pack, and The Compliant Tracking Stack.
- **Part 6, The Weekly Audit Workflow and 8 Common Mistakes.** The 30-minute weekly routine, the reporting cadence, the eight mistakes that produce false readings.

Pairs with the rest of the AI BrandFactory marketing-infrastructure series. This is the measurement layer. The other playbooks are the fix layers.

Part 1, Why AI Visibility Needs a Formal Audit in 2026

For most of search history, visibility was a single-engine question. You ranked on Google or you did not. Bing existed but was a rounding error. Yahoo existed but was a punchline. Tracking visibility meant tracking Google rankings, and Search Console plus Ahrefs plus SEMrush gave you everything you needed.

That model broke in 2024, and by 2026 it is officially dead. Visibility is now a per-engine question across at least eight major surfaces. Each surface has its own index, its own citation algorithm, and its own bias. A practice that ranks beautifully on Google can be invisible on Perplexity. A firm that gets cited in every ChatGPT answer might not appear in Google AI Overviews at all. There is no single "rank" anymore. There is a portfolio of visibility positions.

What changed

Five forces pulled the rug out.

1, AI search adoption crossed the chasm

ChatGPT alone is now estimated to handle more than a billion search-style queries per month. Perplexity reports more than 800 million queries per month as of early 2026. Gemini and Google AI Overviews now appear on more than half of Google searches. Microsoft Copilot is integrated into Bing, Edge, Windows 11, and the entire Microsoft 365 stack. Claude.ai shipped search in late 2025.

The combined query volume routed through AI engines now rivals Google's traditional search query volume. The audience that used to ask Google is now asking eight different things, and the answer they

get is shaped by which one they asked.

2, Each engine indexes and weights sources differently

Perplexity weights recency and academic-style citations heavily. ChatGPT favors well-structured authoritative sites with strong domain reputation. Gemini blends Google's existing index with its own quality scoring. Google AI Overviews specifically privilege sites with strong schema and review profile. Brave AI uses Brave's independent index, which has different coverage than Google. Kagi's index is paid-only and intentionally curated.

The result: a site can be the top citation in Perplexity and not appear in ChatGPT for the same query. A practice can be cited by Google AI Overviews and ignored by Gemini. The "AI search results" are not one set of results; they are eight independent sets that happen to overlap.

3, Click-through rates from AI engines do not match traditional search

When Google AI Overviews appear, the click-through rate to organic results below them drops 30 to 60 percent depending on query type. The traffic that used to flow to organic positions 1 through 10 is now intercepted by the AI Overview answer block. If you are cited in the answer block, you still get some traffic. If you are not, you get almost none.

ChatGPT and Perplexity are even more compressed. Most queries get answered in-engine without the user clicking through to any source. Citation in the answer is the visibility outcome.

4, Branded traffic is increasingly mediated by AI engines

A patient who hears about your practice from a friend used to type your name into Google. Now an increasing share opens ChatGPT and asks "tell me about [practice name]". The answer they get shapes the first impression. If the AI describes your practice accurately and positively, you get a warm prospect. If the AI hallucinates something wrong or describes you negatively, you get a confused or scared one.

You cannot control what the AI says, but you can audit what it actually says today, week over week.

5, Competitors who measure are pulling ahead

Most professional services have done nothing about AI visibility yet. The few that have started are pulling ahead of the field with a 12 to 18 month head start. Once an AI engine has settled on its preferred citations for a query, dislodging the incumbent is hard. The audit is the first step toward establishing your incumbent position before someone else does.

Why a per-engine audit is necessary

You cannot improve what you cannot measure. The traditional "Google ranking" measure no longer maps to what a prospect sees. They might see ChatGPT's answer, Perplexity's answer, Gemini's answer, or Google AI Overviews' answer. Each is its own citation surface. If you measure only Google ranking, you are measuring 30 to 50 percent of the actual visibility surface. The other half is invisible to you.

The audit is the structured check across all eight major surfaces. Done weekly, it surfaces three things:

1. **Where you are visible** (the engines that already cite you, for which queries)
2. **Where you are not** (the gaps; the most actionable column of the audit)
3. **Where competitors are visible that you are not** (the competitive opportunities)

These three columns drive the visibility roadmap. Without the audit, the roadmap is guesswork.

The cost of staying invisible

Concrete cost of NOT auditing in 2026:

- **Lost first-impression queries.** Every prospect who asks an AI about your industry without seeing your name is a prospect you cannot retarget. There is no UTM, no cookie, no GA4 event. The conversation happens in the AI engine and ends there.
- **Negative or hallucinated descriptions of your business going uncorrected.** AI engines occasionally describe businesses inaccurately. Without monitoring, you do not learn until a confused prospect calls.
- **Competitor positioning.** Once an AI engine has decided that Practice A is the leader in your specialty in your city, that placement compounds. Ranking shifts on AI search are slower than on traditional search; first-mover positioning sticks.
- **Wasted optimization effort.** Spending budget on AEO and GEO without measuring per-engine outcomes means you cannot tell what works. The audit is the feedback loop.

The audit takes 30 minutes a week to run manually, or it can be largely automated with tools covered in Part 3. The cost is low. The value is the visibility roadmap and an early-warning system for visibility decay.

Part 2, The AI Engines List: 15 Surfaces and the Manual Search Patterns for Each

This is the operational core of the audit. For each engine: how to access it, the five query patterns to run, what visibility actually looks like on that engine, and the known limitations.

How to use this section

For your business, customize the five generic query patterns into your specific queries. The five patterns are:

1. **Branded query:** your business name alone
2. **Branded reputation query:** your name plus "reviews" or "is X good"
3. **Service plus location query:** "best [your service] in [your city]"
4. **Comparison query:** your name versus a known competitor, OR "alternatives to [a known leader in your space]"
5. **Informational long-tail query:** a specific question your ideal client would ask before discovering you ("[problem] solutions" or "how to [thing your service solves]")

Run all five on every engine. Note where you appear, where you do not, and where a competitor appears that you should. That is the audit row for that engine.

The 5 vertical-specific example query sets

Doctor example queries

- Branded: "Acme Family Practice Boulder"
- Branded reputation: "is Acme Family Practice Boulder good"
- Service plus location: "best family doctor in Boulder Colorado"
- Comparison: "Acme Family Practice vs Boulder Direct Primary Care"
- Informational long-tail: "how to find a primary care doctor that accepts new patients in Boulder"

Lawyer example queries

- Branded: "Smith and Associates Attorneys Phoenix"
- Branded reputation: "Smith and Associates Phoenix reviews"

- Service plus location: "best personal injury lawyer in Phoenix"
- Comparison: "Smith and Associates vs [competitor firm]"
- Informational long-tail: "what to do after a car accident in Arizona"

Home service example queries

- Branded: "Bayside Plumbing San Diego"
- Branded reputation: "Bayside Plumbing San Diego reviews"
- Service plus location: "best plumber in San Diego"
- Comparison: "Bayside Plumbing vs Coastal Plumbing San Diego"
- Informational long-tail: "what to do when water heater is leaking San Diego"

Now the engines.

Engine 1, ChatGPT (with web search)

Access: chatgpt.com (free or Plus account). Web search is enabled by default in 2026 for most queries; for guaranteed search, prefix queries with "search the web for" or use the search button in the input bar.

What visibility looks like:

- Direct citation in the answer with a numbered footnote or inline link to your URL
- Mention of your business name in the answer body (with or without link)
- Quote of your content (sometimes attributed, sometimes not)
- Appearance in a comparison list ChatGPT generates
- Inclusion in the "sources" panel that appears next to web-search answers

The 5 query checks: run each of your five queries through ChatGPT. For each, document:

- Did your business appear?
- Was the appearance a citation (with link), a mention (no link), or a quote?
- Was the description accurate?
- What other businesses appeared?

Known limitations:

- ChatGPT does not always trigger web search; for branded queries especially, it may answer from training data only
- Same query repeated minutes apart can produce different answers
- Response may differ between free, Plus, and Enterprise accounts
- The "sources" list is not always exhaustive; some sources contribute to the answer without being shown

Engine 2, Perplexity

Access: perplexity.ai (free or Pro account). All queries trigger web search by default.

What visibility looks like:

- Numbered citations next to each sentence in the answer
- A "Sources" panel above the answer listing every cited site
- Direct quotes from your content with attribution
- Appearance in the "Related" suggested follow-up questions
- Inclusion in any comparison or list-style answer

The 5 query checks: same five queries. Perplexity is the engine where citation tracking is most explicit; the Sources panel makes it easy to record which URLs were cited.

Known limitations:

- Pro account may produce different answers than free
- Perplexity weights recency heavily; old content gets cited less even if authoritative
- Some queries trigger "Pro Search" which uses a different model and citation pattern
- The default model can be changed (Sonar, GPT-4, Claude); different models produce different citation patterns

Engine 3, Gemini (gemini.google.com)

Access: gemini.google.com. Free and paid (Gemini Advanced) tiers exist; both have web search.

What visibility looks like:

- Inline citations with small superscript numbers
- A "Sources" expansion at the bottom of responses (click to expand)

- Direct mention of your business name with link to your domain
- Inclusion in lists Gemini generates ("here are the top X in Y")
- Appearance in "Double-check response" results when the user fact-checks

The 5 query checks: same five. Gemini's responses for local queries draw heavily from Google Maps and Google Business Profile data, so verify your GBP is current and complete before auditing here.

Known limitations:

- Gemini in gemini.google.com differs from Gemini integrated into Google Search (AI Overviews); audit both
- Personalization based on Google account history can skew results
- Some queries get answered without any citation if Gemini is highly confident

Engine 4, Google AI Overviews (in Google Search)

Access: google.com, perform a search query. AI Overviews appear at the top of the results page for ~50 to 60 percent of queries (varies by query type and country). For local-services and informational queries, AI Overviews appear most reliably.

What visibility looks like:

- Your URL appears in the small grid of source thumbnails to the right of or below the AI Overview answer
- Your business name is mentioned in the answer text
- Your content is paraphrased or quoted in the answer
- Your business appears in the "Places" or "Local pack" section that often follows the AI Overview

The 5 query checks: same five queries. For service-plus-location queries, also note whether your business appears in the local pack (3-pack of map results) below the AI Overview.

Known limitations:

- AI Overviews appear inconsistently; the same query at different times may or may not show one
- Personalization (logged-in Google account, location settings) heavily affects which sources are cited
- Mobile and desktop produce different AI Overview content

- Some queries trigger Google's "AI Mode" (separate full-page experience) instead of inline AI Overviews

Engine 5, Bing Copilot

Access: bing.com/chat or copilot.microsoft.com. Also accessible via the Copilot app on Windows 11, Edge, and Microsoft 365.

What visibility looks like:

- Inline numbered citations, similar to Perplexity
- A "Learn more" panel with cited sources
- Appearance in suggested follow-up prompts
- Mention or quote of your content
- Inclusion in lists or comparisons

The 5 query checks: same five. Bing Copilot's index is Microsoft's; coverage and ranking differ from Google. Sites that rank well on Bing organic search tend to be cited more in Copilot.

Known limitations:

- Three response modes (Creative, Balanced, Precise) produce different citation patterns; audit at least Balanced and Precise
- Microsoft account login affects personalization
- Copilot in Edge sidebar can pull context from the current page, which skews answers if you audit while on your own site

Engine 6, Claude.ai search

Access: claude.ai (Pro or Team account). Web search enabled in Settings; toggle on for the conversation.

What visibility looks like:

- Citations as inline links within the response
- Direct quotes with source attribution
- Appearance in lists Claude generates
- Mention of your business in the response body

The 5 query checks: same five. Claude tends to cite fewer sources per answer than Perplexity but with higher quality matching to the query intent.

Known limitations:

- Search is opt-in per conversation; if not toggled on, Claude answers from training data
- Response patterns differ between Sonnet, Opus, and Haiku model selections
- Lower visibility for purely transactional queries; Claude is more often used for analysis-style queries

Engine 7, Grok (on X)

Access: x.com/i/grok or grok.com. Free and Premium tiers; Premium gets the larger context window and faster responses.

What visibility looks like:

- Direct citation with link
- Quote of recent X posts about your business (Grok pulls heavily from X's real-time stream)
- Mention in answers, especially for time-sensitive or trending topics

The 5 query checks: same five. Grok is heavily influenced by X (Twitter) activity, so businesses with active X presence tend to be cited more often even when the query is not X-related.

Known limitations:

- Real-time bias means citations change frequently as the X conversation moves
- Less coverage for niche professional services that have limited X presence
- Different response style from other engines (more conversational, less structured)

Engine 8, SearchGPT (OpenAI's dedicated search)

Access: chatgpt.com/search or search.chatgpt.com (rolled out broadly in 2025).

What visibility looks like:

- Search-engine-style results with AI-generated summaries above
- Each result has a clickable URL with title and snippet
- AI summary at top synthesizes top sources

The 5 query checks: same five. SearchGPT positioning is a hybrid of Google search results and ChatGPT answer, so visibility shows up both in the AI summary AND in the result list below.

Known limitations:

- Newer than ChatGPT proper; result patterns still stabilizing
- Some queries route to standard ChatGPT instead of SearchGPT
- Smaller index than Google or Bing in some categories

Engine 9, You.com

Access: you.com (free or Pro account). Multiple modes: Smart Mode, Genius Mode, Research Mode.

What visibility looks like:

- Citation in AI Mode answers
- Appearance in the source list panel
- Inclusion in lists or comparisons
- Quote of content with attribution

The 5 query checks: same five. You.com is smaller than the major engines but cited reliably by some demographics (especially privacy-conscious users).

Known limitations:

- Smaller market share than top 5 engines
- Different modes produce different result patterns
- Free tier has rate limits

Engine 10, Brave Search AI

Access: search.brave.com (free, no account required). AI summaries appear for most informational queries.

What visibility looks like:

- AI summary at the top of search results
- Citation links within the summary
- Appearance in the standard search results list

- Inclusion in any "Discussions" or "Goggle" supplementary panels

The 5 query checks: same five. Brave's index is independent of Google and Bing, so coverage can differ. Sites with strong organic SEO that ignore Brave often miss visibility here.

Known limitations:

- Smaller user base than major engines
- AI summary triggers vary by query type
- No personalization (which is by design but means no logged-in audit data)

Engine 11, DuckDuckGo AI Chat

Access: duckduckgo.com/chat (free, no account). Multiple model options (GPT-4o mini, Claude Haiku, Llama, Mixtral).

What visibility looks like:

- AI chat response with inline links
- Citation of sources where applicable
- Mention of business names

The 5 query checks: same five. Different model selection produces different responses; audit at least the default model.

Known limitations:

- Privacy-first design means no personalization or session memory
- Smaller index than major engines
- Citation patterns vary heavily by which underlying model is selected

Engine 12, Phind

Access: phind.com (free or Pro). Developer-focused but used by some technical professional services for research queries.

What visibility looks like:

- Code-style citation with snippets
- Source links in a dedicated panel

- Less common for non-technical queries

The 5 query checks: same five. Most relevant for professional services in technical adjacent fields (medical IT, legal tech, technical home services like data cabling).

Known limitations:

- Heavy developer bias means non-technical queries may not surface relevant sources
- Smaller user base in non-technical professional services contexts

Engine 13, Kagi

Access: kagi.com (paid only, \$10 to \$25 per month). The Quick Answer feature provides AI summaries.

What visibility looks like:

- Quick Answer summary at the top
- Standard search results below with quality scoring
- Citation links in the summary

The 5 query checks: same five. Kagi's user base is small but high-value (premium subscribers, often senior professionals). Visibility here may not move volume but it can move high-conversion traffic.

Known limitations:

- Paid subscription required to test
- Smaller index relative to free engines
- Limited personalization

Engine 14, Mistral Le Chat

Access: chat.mistral.ai (free or Pro). EU-headquartered alternative; popular in European markets.

What visibility looks like:

- AI chat response with citations
- Source links inline

The 5 query checks: same five. Particularly important if your professional services business has any European audience (EU privacy regulations push some users toward European AI alternatives).

Known limitations:

- Smaller user base in US market
- Citation patterns less mature than Anthropic / OpenAI engines
- Multilingual responses; default may not be English

Engine 15, Arc Search

Access: Arc Browser (mobile app on iOS and Android). The "Browse for Me" feature generates an AI summary for any query.

What visibility looks like:

- AI-generated summary page that aggregates citations
- Source links within the summary
- Mention of business names in the answer

The 5 query checks: same five. Mobile-only audit, run from a phone. Best for verifying mobile-context visibility for service businesses.

Known limitations:

- Mobile-only access (cannot audit from desktop)
- Browse for Me feature triggers vary by query
- Smaller user base than major engines but skews younger and tech-aware

The audit results template

After running all 5 queries on all 15 engines, summarize per engine:

Engine	Branded	Branded Rep	Service+Loc	Comparison	Long-tail	Notes
ChatGPT	Cited	Mention	Not present	Competitor cited	Cited	Strong on branded, weak on local
Perplexity	Cited	Cited	Cited	Cited	Cited	Best engine for us
Gemini	...	...	...	...	...	...
(etc)						

This single grid is your weekly visibility report. The empty cells are your roadmap.

The rest of the AI BrandFactory library builds on patterns like this. See the full set at aibrandfactory.com.

Part 3, Automated Monitoring Tools Ranked by Use Case

The manual audit in Part 2 takes 30 minutes per week and costs nothing. That works for a single business. For an agency tracking 5 to 50 clients, or for a practice that needs daily monitoring rather than weekly, automated tools are necessary.

Eight tools cover the territory in 2026. Each occupies a different niche.

Tool 1, Profound

What it does: Tracks brand mentions across ChatGPT, Perplexity, Gemini, Claude, and Copilot. Runs your tracked queries on each engine multiple times per day and reports citation share, sentiment, and ranking against competitors.

Pricing in 2026: Starting around \$500 per month for small business plan; agency plans into 4-figures.

Use when: You want a single dashboard view of share-of-AI-voice across the major engines, with automated daily checks rather than weekly manual.

Limitations: Engine coverage is broad but not exhaustive (Brave AI, Kagi, Grok not always covered). Pricing is steep for single-practice use.

Tool 2, Otterly.AI

What it does: Specifically built for AI search visibility tracking. Monitors brand mentions in ChatGPT, Perplexity, and AI Overviews. Surfaces "missed opportunities" where competitors are mentioned and you are not.

Pricing in 2026: Around \$129 per month for the small business tier; higher for agencies.

Use when: You want a focused tool that does one thing well (AI search visibility tracking) without paying for adjacent features.

Limitations: Smaller engine list than Profound. UI is functional but less polished.

Tool 3, AthenaHQ (Athena)

What it does: AI search analytics platform tracking ChatGPT, Perplexity, Gemini, Bing Copilot, Claude. Includes share-of-voice reporting and citation source analysis (which of your URLs are cited most).

Pricing in 2026: Starting around \$200 per month.

Use when: You need source-level analysis (which specific pages on your site are getting cited) in addition to brand-level tracking.

Limitations: Newer entrant; engine coverage growing but not yet exhaustive.

Tool 4, Peec AI

What it does: Tracks how brands appear in AI search responses. Strong on Perplexity and ChatGPT. Includes prompt suggestion tools (queries you should be tracking that you may have missed).

Pricing in 2026: Around \$100 to \$300 per month depending on tracked query volume.

Use when: You want help identifying queries to track, not just running queries you have already chosen.

Limitations: Focused on the major 3 engines; less coverage of Gemini, Copilot, Brave AI.

Tool 5, Brand24

What it does: General brand-mention monitoring (long-established tool) with AI search added in 2024 to 2025. Covers traditional sources (news, social, blogs) plus AI engines.

Pricing in 2026: Around \$79 to \$249 per month depending on tracking volume.

Use when: You want one platform that tracks your brand across all surfaces (traditional + AI), not just AI search.

Limitations: AI search is an add-on, not the primary focus; depth of AI tracking is shallower than dedicated tools.

Tool 6, HubSpot AI Search Grader

What it does: Free tool from HubSpot. Enter your brand and competitors; the grader runs queries across the major AI engines and produces a comparison report.

Pricing in 2026: Free.

Use when: You want a fast one-time check or a baseline. Useful for proposals and initial client conversations.

Limitations: Free tool; not intended for ongoing tracking. Limited to a small set of queries per check.

Tool 7, Ahrefs Brand Radar

What it does: Part of the Ahrefs suite. Tracks brand mentions in AI Overviews and (via integration) other AI search surfaces. Pairs with Ahrefs' existing organic search and backlink data.

Pricing in 2026: Bundled with Ahrefs subscription; typical Ahrefs starts around \$129 per month.

Use when: You already use Ahrefs for SEO and want to add AI visibility to the same dashboard.

Limitations: Strongest on AI Overviews; less depth on ChatGPT, Perplexity, Gemini compared to dedicated tools.

Tool 8, SEMrush AI Mentions

What it does: Equivalent feature in the SEMrush suite. Tracks brand mentions in AI Overviews and ChatGPT.

Pricing in 2026: Bundled with SEMrush subscription; SEMrush starts around \$140 per month.

Use when: You already use SEMrush. Same logic as Ahrefs Brand Radar; the tool is convenient if it lives in your existing workflow.

Limitations: Less depth than Profound or AthenaHQ for pure AI tracking.

Picking the right tool

Decision tree:

- **Just one practice, want to start cheap:** HubSpot AI Search Grader (free) for baseline, then graduate to Otterly.AI (\$129/mo) for ongoing.
- **Single practice, want depth:** Profound or AthenaHQ.
- **Already pay for Ahrefs or SEMrush:** Activate the AI mentions feature there before adding another tool. The data may be enough.
- **Agency tracking 5+ clients:** Profound for breadth, with Otterly.AI as backup for cross-validation.
- **Just want the manual audit done weekly:** Skip the tools entirely. Run the Part 2 audit on Monday mornings. The 30 minutes is enough for most single practices.

What no tool replaces

The five query patterns from Part 2. Every tool above tracks queries YOU give it. The tools do not magically know which queries matter for your practice. The first investment is always the query list. Get that right with the framework in Part 2; then layer the tooling on top if scale demands it.

Part 4, The Metrics and the Dashboard

A weekly grid of "did we appear" yes/no checks is a useful start. To turn the audit into a steady-state operating dashboard, define five metrics that summarize the audit at a glance.

Metric 1, Share of AI Voice (SoAIV)

Definition: Across all engines and all queries you track, the percentage of citation slots that mention you (versus competitors or no one).

Formula:

$$\text{SoAIV} = \frac{\text{(your citations across all engine-query pairs)}}{\text{(total citation slots across all engine-query pairs)}}$$

A citation slot is one mention opportunity (one engine, one query). If you track 5 queries on 8 engines, that is 40 slots per audit run. If you appear in 12 of them, your SoAIV is 30 percent.

Why it matters: Single number that tracks visibility over time. Easy to chart week-over-week.

Target: For a single-location practice in a competitive metro, 25 to 40 percent SoAIV is strong. For a multi-location brand or a national authority, 50 percent and above.

Metric 2, Citation Count by Engine

Definition: How many of the 5 queries cited you, per engine.

Why it matters: Reveals which engines are working for you and which are not. Two practices can have the same SoAIV with very different per-engine distributions; one might be cited heavily in 2 engines and not at all in 6, the other might be moderately cited across all 8.

Target: Citation in at least 3 of the 5 queries on each top-tier engine (ChatGPT, Perplexity, Gemini, AI Overviews, Copilot).

Metric 3, Sentiment of Citations

Definition: Of the citations you receive, what proportion describe you positively, neutrally, or negatively.

Why it matters: Being cited is not always good. AI engines occasionally cite businesses in negative comparisons ("X firm has poor reviews compared to Y firm"). Sentiment tracking catches this.

Scoring (manual, fast): Read each citation. Score positive (recommended, praised, top-of-list), neutral (mentioned without judgment), or negative (cited as cautionary or unfavorable comparison).

Target: 80 percent or higher positive across citations. Below that and the audit surfaces a reputation issue (covered in Part 5 diagnosis).

Metric 4, Competitor Comparison

Definition: For each query, who else appears? Track the top 3 to 5 competitors.

Why it matters: Visibility is relative. Your 30 percent SoAIV may be best-in-class or worst-in-class depending on what competitors achieve. The gap between you and the top-cited competitor is the actionable spread.

Format: For each engine and query, list "you appeared, competitor A appeared, competitor B did not". Run this monthly; competitor changes are slower than your own.

Metric 5, Trend Over Time

Definition: Week-over-week and month-over-month change in SoAIV per engine.

Why it matters: AI engines update frequently. A practice cited in Perplexity's top results in Q1 may slide by Q2 if competitors invest more or if Perplexity changes its citation algorithm. Trend detection is the early warning.

Format: Line chart of SoAIV per engine over time. If any engine's SoAIV drops more than 20 percent week-over-week, investigate that week.

The dashboard layout

A simple weekly dashboard has these panels:

1. **Headline number:** SoAIV this week vs last week (with arrow)
2. **Per-engine bar chart:** citation count by engine, this week
3. **Citation sentiment pie:** positive / neutral / negative split
4. **Competitor comparison table:** top 3 competitors and their citation counts vs yours
5. **Trend line chart:** SoAIV per engine over the last 12 weeks
6. **Action items:** 3 to 5 bullet points generated from the audit (e.g., "Not cited in Gemini for 3 of 5 queries; investigate Gemini-specific signals")

This fits on one screen. It is the format to share with stakeholders weekly.

Cadence

- **Weekly:** Run the manual audit (or pull tool report). Update the dashboard. Review with team.
- **Monthly:** Update competitor list. Refresh the 5 query patterns if business focus has changed.
- **Quarterly:** Review which engines have grown in volume; add or drop engines from the audit. Reassess SoAIV targets based on market changes.
- **Annually:** Full reset of the audit framework if the AI engine landscape has materially shifted (which it does, every 12 to 18 months).

Part 5, Diagnosis: When You Are Not Cited, Why Not and How to Fix

The audit surfaces gaps. This part diagnoses each gap and points to the fix.

The diagnosis flowchart

For each gap (an engine + query pair where you should appear and do not), walk through these five questions in order. Stop at the first "yes."

Question 1, Does your site have a page that should match this query?

Run the query in standard Google search. Click the top 5 results. Do they match the query intent? Now look at your own site. Do you have a page that addresses the same intent, with comparable depth?

If no: this is a content gap. The fix is to publish the page. AI engines cannot cite content that does not exist.

If yes: continue to question 2.

Question 2, Does that page rank in standard Google search?

Search the query in Google. Where does your page appear in organic results?

If not in top 20: the page has an organic search visibility problem. AI engines lean heavily on traditional ranking signals as input to citation decisions. A page that ranks position 50 organic is rarely cited in AI Overviews regardless of content quality. Fix the organic ranking first; AI citation usually follows. See **The AEO and GEO Playbook** (P16) for the AEO-specific signal map.

If in top 20: continue to question 3.

Question 3, Does the page have proper schema markup?

Inspect the page (View Source or use Schema.org Validator). Check for:

- Organization or LocalBusiness schema on the homepage
- Service schema on service pages

- Article schema with author Person on content pages
- Review and AggregateRating schema where applicable
- BreadcrumbList on every non-home page

If schema is missing or broken: AI engines deprioritize sites with weak structured data. Schema is the explicit machine-readable description of what the page is about; without it, the AI has to infer. Fix the schema. See **The Local Schema Pack** (P17) for the complete schema specification.

If schema is in place and validates: continue to question 4.

Question 4, Does the business have authority signals AI engines weight?

Five signals AI engines look for as authority proxies:

1. **Review profile:** Google Business Profile rating + count, plus reviews on category-specific platforms (Healthgrades, Avvo, Yelp depending on vertical). Strong signal.
2. **Citation density:** how often is your business mentioned across the web (news, directories, blog mentions, social posts) without you having to ask? AI engines weight unsolicited mentions heavily.
3. **Backlink profile:** domain authority and quality of inbound links from authoritative sites in your industry.
4. **Content depth on the topic:** a single thin page about a topic ranks lower in AI engines than a topic cluster (5 to 10 related pages) demonstrating expertise.
5. **Person schema linkage:** named experts (doctors, attorneys, certified technicians) with Person schema, LinkedIn presence, and authoritative bios.

If any of these are weak: address the weakest first. Review profile is usually the highest-impact fix because it affects multiple engines simultaneously. Citation density is the slowest to build but the most defensible long-term.

If authority signals are strong: continue to question 5.

Question 5, Is the engine specifically routing this query elsewhere?

Some engines are biased against certain query types. Examples:

- **Perplexity** weights academic and authoritative news sources for some informational queries; commercial sites struggle to be cited for "how does X work" queries even with good content

- **Gemini** privileges Google's existing index heavily; if Google ranks you poorly, Gemini cites you poorly
- **ChatGPT** sometimes answers from training data without web search, which means recent updates are invisible
- **AI Overviews** appear inconsistently; same query at different times triggers or skips the AI Overview

If this is the issue: the gap may not be addressable. Some queries on some engines are not winnable for commercial sites. Document and move on; focus on engines and queries where you can move the needle.

The signal map

The five signals from Question 4, mapped to the playbook that addresses each:

Signal	Where it is fixed
Review profile	Build review velocity through compliant ask flows; cover Google plus vertical platforms
Citation density	Earned media, PR, content distribution, partnerships
Backlink profile	Standard SEO (digital PR, guest posts, partnerships, content link earning)
Content depth (topic clusters)	The AEO and GEO Playbook (P16) plus content strategy
Schema (Organization, Person, Service, Review)	The Local Schema Pack (P17)

Three diagnostic shortcuts

These three quick checks identify the most common root causes without going through the full flowchart.

Shortcut 1, "Cited everywhere except Google AI Overviews"

If you appear in ChatGPT, Perplexity, Claude, and Copilot, but never in Google AI Overviews, the issue is almost always Google-specific:

- Weak Google Business Profile (incomplete, low review count, inconsistent NAP)
- Weak organic Google ranking despite good content (technical SEO issues)
- Missing or invalid LocalBusiness schema specifically as Google reads it

Fix Google Business Profile and Local Schema. AI Overviews citations usually follow within 4 to 8 weeks.

Shortcut 2, "Cited in some engines but description is wrong or outdated"

The engine knows about you but describes you incorrectly. Three causes:

- **Outdated training data:** the engine learned about you from old content; new accurate content has not been indexed yet. Solution: publish current canonical content (a clearly dated "About" or "Services" page) and wait. ChatGPT and Claude refresh slowly; Perplexity refreshes faster.
- **Inconsistent information across the web:** your site says X, an old directory says Y, a competitor says Z. The engine averages. Solution: audit and unify your business listings (NAP consistency across Google Business Profile, Yelp, Bing Places, BBB, vertical directories).
- **Negative signals dominate:** a few negative reviews or critical articles outweigh positive signals. Solution: review velocity (build positive signal density) plus reputation response (covered in adjacent playbooks).

Shortcut 3, "Cited for branded queries but never for service-plus-location"

The engines know you exist but do not associate you with the service or location category. Three causes:

- **Weak service-specific landing pages:** your homepage talks about you but no dedicated page targets "[service] in [city]"
- **Weak local-pack signals:** Google Business Profile not set up correctly, missing service categories, missing service-area definitions
- **Generic content:** pages that could be from any practice in any city; nothing local-specific

Fix the service pages and the GBP categories. Cited-for-service rate usually doubles within 6 to 8 weeks of doing this right.

When to accept the gap

Some gaps are not worth fixing.

- **Engines with low traffic relevance:** if your customer base does not use Brave or DuckDuckGo, gaps there are not financially meaningful
- **Queries with low commercial intent:** appearing for "what is [service]" matters less than for "best [service] in [city]"
- **Queries where the cited competitor is genuinely better positioned:** if a competitor has 20 years of authority and you opened last quarter, expecting parity in 6 months is unrealistic

The audit produces a long gap list. Prioritize ruthlessly. Fix the gaps where the fix is achievable AND the engine plus query has commercial value.

Part 6, The Weekly Audit Workflow and 8 Common Mistakes

This is the operational layer. How to make the audit a routine, not a project.

The 30-minute weekly audit (manual)

Same time every week. Same person every week (or rotated with documentation). Same engines, same queries.

Setup (one-time, takes 1 hour)

1. Document your 5 query patterns in a shared doc. Be specific with the actual query strings, not generic templates.
2. Document your top 3 to 5 competitors by name (the ones you want to track against).
3. Set up a fresh browser profile or use Incognito for every audit run, to minimize personalization noise.
4. Create a spreadsheet with the audit grid: rows for engines, columns for queries, plus columns for sentiment and competitor notes.

Weekly run (takes 30 minutes)

1. Open the spreadsheet from last week. Save a copy with the new date.
2. Open Incognito. Sign out of all accounts that personalize results.

3. Run query 1 on engine 1. Record cited / mentioned / not present. Note any sentiment issue. Note which competitors appear.
4. Repeat for queries 2 through 5 on engine 1.
5. Repeat for engines 2 through 8 (the top tier). Skip engines 9 through 15 unless your weekly time budget allows.
6. Calculate this week's metrics: SoAIV, citation count by engine, sentiment split.
7. Compare to last week. Note any drops greater than 20 percent on any engine.
8. Write 3 to 5 action items based on the gaps.

This is 30 minutes once you have done it twice. The first run takes 60 minutes; you settle into rhythm by run 3.

The reporting cadence

- **Weekly internal:** 1-page report to the marketing team. Includes the dashboard panels from Part 4.
- **Monthly stakeholder:** The weekly report rolled up, plus the action items completed and their visibility impact.
- **Quarterly board / executive:** Trend chart of SoAIV over the quarter, the top 3 wins, the top 3 still-broken gaps, the budget ask for fixes.

The hardest part of this routine is doing it consistently. The audit is most valuable as a trend line, which requires running it the same way at the same cadence. Skipping a week is fine; skipping six weeks breaks the trend.

When to automate

Move from manual to automated when one of three conditions hits:

1. The audit takes more than 60 minutes per week consistently
2. You are tracking 3 or more brands (agency case, or multi-location brand with separate visibility tracking per location)
3. Stakeholders are asking for daily or near-real-time visibility (regulated industries, post-crisis monitoring)

For 1 and 2, a tool from Part 3 pays for itself quickly. For 3, custom monitoring (Profound or AthenaHQ at the higher tier) is the answer.

The 8 audit mistakes

Mistake 1, Auditing while logged in

Personalization heavily skews results. A logged-in Google account, a logged-in ChatGPT account, an Edge browser with Microsoft account sync all return personalized answers that are not what a fresh prospect sees. Always audit in Incognito or a clean browser profile.

Mistake 2, Running queries inconsistently

If you run "best plumber San Diego" one week and "plumbers in San Diego" the next, you cannot compare. Lock the exact query strings and never paraphrase mid-stream.

Mistake 3, Ignoring engine variants

Gemini in gemini.google.com is not the same as Gemini in Google Search AI Overviews. ChatGPT default mode is not the same as SearchGPT. Audit each surface separately and label clearly.

Mistake 4, Counting any mention as a win

A negative mention is not a win. A neutral mention is barely a win. Sentiment matters. The audit should distinguish positive citations from negative or neutral mentions, otherwise the SoAIV number flatters you while the underlying signal degrades.

Mistake 5, Not tracking competitors

Visibility is relative. Knowing your SoAIV without knowing competitor SoAIV is half the picture. Always track the same metrics for the top 3 competitors, even if it doubles the audit time.

Mistake 6, Auditing too few queries

Five queries gives you a baseline. Two queries gives you noise. If your service breadth is wide, audit at least 5 queries per major service line, not 5 total. A multi-specialty medical practice might audit 15 queries (5 per specialty, 3 specialties).

Mistake 7, Auditing too few engines

Skipping any of the top 5 (ChatGPT, Perplexity, Gemini, AI Overviews, Copilot) creates a blind spot. If time is short, drop tier-3 engines (Phind, Mistral, Arc) before dropping any tier-1.

Mistake 8, No action loop

The audit produces gaps. The gaps demand action. If the action loop does not exist (or stalls), the audit becomes a reporting exercise instead of a visibility-improvement program. The Friday 30-minute audit must connect to a Monday 30-minute action review where the gaps from last week become this week's work.

The pre-audit checklist

Before considering the audit framework deployed:

- ☐ 5 query patterns documented and specific (not generic templates)
- ☐ Top 3 to 5 competitors documented
- ☐ Audit spreadsheet template created with all engine and query rows
- ☐ Incognito browser workflow documented (so anyone can run the audit, not just one person)
- ☐ Weekly recurring 30-minute calendar block on the team's calendar
- ☐ Monthly recurring 60-minute action-review meeting on the team's calendar
- ☐ Stakeholder reporting template defined (1-page weekly format)
- ☐ If tooling: tool selected, queries loaded, dashboard configured
- ☐ Baseline audit completed (the first week's data, to compare future weeks against)
- ☐ Action loop process defined (who owns gap fixes, how priority is decided)

The compound effect

Done weekly for a year, the AI Visibility Audit produces:

- A trend line showing visibility direction over 52 data points
- A gap inventory that drives the visibility roadmap
- An early-warning system that surfaces visibility decay within 1 week of onset
- A stakeholder narrative that justifies investment in AEO, schema, content, and reputation work
- A competitive intelligence stream (which competitors are gaining or losing visibility on which queries)

The audit by itself does not improve visibility. It surfaces the data that drives the work that improves visibility. Treated as a discipline, it becomes the central operating measurement for the marketing function in 2026.

There is more where this came from. For deeper playbooks on AI search visibility, structured data, conversion tracking, and content distribution, visit www.aibrandfactory.com.

The AI Visibility Audit is the measurement layer. The fix layers are **The AEO and GEO Playbook**, **The Local Schema Pack**, and **The Compliant Tracking Stack for Professionals**, all available at the same library.

Built by AI BrandFactory

This playbook is part of AI BrandFactory's open toolkit. Production-grade systems, frameworks, and templates for founders, agency operators, marketers, and developers shipping AI-powered businesses in 2026.

What else is in the library

50+ playbooks across content, copy, ads, SaaS, agency operations, AI systems, and more. All free, all production-grade, all packaged for solo operators.

Browse the full set: files.aibrandfactory.com/playbooks

What we built it for

We ship lead magnets that beginners can use and pros can adapt. Every playbook in the library follows the same standards: no AI fluff, no consultancy speak, real source attribution, working code or templates where applicable, and beginner-friendly explanations layered with operator-grade depth.

How to use this

- Read it through once
- Pick the 1 to 2 things that apply to your project right now
- Ship them this week
- Come back for the next layer when you are ready

License

Free to use with attribution. Adapt freely. Cite back to AI BrandFactory when you publish, train, or remix.

Stay in touch

aibrandfactory.com for new playbooks every week.

github.com/AI-BrandFactory for the open-source repos.

Questions, feedback, or you have shipped something cool with this?

piyush@winmassiveimpact.com.

Built with care by the AI BrandFactory team.