

How AI assistants actually recommend eye doctors, and what the evidence says you should do about it

A review of the peer-reviewed research, the vendor claims that do not survive it, and an original study of what ChatGPT cites when someone asks it for the best optometrist in town.

Streben Marketing · July 2026 · For optometry practice owners

00 THE SHORT VERSION

Patients are starting to ask ChatGPT and Google's AI which eye doctor to see. So the obvious question for a practice is: how do you become the one it recommends? The marketing industry has a confident answer, built around your website. After reading the research and running our own test, we think that answer is mostly wrong.

For the query a patient actually uses to pick a doctor, **"best optometrist in [city]," the practice's own website is a minority of what AI cites, and often none of it.** The levers that move an AI recommendation sit almost entirely off your website: your Google Business Profile, your review volume and rating, and your presence on the directories the AI already trusts. Your website earns citations on a narrower set of questions (specific conditions and services), and it has to be technically reachable by the AI crawlers, which we found is not a given.

Most of what is sold as "AI optimization," FAQ schema, llms.txt files, and "make your pages more specific," is either untested, inert, or traceable to studies that do not say what the seller claims. Some of the most-quoted statistics in the field are fabricated. This paper shows the evidence, names what holds up, and gives a practice a short, honest list of what to actually do.

One caveat before the evidence, so none of this is read as "SEO is dead." This paper is about the AI *recommendation* layer, and you will see it largely ignores the on-page work, schema, content structure, keyword targeting, that classic search still rewards. That does not make those things wasted. They are how you get found in ordinary Google and Bing results, which remain a large share of how patients reach a practice, and Google and Bing weigh exactly the signals the AI layer shrugs off. The right posture is additive, not either-or: keep your standard SEO best practices for organic search, and add the off-site work described here for the AI layer. Two systems choose differently, and you want to show up in both.

01 WHY THIS IS WORTH AN OPTOMETRIST'S ATTENTION

Adoption is real and fast, even if you would not know it from your own analytics. In a survey of nearly 1,000 US patients, **47% said they had used AI to find a new provider**, up from 31% nine months earlier.¹ A separate representative survey found **45% of consumers had used AI for local business recommendations, up from 6% a year earlier.**²

Two things temper it. First, people verify: 88% fact-check the AI's recommendation, and **42% say they always go to native review platforms to check** before booking.² Second, the click-through is low, so the traffic barely shows up in your website reports even when the AI is influencing the choice. The influence happens upstream of your site, which is the theme of this entire paper.

02 HOW AN AI ASSISTANT ACTUALLY PICKS A LOCAL PRACTICE

The mechanism is more legible than the hype suggests, and it is different from classic SEO. Two facts drive everything below.

First, for genuinely local queries, Google usually does not serve an "AI Overview" at all. It serves the local pack. In a 540-query study across six industries including optometry, local-intent queries returned an AI Overview only **15% of the time, while the local 3-pack appeared 93% of the time.**⁷ A 500-million-keyword analysis found AI Overviews on local keywords actually *declining*, to about 0.01%.⁸ So for "eye doctor in [city]" typed into Google, the surface you are competing on is still Google Maps and the Business Profile, not a generative answer (asking ChatGPT or Perplexity directly is the separate case Section 3 measures).

Second, when AI does answer, it increasingly cites its own place data, not your web page. Across 1.3 million Google AI Mode citations, **google.com was the single most-cited domain at 17.4%, more than the next six domains combined**, and those links resolve to Google Business Profiles and knowledge panels.⁴ An independent analysis of 32 million observations measured an **8.4x surge** in Google citing itself this way, concentrated in local verticals including healthcare, and summarized it bluntly: "users are no longer being sent to your homepage. They are being shown a Google-hosted card containing your data."⁵

The other assistants each lean heavily on a structured data source, though (as our own study shows for ChatGPT) they also run open web searches:

Where each assistant gets its local answer. Sources at right.

ASSISTANT	PRIMARY STRUCTURED SOURCE	EVIDENCE
Google AI Mode / Overviews	Google Business Profile / Maps	#4, #5
Gemini	Google Maps (dominant)	#6
ChatGPT	Place data + a licensed Yelp feed	#6, #9
Perplexity	Yelp data partnership	#10

Gemini is the clearest case: one analysis of 350,000 locations found it grounded **solely in Google Maps, with 100% data accuracy versus about 68% for ChatGPT and Perplexity.**⁶ And as of July 23, 2026, ChatGPT reads a licensed Yelp feed of 330 million reviews and 8 million listings directly.⁹ None of these inputs is your website's marketing copy.

03 **OUR STUDY: WHAT CHATGPT ACTUALLY CITES, BY HOW SPECIFIC THE QUERY IS**

The measurement that would settle the "does my website get cited" question does not exist publicly. Every directly relevant study we could find has an important limitation (Section 6). So we ran our own.

METHOD

We queried ChatGPT (via the DataForSEO scraper, natural web-search behavior, US) across **9 US metros** and **3 query types**, running each query **7 times** to measure run-to-run variance, for 189 planned responses (187 returned usable data; 2 were lost to transient API timeouts). For each answer we separated the *cited web sources*, the pages ChatGPT grounded on, from the *named businesses*, the practices it listed from place data, and classified every source (first-party practice site, directory, forum, Google entity, medical authority). The three query types sample three points on a broad-to-specialized spectrum, not the whole patient journey: generic ("best optometrist in [city]"), condition ("dry eye specialist in [city]"), and sub-specialty ("scleral lens specialist in [city]"). Full metro list, coding rubric, and per-run data are held in the study files and available on request.

RESULT: THE SPECIFICITY GRADIENT

The pattern is monotonic and near-unanimous across all nine metros. Whether ChatGPT cites a practice's *own website* as a grounding source tracks closely with how specific the query is. On the generic query, **zero of 63 runs cited any practice website**; the answer grounded on Reddit and directories, even though ChatGPT still named five to nine local practices per answer from place data. On the condition query, first-party sites appeared in **89% of runs**; on the sub-specialty query, **100%**.

Table 1. Cited-source distribution by query type ($n = 187$ runs across 9 metros). The first four columns are the share of cited web sources (rows omit a small news remainder, $\leq 1\%$); the last column is the share of runs that cited any first-party site.

QUERY TYPE	FIRST-PARTY SITE	DIRECTORY	FORUM	AUTHORITY	RUNS CITING ANY FIRST-PARTY SITE
"best optometrist in [city]"	0.0%	28.8%	70.2%	0.0%	0% (0/63)
"dry eye specialist in [city]"	67.1%	11.4%	7.6%	13.3%	89% (56/63)
"scleral lens specialist in [city]"	79.5%	1.4%	4.1%	15.1%	100% (61/61)

Read the last column first. It is the most useful number in this paper: **for the generic "best optometrist" query, editing your website is unlikely to move the needle**, because the practice's own site was not what any of those answers were visibly built from. That answer is assembled from place records, which name you, plus Reddit and directory pages, which rank you; community forums alone were 70% of the generic query's cited sources. You cannot manufacture Reddit threads, but you can influence the reputation they reflect and the place records, reviews, and directory profiles the model leans on, none of which live on your marketing pages. (This measures what ChatGPT visibly cited, not the full internal path; a website can still shape an answer indirectly by feeding those records.)

The specialty queries invert this. Ask for a "dry eye specialist" or a "scleral lens specialist" and the grounding sources become the practices' own service pages, 67% then 80% of all citations. This is the query class where a practice's own pages get cited; specific service pages are what earn that, though our data cannot isolate which page qualities (structure, freshness, authority, exact-match terms) drive it. The practical instruction falls straight out of the table: **compete for your reviews and profile on the generic term; compete with your website on the specific ones.**

Three honest limits. First, this is one model (ChatGPT), one vertical (optometry), one week; the DataForSEO scraper targets the search-enabled surface, so it does not estimate how often ChatGPT answers with no web search at all (only 2 of 187 runs did here, whereas the independent St. Gallen study in Section 10 found ChatGPT ran no web search on 57.8% of its runs). That gap sharpens the finding rather than undercutting it: much of the time no practice website is cited because no search happens at all, and when a search does happen, the specificity gradient above decides whose site gets used. Second, run-to-run variance is real (the generic query averaged 2.6 different top picks across 7 runs versus 1.9 for sub-specialty); a follow-up study varied the wording four ways per tier and the gradient held at both ends (see below); and 0 of 63 is a strong signal but not a proven absolute zero, its 95% upper bound is roughly 6%, on runs clustered within 9 metros rather than fully independent. Third, the one metro that broke pattern is instructive. In Portland, where large academic systems (OHSU, Legacy) dominate, the *condition* query was captured by those institutions and independent practices fell to 29%, though independents still took 100% of the *sub-specialty* query. In academic-heavy markets, a small practice wins narrower, not broader.

OUT-OF-SAMPLE SPOT CHECK (LIVE)

A study you cannot rerun is just an assertion, so we re-ran the two ends of the gradient live on two metros that were *not* in the original nine: Minneapolis, and Raleigh (the second named on the spot by a skeptical reader, to remove any suspicion of cherry-picking). Both reproduced it exactly. The generic "best optometrist in [city]" query cited **zero practice websites** in both, grounding on Reddit and directories, in Raleigh's case four Reddit threads and nothing else, while ChatGPT still named a dozen-plus local practices from place data. The "scleral lens specialist in [city]" query cited **practice websites** in both (Raleigh cleanly: fusioneyecare.com, raleighcontactlens.com, truevisioneye.com; Minneapolis with one arguable network domain among the three). This is one run per query, subject to the run-to-run variance noted above: a directional spot check, not a replication estimate, and too few runs to bound reliability. The exact prompts, surface, and coding rules are in the study files, so it can be re-run.

PHRASING ROBUSTNESS, AND THE "BEST" EFFECT (TWO FOLLOW-UPS, 305 RUNS)

Two follow-ups tested whether the gradient is a property of the query or of the exact wording. First, four phrasings per tier across three metros (179 runs): it held at both ends, every generic wording near zero first-party (pooled **2%**), every sub-specialty wording high (pooled **98%**). Second, a controlled test (126 runs) that toggled *only* the word "best" against a fixed base isolated a real second lever: **adding "best" costs roughly one specificity tier of first-party citation.**

Table 2. First-party citation rate, descriptive base vs. adding "best" (toggle only), pooled over 3 metros, 21 runs per cell.

QUERY LEVEL	DESCRIPTIVE BASE	+ "BEST"	CHANGE
"optometrist in [city]"	5%	0%	-5
"dry eye specialist in [city]"	90%	0%	-90
"scleral lens specialist in [city]"	95%	86%	-10

Why did two specialties diverge so sharply, dry eye falling to zero while scleral barely moved? The cause is not the word "best" itself, but what "best" can find. A superlative is a request for a *ranking*, and the model answers a ranking request from ranking and opinion sources when they exist, and only from the practice's own pages when they do not. Those ranking sources exist in abundance for a common, self-diagnosed condition like dry eye (Reddit threads, "best dry eye doctor" listicles, directory rankings) and almost not at all for a rare, referral-driven fitting like scleral lenses, which patients reach through their doctor rather than by searching. The cited-source mix makes it visible: under "best," dry-eye answers drew **70% of their sources from forums and none from practice sites**, while scleral answers drew **3% from forums and 64% from practice sites**. Same superlative, opposite corpus. "Best" jumps to the reputation content where it exists (and the practice's pages drop out), and falls back to those pages where it does not. That is also why the generic tier never moves: even its plain form is already directory-grounded, so there is nothing left for "best" to displace. The condition collapse was unanimous across all three metros, and in it the same doctor usually stayed the top *named* result even as the *cited source* flipped to Reddit and Healthgrades.

The lesson: the more a patient's query asks for a ranking ("best..."), the more the answer runs on reviews and reputation instead of your website, and that is exactly the highest-intent search. Your pages win the descriptive phrasings and the genuinely sub-specialised ones; the "best [common service]" query is a reviews-and-profile game.

04 REVIEWS ARE THE GATE, AND THE NUMBERS ARE SPECIFIC

Two independent parties, using different samples and methods, measured the same thing: AI assistants filter on rating before they choose, and the numbers line up almost exactly.

4.3★

Average rating of practices
ChatGPT recommends⁶

4.1★

Average rating Perplexity
recommends⁶

3.9★

Average rating Gemini
recommends⁶

The 350,000-location study that produced those numbers put it plainly: reviews function **"as a filter rather than a ranking signal."** The AI gates on rating, then chooses from what is left.⁶ A restaurant study by a different company reports the same 4.3 / 4.1 / 3.9 as effective floors ("at or above") on a separate sample.¹¹ To be exact, the tile numbers are averages of the practices each assistant recommends, not a measured hard cutoff; but because those averages sit high and recommendations rarely fall below them, rating behaves as a level to clear. And it is *volume*, not just rating, that separates winners: in a statistically corrected study of 1,120 query executions, businesses that appeared in Google's AI Mode had **49% more reviews at the median** than those that appeared only in the traditional local pack. For plumbers the gap was 337% (546 reviews versus 125).¹²

This is why review generation is not a "nice to have" for AI visibility. On the current evidence it is one of the strongest observable correlates of inclusion, and it is one your website does not control. (Correlate, not proven cause: review count also tracks business age, prominence, and market share, so it is a strong signal to act on, not a guaranteed lever.)

05 HOW NARROW THE AI FUNNEL REALLY IS

Before a practice spends heavily to "win AI," it should see the size of the prize. The same 350,000-location study measured how often each platform recommends any given business:

Table 3. Share of locations recommended, by surface.⁶

SURFACE	LOCATIONS RECOMMENDED
Google local 3-pack	35.9%
Gemini	11%
Perplexity	7.4%
ChatGPT	1.2%

The study's own framing is that AI visibility is **3 to 30 times harder** than ranking in traditional local search.⁶ (Table 3 comes from one vendor's dataset, SOCi's, though it is the largest of its kind, 350,000 locations and 3.2 million AI queries, and we confirmed the headline figures against their primary report; read the exact percentages as directional and the pattern as solid.) A separate 322-market analysis found AI local packs surface only about a third as many distinct businesses as the traditional 3-pack, and show no call button.¹³ The honest promise to a practice is not "we will make you the AI answer." It is "we will make you the practice the AI keeps finding, and get you through the review gate," which is a reasonable, measurable objective rather than a guaranteed outcome.

06 WHAT THE SCIENCE SAYS ABOUT "OPTIMIZING" YOUR CONTENT

There is a genuine academic literature on generative engine optimization, and it is far more skeptical than the agency version. The single most important structural fact: **every study that reported a positive optimization effect on a real commercial engine did so by skipping the retrieval step**, handing the engine the document and then measuring what it did with it. Those studies answer a narrower question, how content affects use once a page is already retrieved, not whether your page gets retrieved at all. The one large study that put retrieval back in, across 171,003 documents, found that rewriting the body of a page to "optimize" it **reduced** its presence in results by roughly 9%, and its final citation rate by 6%.¹⁴

The careful replications agree. A benchmark of 1.9k queries with proper significance testing found only **3 of 54 optimization tactics significantly helped, and none helped for question-answering**. Moving a document to the top of the retrieval set beat every content tactic.¹⁵ An MIT and Columbia study of 13,747 queries found that **10 of the 15 hand-written optimization heuristics "yield negligible or even negative changes in ranking."**¹⁶ Keyword stuffing was null to negative in the tests that isolated it, and an "authoritative tone" was among the weakest and least stable levers.¹⁷

THE PAPER EVERYONE CITES DOES NOT SAY WHAT THEY CLAIM

The foundational "GEO" paper (Princeton and collaborators, KDD 2024) is real and honest, but it was a *simulated* pipeline scored by an older model, and its famous "+40%" is a source's share of a five-source answer it was *already in*, not a 40% increase in getting found.¹⁷ It tested nine content tactics. **None was FAQ or Q&A formatting.** The widely repeated claim that "Princeton research found FAQ content increases AI citation 37 to 40%" is fabricated: those numbers belong to a different tactic in the paper, reattached to something it never tested, and circulated with the correct link to the real paper attached.¹⁸

07 WHAT TO IGNORE, AND WHY

Three tactics are sold hard to local practices and do not hold up:

FAQ schema markup. The only controlled test of structured data on AI citation, a difference-in-differences study of 1,885 pages that added schema against 4,000 controls, found **no meaningful uplift on any platform** (and slightly negative on Google's AI Overviews).¹⁹ A 216,524-page analysis found FAQ-schema pages were cited slightly *less* than pages without it.²⁰ A clean demonstration showed why: a researcher put a fake address inside deliberately *invalid* JSON-LD, and ChatGPT and Perplexity both repeated it, because the model reads the markup as "slightly weirdly punctuated prose," not as structured data.²¹ Writing good, plain answers to real patient questions still helps. The *markup* around them does close to nothing for AI citation specifically; it may still serve conventional search and entity clarity.

llms.txt files. No major AI provider has committed to reading them. A study of 137,210 domains with server-log data found **97% of llms.txt files received zero traffic**, and the largest source of the requests that did arrive was SEO audit tools, not AI.²² Google states directly that you do not need any such file to appear in its results.²³

Suspiciously precise "ranking factors." The circulating "optimal passage length is 134 to 167 words" and "multi-modal content sees 156% higher selection" trace to vendor blogs reporting their own unpublished analyses, with no dataset, sample, or method released to check.²⁴ The most-quoted traffic claim in the field, "AI traffic is 4.4x more valuable," rests on an unspecified "conversion rate" with no disclosed sample size, data window, or method, and was measured on the vendor's own product category.²⁵ A field pattern worth internalizing: the warning sign is a precise-sounding number with no disclosed sample, method, or date, not precision itself. Real studies can be precise; fabrications hide their provenance.

08 THE FAILURE MODE NOBODY CHECKS: CAN THE AI EVEN REACH YOUR SITE?

An AI search assistant generally cannot cite the current version of a page its crawler is blocked from fetching, and this fails silently. Modern security edges and some managed hosts block the AI search crawlers by default or by

misconfiguration, and it shows up nowhere in Google Search Console, your WordPress admin, or your robots.txt file.

A REAL, ANONYMIZED EXAMPLE

Auditing a portfolio of eye-care and healthcare sites, we found one whose security edge returned an "allowed" response to a browser and to Google, and to a nonsense fake crawler, while returning a hard block to **ChatGPT's search crawler, Claude's, and Perplexity's**. That practice is invisible to every AI search assistant, and its own robots.txt file claimed the opposite. The block was a targeted rule at the content-delivery edge, not a plugin the owner could see. It is a quick fix once found, and high-impact until it is.

The distinction that matters: blocking a *training* crawler (GPTBot, Google-Extended) costs little in AI search visibility, because the search crawlers are separate. Blocking a *search* crawler (OAI-SearchBot, Claude-SearchBot, PerplexityBot) can materially reduce whether your current page is retrieved and directly cited, though cached or third-party copies may still surface.²⁶ Every practice should have this checked at onboarding and after any host change. It is the cheapest high-stakes item on the list.

09 WHAT A SINGLE-LOCATION PRACTICE SHOULD ACTUALLY DO

Stripped to what survives the evidence, the list is short, unglamorous, and mostly not about your website:

1. **Get and keep reviews.** Rating clears the bar (recommended practices average 4.3 / 4.1 / 3.9), and review volume is the strongest visible correlate of inclusion. This is the single highest-leverage action, and it is off your website.
2. **Complete and maintain your Google Business Profile, Yelp, Healthgrades, and Zocdoc listings.** These are the structured records AI reads and the directories it cites. An unclaimed or stale Yelp listing is now a direct AI liability, because ChatGPT reads a licensed Yelp feed.
3. **Verify the AI search crawlers can reach your site.** Once, at onboarding, and after any host or CDN change.
4. **Write genuine, plain-language answers to the specific conditions and services you treat,** because those are the queries where your website does earn citations. Keep it factual and specific. Skip schema added solely to chase AI citations.
5. **Keep pages fresh.** Recently updated content is cited about 1.7 times as often as stale content (6.0 vs 3.6 citations per page in a 216,000-page analysis)²⁰, so a real publishing and refresh cadence matters more than a one-time rewrite.
6. **Measure honestly.** A single AI query is noise: the top result changes run to run. Track appearance rates across many runs over weeks, not a single screenshot.

Everything on that list is measurable, and none of it requires believing a fabricated statistic.

10 A NOTE ON MEASUREMENT, BECAUSE IT IS EASY TO FOOL YOURSELF

AI answers are unstable in a way that breaks most "AI visibility" reporting. A 45-day study across four engines found that two runs of the *same* prompt minutes apart shared only about a quarter to a third of their cited sources, and that **a single run is "essentially uninformative," with a 95% confidence interval of plus or minus 0.72 on a zero-to-one scale.** Reliable ongoing monitoring needs at least 7 runs per prompt per day.²⁷ Worse, ChatGPT did not even perform a web search on 57.8% of runs. And because the same question can be phrased many ways, a fixed prompt list mostly measures which phrasing you happened to pick.²⁸ Any vendor selling you a single-number "AI visibility score" from one weekly query is selling you noise. Our own study meets the seven-run bar for a point-in-time snapshot; it is not a multi-day stability study, and it varies reruns but not phrasings, so it does not capture how much a reworded prompt would move the result.

REFERENCES

Each entry notes what the source actually measured. Every external entry has been fetched and confirmed to resolve and support its claim; the academic sources were checked against full text.

1. rater8, *2026 Patient Choice Report*. SurveyMonkey, Apr 2026, ~1,000 US patients. [MediaPost coverage](#) · [rater8.com](#) 47% had used AI to research a provider, up from 31% in rater8's prior wave nine months earlier (stated in MediaPost's coverage, attributed to the report). rater8's own summary page leads with a separate 36% "AI influenced the choice" figure, so it is cited here for the 47%/31% usage stat specifically.
2. BrightLocal, *Local Consumer Review Survey: AI Trust*, 2026, n=1,002. [brightlocal.com](#) 45% used AI for local recs; 42% always verify on review platforms.
4. SE Ranking, *Google links in AI Mode*, Feb 2026, 68,313 keywords, 1.32M citations. [seranking.com](#) google.com = 17.4% of all AI Mode citations.
5. Profound, *google.com is now AI Mode's #2 cited domain*, 32M+ observations, 2026. [tryprofound.com](#) 8.4x surge in Google self-citation via Business Profile cards.
6. SOCi, *2026 Local Visibility Index*, 350,000+ locations / 2,751 brands / 3.2M AI queries, Jan 28 2026. [soci.ai](#) · [press release](#) · per-platform detail [via Search Engine Land](#) Recommendation rates 35.9/11/7.4/1.2%, avg recommended ratings 4.3/4.1/3.9, reviews as a confidence threshold, ~30x selectivity. Primary source confirms 35.9%/1.2%/30x/4.3/sample; SEL carries the Gemini/Perplexity rates and 100-vs-68% accuracy.
7. Whitespark, *Prevalence of AI Overviews in local search*, 540 queries, 6 industries. [whitespark.ca](#) Local-intent AIO 15%, local pack 93%.
8. seoClarity, *AI Overviews impact*, 500M+ keywords. [seoclarity.net](#) AIO on local keywords ~0.01%.
9. OpenAI + Yelp licensing deal, announced Jul 23, 2026. [via Search Engine Land](#) ChatGPT reads a licensed Yelp feed; ~330M reviews / 8M listings per Yelp's announcement.
10. Perplexity + Yelp Fusion data partnership (since Mar 12, 2024). [maginative.com](#) Perplexity local answers use licensed Yelp Fusion data.
11. Uberall, *The Discovery Gap* (restaurants), May 2026. [uberall.com](#) Reports the same 4.3/4.1/3.9 rating floors (restaurants, separate sample).
12. Local SEO Data, 1,120 query executions, May 2026, bootstrap CIs + corrections. [prweb.com](#) 49% median review-count gap (plumbers 337%).
13. Sterling Sky / Places Scout, *State of Local SEO 2026*, 322 markets. [sterlingsky.ca](#) AI local packs surface ~32% as many businesses; no call buttons.
14. Kim et al., *SAGEO Arena*, 2026, arXiv 2602.12187, 171,003 docs, retrieval reinstated. [arxiv.org](#) Body optimization reduced presence ~9%, citation ~6%.
15. Puerto et al., *C-SEO Bench*, NeurIPS 2025, arXiv 2506.11097. [arxiv.org](#) 3 of 54 tactics helped; none in QA.
16. Bagga et al. (MIT/Columbia), *E-GEO*, arXiv 2511.20867, 13,747 queries. [arxiv.org](#) 10 of 15 hand-written heuristics yield negligible or negative ranking changes.

17. Aggarwal et al., *GEO: Generative Engine Optimization*, KDD 2024, arXiv 2311.09735. [arxiv.org](https://arxiv.org/abs/2311.09735) Simulated pipeline; "+40%" is share-of-answer; no FAQ condition.
 18. Fabrication trace: "Princeton FAQ 37-40%." The paper (#17) tested nine tactics, none FAQ; figures belong to Statistics Addition. Circulated with the correct arXiv link attached.
 19. Ahrefs, *Does schema affect AI citations*, May 2026, difference-in-differences, 1,885 pages vs 4,000 controls. ahrefs.com No major uplift on any platform.
 20. SE Ranking, *How to optimize for ChatGPT*, 216,524 pages, Nov 2025. seranking.com FAQ-schema pages cited slightly less (3.6 vs 4.2); fresh content cited ~1.7x (6.0 vs 3.6).
 21. Williams-Cook, *Schema and LLMs*, Search Engine Journal, Jun 2026. searchenginejournal.com Fake address in invalid JSON-LD repeated by ChatGPT and Perplexity.
 22. Ahrefs, *llms.txt study*, Jun 2026, 137,210 domains. ahrefs.com 97% of llms.txt files got zero traffic.
 23. Google, AI optimization guide. developers.google.com No special files/markup required.
 24. "134-167 words" and "156% multi-modal": vendor-blog figures (e.g. wellows.com) from unpublished in-house analyses; no dataset, sample, or method released to verify.
 25. Semrush, "AI traffic is 4.4x more valuable." semrush.com Value = unspecified "conversion rate"; no sample size, date, or method; own category.
 26. OpenAI ([bots docs](#)), Anthropic, Perplexity crawler documentation. Search crawlers vs training crawlers; blocking search crawlers removes you from cited AI answers.
 27. Schulte et al. (Univ. St. Gallen), *Don't Measure Once*, Apr 2026, arXiv 2604.07585. [arxiv.org](https://arxiv.org/abs/2604.07585) Single run "essentially uninformative"; need 7+ runs; ChatGPT skipped search 57.8% of runs.
 28. Jack et al. (unusual.ai), arXiv 2605.27440. [arxiv.org](https://arxiv.org/abs/2605.27440) Paraphrase variance exceeds rerun variance.
-

About this research. The original study (Section 3) is complete (n = 187 runs across 9 US metros; figures in Table 1), and every external reference has been fetched and confirmed to resolve and support its claim, with the six academic sources checked against full text. The Section 8 crawler example is anonymized. Prepared by Streben Marketing, 2026.