

THE 2026 STATE OF AI CITATIONS · FIRST EDITION

We asked three AIs who to buy. They couldn't agree.

ChatGPT, Perplexity, and Gemini now answer the questions your customers used to Google. We put the same 39 buyer questions to all three — 117 answers — and measured exactly who gets recommended, and why.

Prompt asked to all three engines: **"Best project management software for agencies"**

● Perplexity

- 1 Productive
- 2 Teamwork.com
- 3 Ravetree
- 4 **ClickUp** SHARED
- 5 **Wrike** SHARED
- 6 **Asana** SHARED
- 7 Monday.com
- 8 Jira

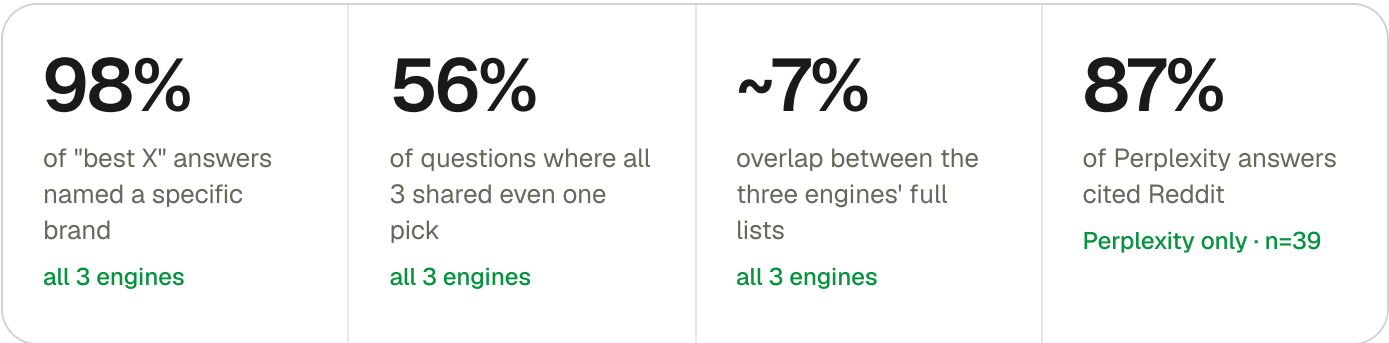
● ChatGPT

- 1 **Asana** SHARED
- 2 Trello
- 3 Monday.com
- 4 Basecamp
- 5 **Wrike** SHARED
- 6 **ClickUp** SHARED
- 7 Notion
- 8 Smartsheet
- 9 TeamGantt
- 10 Zenkit

● Gemini

- 1 Monday.com
- 2 **ClickUp** SHARED
- 3 **Asana** SHARED
- 4 Teamwork.com
- 5 **Wrike** SHARED
- 6 Adobe Workfront

Three names shared by all three engines, out of fifteen named. *Everything else — different on every engine.*

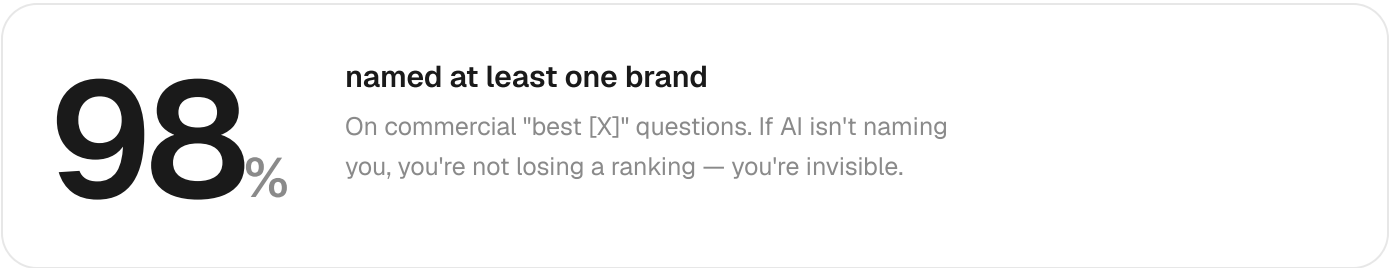


What each number covers. The first three use all three engines. Everything source-level in this report — Reddit, domain concentration, and the whole page-anatomy section — is **Perplexity only, n=39**, because ChatGPT and Gemini don't expose their citations at all. Nobody can tell you what those two cite, including us. Every claim below says which scope it's on.

FINDING 01

AI is a recommendation engine now

Ask any of the three engines a commercial question and it names specific companies almost every time. The era of "ten blue links to choose from" is over — buyers get a short list of names. The only question that matters is whether one of them is yours.

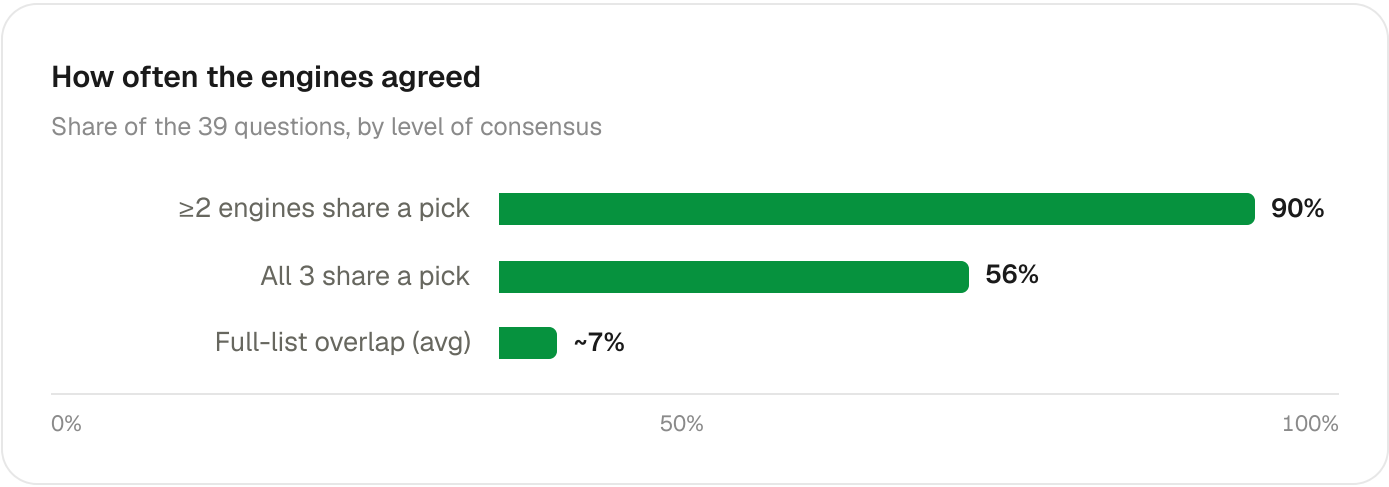


FINDING 02 · THE BIG ONE

The engines don't agree

Run the same question through all three and they largely recommend **different** companies. All three shared at least one common pick in only 56% of questions. At least two agreed 90% of the time — but their full lists overlapped just ~7%.

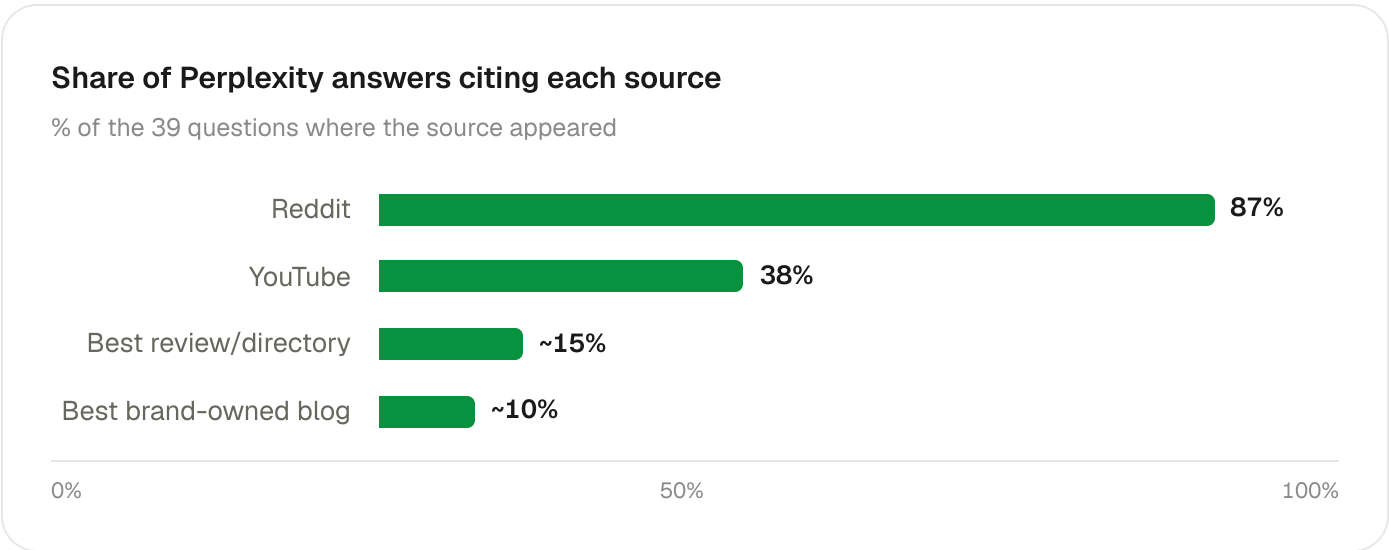
"Rank on AI" is really "rank on each AI." You can be ChatGPT's #2 and invisible on Gemini.



FINDING 03

Perplexity cites Reddit, not your blog

Across every category, Perplexity's most-cited source was Reddit — on 87% of questions — followed by YouTube (38%). Brand-owned blogs and even professional review sites trailed far behind. For at least one major engine, getting cited is less about your blog and more about showing up in the conversations and videos AI already trusts.



Source-level data is Perplexity-only — ChatGPT and Gemini don't expose their citations, so we can't yet say whether they lean the same way. Treat this as Perplexity-specific.

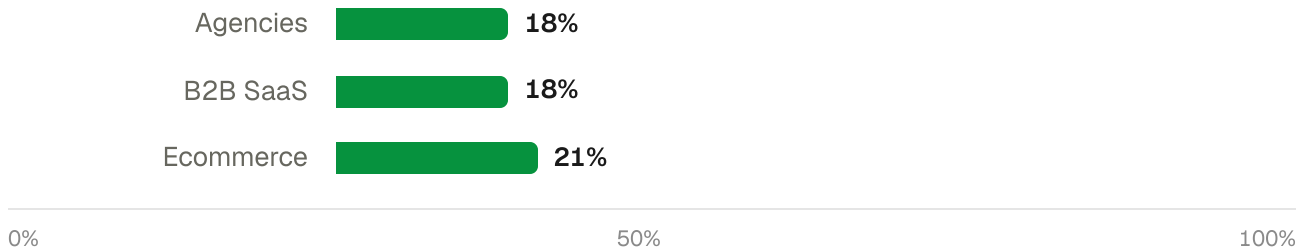
FINDING 04

It's a long tail, not winner-take-most

No handful of sites owns AI's citations. The top three domains earned only about a fifth of citations in each category — which means a newcomer can get cited without dethroning an incumbent.

Citation share held by the top 3 domains

By category (Perplexity)



FINDING 05

The cited page: long, listicle, rarely original

The pages AI cites share a shape — long "best-of" listicles, median ~3,200 words. But only about a third contain original data. That last number is the opening: two in three cited pages have **no** original research. It's the emptiest lane in GEO.

What the top cited page had

Share of the 39 top-cited pages (Perplexity)





0%

50%

100%

Median cited-page length: **~3,200 words** · AI's cross-category favorites: Toptal, Finsweet, HubSpot, ActiveCampaign, Flow Ninja.

WHAT TO DO ABOUT IT

The playbook

01

Win each engine separately

Track your citation share on ChatGPT, Perplexity, and Gemini as three different scoreboards — because they are.

02

Show up where the AI looks

For Perplexity, that means Reddit threads and YouTube, not just your blog.

03

Publish what gets cited

Long, genuinely useful comparison and "best-of" pages with clear structure — not thin marketing copy.

04

Be the 31%

Put original data — your own tests, benchmarks, results — into your content. Two-thirds of cited pages don't, so it's how you get quoted.

TRANSPARENCY

Four bugs we caught before publishing

A report is only as good as its worst number. Before publishing we stress-tested our own analysis and fixed four things. We're listing them because a study about *trustworthy content* that hides its own corrections wouldn't deserve to be one.

What we found and fixed

FIXED

Our page analysis truncated long pages, undercounting "original data" pages at 5%. The real figure is 31%.

FIXED A set-math error reported cross-engine agreement as 0%. The corrected figure was 62%.

FIXED

Brand extraction initially grabbed client names and technologies as "recommendations." We tightened it.

FIXED

A final pass recomputed every headline figure straight from the released CSV — and the agreement numbers were *still* wrong. All-three-agree is 56% (not 62%), at-least-two is 90% (not 92%), full-list overlap is ~7% (not ~8%). Those recomputed figures are the ones in this report. We'd already fixed that number once; recomputing from raw data is the only reason we caught it twice.

THE DATA

Every number, in one place

METRIC	SCOPE	VALUE
Answers named a specific brand	3 engines	98%
All 3 engines shared ≥1 pick	3 engines	56%

METRIC	SCOPE	VALUE
≥2 engines shared ≥1 pick	3 engines	90%
Full-list overlap (avg Jaccard)	3 engines	~7%
Answers citing Reddit	Perplexity	87%
Answers citing YouTube	Perplexity	~38%
Top-3 domain citation share	Perplexity	18–21%
Cited pages that are listicles	Perplexity	77%
Cited pages with original data	Perplexity	31%
Median cited-page length	Perplexity	~3,200 words
Mean cited-page length (for contrast)	Perplexity	~4,300 words
Questions × engines	total	39 × 3 = 117

Two notes for anyone recomputing. Medians, not means. The typical cited page is ~3,200 words; the mean is ~4,300, dragged up by a few very long outliers. We report the median because it describes the typical page — if you compute the mean and get 4,300, that's why, not a discrepancy. **The headline survives name-folding.** Engines name the same product differently — Perplexity says "Monday Work Management" where the others say "Monday.com" — and naive string matching would manufacture disagreement. We re-ran everything with 8 such variant groups folded: all-three-agree moved **+0.0 points** (56.4% → 56.4%), overlap +0.2. The disagreement finding is not an artifact of naming.

CHANGELOG

What changed, and when

This report is versioned. When a number changes, it gets listed here rather than quietly edited.

v1.1 — 2026-07-15

CORRECTED

Four figures had gone stale between an internal recompute and this document: all-three-agree 62% → 56%, at-least-two 92% → 90%, full-list overlap ~8% → ~7%, Perplexity/Reddit "9 in 10" → 87%. The dataset never changed; the report had simply not been updated to match it.

CORRECTED

The worked example showed Wrike as a Gemini-only pick. All three engines named it — the set shared by all three is Asana, ClickUp

and

Wrike. The side-by-side also showed an idealised ordering and now reflects what each engine actually returned.

CORRECTED

The headline read "117 buyer questions." It is 39 questions × 3 engines = 117 *answers*

.

ADDED

Scope labels on every headline figure — three of the five are Perplexity-only and the summary didn't say so.

ADDED

The name-folding robustness test, the median-vs-mean note, and

`verify-numbers.py`

, which recomputes all 48 figures in this series from the CSVs and fails if any drifts.

v1.0 — 2026-07-02

First edition.

QUESTIONS THIS REPORT ANSWERS

Quick answers, straight from the data

How often do ChatGPT, Perplexity and Gemini agree on a recommendation?

On 39 identical buyer questions in July 2026, all three engines shared at least one pick 56% of the time, at least two agreed 90% of the time, and their full lists overlapped just ~7%.

Does AI actually name brands in its answers?

Yes. 98% of “best X” answers named at least one specific brand — buyers get a short list of names, not links.

What sources does Perplexity cite most?

Reddit, in 87% of answers, followed by YouTube at 38% (Perplexity only, n=39 — ChatGPT and Gemini don’t expose their citations).

What kind of page gets cited by AI?

Long “best-of” listicles — median ~3,200 words, 77% listicles — and only 31% carried original data. Original research is the emptiest lane.

Does AI recommend you?

Run your brand through the free AI-visibility audit and see exactly what ChatGPT, Perplexity, and Gemini say about you — and whether they name you or your competitor.

[Run my free audit →](#)

[Download the full dataset](#)

Harper*Flow*

Content engineered to get quoted across every engine — and fact-grounded so AI represents you accurately.

[Download the dataset](#)[All six reports →](#)

Method: 39 buyer-intent questions across agencies, B2B SaaS & ecommerce, each run through Perplexity (sonar), ChatGPT (gpt-4o-mini) & Gemini (2.5-flash), July 2026. Directional, not a census. Source-level findings are Perplexity-only.

[All GEO research](#) · [Start free](#) · [Contact](#) · [Privacy](#)