

The Answer Index

What four AI answer engines actually cite across eight industries, measured on a fixed panel of 480 questions and published in full.

480

QUESTIONS

4

ANSWER ENGINES

8

INDUSTRIES

16,069

CITATIONS RECORDED

Six things the data shows, each with its number attached

We asked ChatGPT, Gemini, Perplexity and Google AI Overviews the same 480 questions and recorded every source each one cited. This page is the report in one sitting. Every figure carries its denominator on the page where it is established.

01 No engine resembles another engine as much as it resembles itself
The strongest overlap between any two engines is 0.181. Perplexity re-asked the same questions 38 minutes later agreed with itself at 0.900. These are four supply chains, not four views of one ranking.

02 ChatGPT mostly does not search, and what it searches for is predictable
It cited a source in 15.0% of 480 answers. Sixty-nine of those 72 sit in two question types, who provides it and what it costs. Six question types produced zero citations in every industry.

03 Appearing in most answers is not the same as taking most of the links
Reddit appears in 68.1% of Perplexity answers and holds 4.45% of its citations. Both are true. Quoting either alone misleads.

04 Google AI Overviews cite video at a rate no other engine approaches
YouTube appears in 47.0% of rendered AI Overview panels, against 8.1% of Gemini answers and 2.3% of Perplexity answers.

05 The citation surface is overwhelmingly tail
7,775 distinct domains produced 16,069 citations. 67% of those domains were cited exactly once in the entire corpus.

06 Who holds the answer varies enormously by industry
Sites owned by the companies that sell the thing take 64.5% of citations in healthcare and dental and 27.0% in education and training. The spread is the finding.

What we measured

Four engines, 480 questions, one capture on 2026-09-01. Every question was asked of every engine, and every cited source was recorded.

The questions live in a fixed panel published in full alongside this report. Eight industries carry sixty questions each, split evenly across ten question types: who provides it, how two options compare, what it costs, how to choose, how it works, what is wrong, when to act, which credentials matter, what a term means, and what else could be done instead.

Because the question mix is identical in all eight industries, a difference between two industries is a property of the industry rather than of the questions asked. That control is what makes the industry comparisons in this report mean anything.

Of 1,920 probes, 1,919 returned an answer; one failed on a vendor server error. 1,902 are citation-eligible, the 17 excluded being Google AI Overview probes where no panel rendered. 1,458 of the eligible cells carried at least one citation.

A source is counted once per answer, never once per link, so every figure answers how often a source gets into the answer rather than rewarding engines that emit more links.

Engine profiles

Citation rate is the share of citation-eligible cells carrying at least one source. Google AI Overviews is measured over 462 rendered panels; the other three over 480 answers each.

ENGINE	ANSWERS	CITED ANYTHING	CITATIONS	SOURCES PER ANSWER	DISTINCT SOURCES
Perplexity	480	100.0%	7,350	15.3	4,788
Gemini	480	99.2%	5,042	10.6	3,579
Google AI Overviews	462	93.1%	3,164	7.4	1,977
ChatGPT	480	15.0%	513	7.1	368

SCOPE

This measures answer-engine APIs, not consumer apps. ChatGPT here is gpt-4.1 with web search through a data vendor, not the product people use. Google AI Overview panels come from a SERP API whose panel-presence reporting we have seen err in both directions, so every AI Overview figure here is directional.

01 No engine resembles another engine as much as it resembles itself

The strongest agreement between any two engines is a mean per-question overlap of 0.181. Perplexity, re-asked the same questions 38 minutes later, agreed with itself at 0.900.

An engine is roughly five times closer to its own answer half an hour later than to any other engine at the same instant. These are not four views of one ranking. They are four different supply chains, and a brand competing for a place in one is largely not competing in the others. Any single blended score across engines describes a surface no one uses.

Cross-engine agreement			
Mean per-question Jaccard overlap of cited-domain sets, on questions both engines answered. Self-agreement measured across two capture rounds 38 minutes apart, for the two engines re-run.			
PAIR	OVERLAP		QUESTIONS
Perplexity and Google AI Overviews	0.181		462
Perplexity and Gemini	0.151		480
Gemini and Google AI Overviews	0.151		462
Perplexity and ChatGPT	0.020		480
Gemini and ChatGPT	0.013		476
Google AI Overviews and ChatGPT	0.011		430
Perplexity against itself, 38 minutes later	0.900		240
Google AI Overviews against itself, 38 minutes later	0.622		201

Only Perplexity and Google AI Overviews were re-run. ChatGPT and Gemini have no measured stability, and no claim is made about theirs. A 38 minute gap establishes that the instrument reproduces; it says nothing about drift over days or weeks.

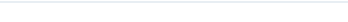
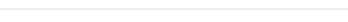
02 ChatGPT mostly does not search, and what it searches for is predictable

ChatGPT returned a citation in 15.0% of its 480 answers. Sixty-nine of those 72 answers sit in just two of the ten question types.

It searched when asked who provides something and what it costs. It answered from memory when asked how to choose, how something works, what is wrong, when to act, which credentials matter, and what a term means. Six question types produced zero citations across all eight industries. When it did search, its answers ran 6,194 characters on average against 2,303 when it did not.

ChatGPT citation rate by question type

Share of answers carrying at least one cited source. Denominator is 48 answers per question type, six in each of eight industries.

QUESTION TYPE	CITED	ANSWERS
Provider discovery	72.9% 	35 / 48
Cost	70.8% 	34 / 48
Comparison	4.2% 	2 / 48
Alternatives	2.1% 	1 / 48
Selection criteria	0.0% 	0 / 48
Explainer	0.0% 	0 / 48
Problem diagnosis	0.0% 	0 / 48
Timing	0.0% 	0 / 48
Trust and credentials	0.0% 	0 / 48
Definition	0.0% 	0 / 48

WHAT THIS MEANS FOR MEASUREMENT

Any ChatGPT citation share must state its denominator. Computed over all 480 answers it is one number; over the 72 that cited anything it is a different one. Both are legitimate. Conflating them is not, and the gap between them is the finding.

03 Appearing in most answers is not the same as taking most of the links

Reddit appears in 68.1% of Perplexity answers and takes 4.45% of Perplexity citations. Both numbers are true, and reporting either alone misleads.

Perplexity cites about 15.3 sources per answer, so a domain can be present almost everywhere while holding a small share of the total. Presence answers whether you are in the room. Mass answers how much of the room you hold. Most reporting in this category quotes the first and implies the second.

Presence against mass

Presence is the share of that engine's answers containing the domain. Mass is the domain's share of every citation the engine made. Denominators are the engine's citation-eligible cells and its total citations.

ENGINE	REDDIT PRESENCE	REDDIT MASS	YOUTUBE PRESENCE	YOUTUBE MASS	CITATIONS
Perplexity	68.1%	4.45%	2.3%	0.15%	7,350
Gemini	11.5%	1.09%	8.1%	0.77%	5,042
Google AI Overviews	38.3%	5.59%	47.0%	6.86%	3,164
ChatGPT	2.9%	2.73%	0.0%	0.00%	513

ChatGPT figures are computed over all 480 of its answers for presence and over its 513 citations for mass. Its citation base is small because it searched rarely. Reddit is nonetheless the only domain to appear in the top eight sources of all four engines.

WHY THE DISTINCTION MATTERS

A source that is in the room in two thirds of answers while holding under five percent of the links is a different competitive fact from one that holds a third of the links. The first is a fixture. The second is an incumbent. Reddit, on this evidence, is a fixture.

04 Google AI Overviews cite video at a rate no other engine approaches

YouTube appears in 47.0% of rendered AI Overview panels, against 8.1% of Gemini answers and 2.3% of Perplexity answers.

Nearly half of the AI Overview panels we captured attached a video result. A brand with no video has no mechanism for entering that half. This is the one finding in the report that depends on no classification judgement, because YouTube is matched by an exact published rule. It is also the finding most exposed to the vendor caveat, since it rests on what the SERP API reported inside the panel.

YouTube presence by engine			
Share of each engine's citation-eligible answers containing youtube.com. Google AI Overviews over 462 rendered panels.			
ENGINE	PRESENCE		ANSWERS
Perplexity	2.3%		11 / 480
Gemini	8.1%		39 / 480
Google AI Overviews	47.0%		217 / 462
ChatGPT	0.0%		0 / 480

05 The citation surface is overwhelmingly tail

7,775 distinct domains produced 16,069 citations. 5,201 of those domains, 67% of them, were cited exactly once in the entire corpus.

The average domain is cited 2.1 times across 1,458 answers. This is not a short list of authoritative sites with a tail attached; it is almost entirely tail. Engines differ sharply in how wide they cast: Perplexity drew on 4,788 distinct domains at about 15.3 per answer, Google AI Overviews on 1,977 at about 7.4 per panel. The practical consequence is that a given site's realistic expectation of being cited on any single probe is low, and the useful question is whether the kind of page it publishes is the kind the engine reaches for.

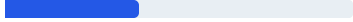
06 Who holds the answer varies enormously by industry

Sites owned by the companies that actually sell the thing take 64.5% of citations in healthcare and dental and 27.0% in education and training.

Pooled across all eight industries and four engines, provider-owned sites take 54.0% of citations. The spread matters more than the pooled figure. The composition of an AI answer is a property of the market, not a constant, which means no row in this table can be assumed to apply to an industry that is not in it.

Provider-owned share of citations by industry

Provider-owned means the site of a company selling the thing being asked about, whether a national brand or an independent operator. Pooled across all four engines. Denominator is every citation recorded in that industry.

INDUSTRY	PROVIDER-OWNED		CITATIONS
Healthcare and dental	64.5%		2,035
B2B SaaS	64.4%		2,060
Professional services	64.3%		2,113
Home services	60.0%		1,982
Automotive	57.6%		1,952
Retail and consumer	49.8%		1,904
Consumer finance	43.7%		1,955
Education and training	27.0%		2,068

READ WITH CARE

This finding depends on classifying 7,775 domains, which was done by a local model. A blind audit of 60 domains against independent classification found a 28% disagreement rate, 32% when weighted to the citation mass of the corpus. The eleven-class breakdown is therefore withheld. The provider versus non-provider split survives because only 7 of the 60 audited domains crossed that boundary, and six of the seven moved toward provider, so the figures above are a floor. A broader definition that also counts directories and agency comparison content as commercial reaches 71.4%. We report the narrower one, because a directory is not the company you are asking about.

The most cited sources, by engine

Presence, not mass. Each figure is the share of that engine's answers in which the domain appeared at least once. At the top eight, the four lists share only Reddit. At the top fifteen, they share three domains.

Perplexity N = 480 answers		Gemini N = 480 answers		Google AI Overviews N = 462 rendered panels		ChatGPT N = 72 answers that cited anything	
reddit.com	68%	reddit.com	12%	youtube.com	47%	reddit.com	19%
forbes.com	13%	youtube.com	8%	reddit.com	38%	legalclarity...	14%
linkedin.com	13%	forbes.com	6%	quora.com	10%	angi.com	11%
nerdwallet.c...	10%	facebook.com	6%	facebook.com	10%	threebestrat...	11%
en.wikipedia...	6%	wikipedia.org	5%	nerdwallet.c...	6%	forbes.com	11%
investopedia...	6%	experian.com	4%	linkedin.com	5%	expertise.com	8%
consumerrepo...	6%	quora.com	3%	experian.com	4%	nerdwallet.c...	8%
pmc.ncbi.nlm...	6%	medium.com	3%	investopedia...	4%	moneygeek.com	7%
cnbc.com	5%	research.com	3%	instagram.com	3%	zocdoc.com	7%
finance.yaho...	4%	bankrate.com	3%	consumerrepo...	3%	bankrate.com	7%
medium.com	4%	geico.com	3%	bankrate.com	3%	techradar.com	6%
angi.com	4%	nih.gov	3%	progressive...	2%	bestprosinto...	6%

How many of each engine's top fifteen sources appear in another engine's top fifteen

Read across a row. Shared by all four engines: bankrate.com, forbes.com, reddit.com.

	PERPLEXITY	GEMINI	GOOGLE AI OVERVIEWS	CHATGPT
Perplexity	15	5	9	6
Gemini	5	15	8	3
Google AI Overviews	9	8	15	4
ChatGPT	6	3	4	15

ChatGPT's list is computed over the 72 answers that carried any citation, which is why its shares run higher on a much smaller base. The full table for every domain and every engine, with counts and denominators, is in the open data package.

SECTION THREE

What it implies

Everything before this page is measurement. Everything on it is interpretation, and the two are kept apart deliberately so they never blur. Every number here is established in section two.

There is no single AI search visibility score, and one should not be built

Four engines that agree with each other between 0.011 and 0.181, and with themselves at 0.622 and 0.900, cannot be averaged into a ranking that describes any of them. A blended visibility number across engines is an artifact of its own pooling. A brand asking to be cited is asking four separate questions, and only the per-engine answer is stable enough to act on.

Contract on conditions, never on the citation

Perplexity reproduced 90% of its own citations 38 minutes later. Google AI Overviews reproduced 62%. A single probe is one observation under one set of conditions, and the absence of a citation on one probe is not a diagnosis. Anyone who commits to producing a citation is committing to something they do not control. What can be committed to is producing the reading and improving the conditions the brand does control.

Video is a budget decision, not a content decision

If nearly half of rendered AI Overview panels attach a video, a brand with no video has no mechanism for entering that half. That is a staffing and production conclusion rather than a copywriting one, and it is the finding least vulnerable to how sources were classified.

Do not read the Reddit number as an instruction to post on Reddit

Reddit is present in two thirds of Perplexity answers and holds under five percent of its citations. That tells you Reddit is in the room. It says nothing about whether a post placed there would be selected, whether selection would persist, or whether it would move anything downstream. We measured citation, not causation, and we did not run the intervention. Seeding forums is also against platform policy and carries reputational risk that no measurement here offsets.

Treat the industry spread as a test to run, not a table to copy

Provider-owned share runs from 27% to 65% across eight industries. Eight industries are eight instances, not a sample of all industries. The useful conclusion is that answer composition is market-specific, which means the only defensible move is to run the reading on your own market rather than assume any row here transfers.

The test before you act on any AI search reading, including this one: name the engine, name the question type, name the denominator, and say what the number would look like if the same question were asked again half an hour later.

A reading that cannot answer all four is not yet evidence, and most of what currently circulates about AI search visibility cannot answer any of them.

Limits

Stated here so that no reader has to discover them. A finding is only as strong as the weakest thing it rests on, and these are the weakest things.

- This measures APIs, not consumer products. ChatGPT here is gpt-4.1 with web search through a data vendor. The consumer ChatGPT app may behave differently, and nothing here should be read as describing it.
- The Google AI Overview payload comes from a SERP API whose panel-presence reporting we have previously seen wrong in both directions. Every AI Overview figure is directional and pending a browser-screenshot control run.
- Stability was measured for two engines, not four, on half the panel, across a 38 minute gap. That establishes the instrument reproduces. It does not establish that rankings hold over days or weeks, and it says nothing about ChatGPT or Gemini.
- Source classification carries a measured 28% disagreement rate by domain and 32% weighted to citation mass. The eleven-class breakdown is withheld for that reason. Only the provider versus non-provider split is reported, as a floor.
- Gemini returns citations as bare hostnames with no page path, so it contributes no page-level information and is excluded from any page-level analysis.
- A citation is not a click. This study makes no claim about referral traffic to any cited domain. A citation is also not a recommendation. A brand can be named in an answer without being linked, and linked without being recommended. We measured links.
- Eight industries, United States, English, one capture. Nothing here generalises beyond that.

HOW THE FINDINGS WERE TESTED

Two rounds of analysis produced 56 candidate findings. Each was handed to an independent reviewer instructed to refute it rather than confirm it, by recomputing the number from the raw data and hunting for a confound or a denominator mismatch. Forty were refuted and do not appear in this report. The six that do are the ones that survived.

Method, data and citation

The panel is 480 questions: eight industries, sixty each, split evenly across ten question types. The identical mix in every industry is the study's control. Questions are written in the voice a person would use with an assistant, not as search keywords, because a keyword panel measures what an engine does with a keyword and gets reported as what an assistant does with a question. No branded queries are included, and local intent is named in the question text.

A domain is counted once per answer, never once per link. Every table states its denominator, and the analysis code asserts at runtime that no table mixes two and that every printed share recomputes from its own numerator and denominator.

The panel, every cited domain per engine with counts and denominators, and the full methodology are published as an open data package so the study can be re-run and checked by anyone.

Panel	v2.0, 480 questions, sha 7a15060d8b5ec5f6, published verbatim
Capture	2026-09-01, single round, four engines
Engines	gpt-4.1 web search, gemini-2.5-flash, perplexity sonar, Google SERP AI Overview
Locale	United States, English, country level
Metros in panel	Phoenix, Columbus OH, Sacramento, Kansas City
Probes	1,920 planned, 1,919 measured, 1 vendor error
Citations	16,069 across 7,775 distinct domains
Independence	No client account, market or query appears in this corpus

HOW TO CITE THIS REPORT

Hendricks, Brandon Lincoln (2026). The Answer Index: what AI answer engines cite across eight industries. Hendricks. hendricks.ai/research/the-answer-index

Data: Hendricks, Brandon Lincoln (2026). The Answer Index [Dataset]. Zenodo. doi.org/10.5281/zenodo.22242103

When quoting a figure, carry its denominator. Every number in this report has one, and a figure reported without it cannot be checked. Corrections are published at hendricks.ai/corrections with a date and the superseded value, never edited silently.

Hendricks is a Search Intelligence Engineering firm. It maps the questions that drive a market, measures whether a brand enters the consideration set across Google and AI search, and engineers the conditions that shape visibility and trust. Absence of a citation is not a diagnosis, and no one can commit to producing one.