

MEASUREMENT

Share of Model: The AI-Visibility Metric That Replaces Keyword Rankings

By Scott Tischler · August 27, 2026 · 9 min read

Ranking #1 stopped telling you the truth

For two decades, the scoreboard of digital marketing was the keyword ranking. You picked a phrase, you climbed toward position one, and the report to the boss was a row of green arrows. I ran and reviewed those reports for more than twenty years. The uncomfortable thing I have to tell you now is that, for a fast-growing share of the questions your customers ask, the ranking has quietly stopped measuring the thing you actually care about.

When someone opens ChatGPT and asks which vendor to shortlist, or asks Perplexity for the best way to solve a problem, or types a question into Google and reads the AI Overview at the top, there is frequently no list of ten links to rank on. There is one synthesized answer. Either the model names you inside that answer, or it doesn't. Your position on a page of blue links is irrelevant to a person who never sees the page. So the question every marketing team should be asking is no longer *where do I rank?* It's **how often does the machine say my name?**

That question needs a metric. The one I use, with clients and internally at Alrecommend.ai, is Share of Model.

What Share of Model actually measures

Share of Model is the percentage of AI-generated answers — across a fixed set of buyer questions, run through every major answer engine — in which your brand is named inside the answer itself.

It is deliberately built as the direct descendant of share of voice, the metric brand advertisers have used for decades to measure how much of a category's attention they own. Share of voice asked: of all the advertising noise in this category, how much is mine? Share of Model asks the 2026 version: of all the answers the machines give to questions in my category, how many include me? The unit of measurement moved from the billboard and the search result to the sentence the AI generates.

Two words in that definition are carrying weight, and they are where most homemade measurement goes wrong.

"Named inside the answer itself." Being buried in a list of eleven footnote citations is not the same as being named in the prose the user reads. Models frequently retrieve dozens of sources and mention two brands. The mention is the recommendation; the citation is plumbing. Share of Model counts the mention, and treats a link the user has to expand a source tray to find as, at best, a weaker signal.

"Across every major answer engine." Your visibility is not one number, because the engines do not agree. A brand that ChatGPT loves can be invisible in Gemini and middling in Perplexity, because each blends training data and live retrieval differently. A single blended figure is useful for a trend line, but the moment you want to act, you need the per-engine breakdown.

How to measure it without fooling yourself

The method is not complicated, but it is unforgiving of shortcuts. Here is the process I trust.

1. Build a fixed prompt set

Write 30 to 50 questions a real prospect would ask, and then freeze them. They should span three altitudes: problem-level ("how do I get my business recommended by AI"), category-level ("best AI visibility firms"), and brand-level ("is [your brand] any good"). The freezing matters more than the writing. If you change the questions every month, you are measuring your question-writing, not your visibility. A stable set is the only way the number means anything over time.

2. Use a fresh session for every prompt

Run each question in a clean session with no memory of the previous ones. Answer engines are context machines; if you ask about your own brand and then ask a category question in the same thread, the earlier turn contaminates the later answer and flatters your score. One prompt, one fresh session, every time.

3. Run all the engines that matter to your buyer

At minimum: ChatGPT, Gemini, Perplexity, and Google's AI Overviews, plus Claude and Grok where your audience uses them. Don't measure only the engine you personally prefer. Measure where your customers actually ask.

4. Score with a strict, written rule

For each answer, record three things: is the brand named in the answer body (yes or no), at what position relative to competitors, and is the framing accurate. Share of Model is the simple ratio —

brand-named answers divided by total answers — but the position and accuracy columns are where the real story lives. Being named last, hedged, and slightly wrong is a different problem from not being named at all, and it needs a different fix.

The output is a grid: prompts down the side, engines across the top, a clean yes/no in every cell. That grid is your baseline and, re-run on a cadence, your scoreboard.

Read three cuts, not one average

A single "we're at 22%" number is where analysis goes to die. The average almost always hides the decision. Read the data three ways.

- **Overall share** is your headline trend line — useful for the boardroom, useless for the work.
- **Share by engine** tells you where the gap is. If you're at 40% in Perplexity and 5% in Gemini, you don't have a content problem, you have a Gemini problem, and the fix is specific.
- **Share by prompt cluster** tells you which parts of your category you own and which you've ceded. You might dominate the brand-level questions and be invisible on the category-level ones — which means the model knows who you are but doesn't yet believe you're a leading answer to the underlying problem.

Those three cuts turn a vanity number into a work list.

How you actually move the number

Here is the good news and the discipline at once: you don't move Share of Model by gaming it. You move it by becoming genuinely easier for a model to cite. Three levers do most of the work, and I've written about each at length elsewhere on this site.

Retrievability. If a model can't fetch or parse your best content, none of the rest matters. Server-rendered HTML, clean structure, and AI crawlers like GPTBot, ClaudeBot, and PerplexityBot explicitly allowed — this is the unglamorous floor beneath everything.

Extractability. Models lift passages they can quote cleanly. Lead with the answer, define your terms, make every important claim true and complete on its own. If a model quoted just two sentences from your page, would they be accurate without the surrounding paragraph? If not, tighten them.

Off-site consensus. This is the lever that separates durable visibility from a lucky month. Models build their picture of you from the whole web — reviews, publications, forums, comparison pages — not from your homepage. When independent sources describe you consistently, the model repeats the story as

fact. When your story is thin or contradictory, the model hedges or omits you. You cannot fully control the answer from your own domain, which is exactly why this work is hard to fake and worth the most.

My honest, practitioner's read — offered as informed opinion, not a guarantee — is that Share of Model will behave like a compounding asset. The brands teaching the models a clear, corroborated story now are building a position later entrants will have to overwrite, and overwriting an established consensus is far harder than establishing one. That is the whole case for starting to measure today, even if the first number stings.

What to do this week

You do not need a platform to begin. Take your twenty most important buyer questions, run them through ChatGPT, Gemini, Perplexity, and Google's AI Overviews in fresh sessions, and fill in the grid by hand. You now have a Share of Model baseline that most of your competitors do not have and, in most cases, have never thought to calculate. Then pick the single lowest cell — the engine or the cluster where you're most invisible — and fix the one thing underneath it. Re-run next month. Watch the line. That loop, run patiently, is how you stop optimizing for a ranking nobody reads and start owning the answer everybody does.

Key takeaways

- ♦ Keyword rank measures position on a page of links; Share of Model measures how often AI answers actually name you — a different, and now more important, thing.
- ♦ Share of Model is the percentage of answers, across a fixed prompt set and every major engine, in which your brand appears in the answer body itself.
- ♦ Measure it with a stable prompt set, fresh sessions per prompt, all engines, and a strict scoring rule: named in the answer body counts, a buried citation does not.
- ♦ Track three cuts — overall share, share by engine, and share by prompt cluster — because the averages hide where you're actually winning or losing.
- ♦ You move Share of Model by fixing retrievability, making claims extractable, and building the off-site consensus models read — not by chasing a ranking.
- ♦ Treat it as a program metric reviewed on a cadence; AI answers drift, so a number you earned last month is not a number you keep.

Frequently asked questions

What is Share of Model?

Share of Model is the share of AI-generated answers, across a fixed set of buyer questions and every major answer engine, in which your brand is named inside the answer itself. It is the AI-era equivalent of share of voice: instead of measuring how much of the ad space or the search results you occupy, it measures how much of the machine's answer you occupy.

How is Share of Model different from keyword rankings?

A keyword ranking tells you where a link sits on a results page. Share of Model tells you whether the AI naming a solution names you. Since a growing share of buyers now get a synthesized answer instead of a list of links, being ranked #3 on a page nobody reads is worth far less than being one of the two brands the model actually mentions.

How often should I measure Share of Model?

Re-run your prompt set on a regular cadence — monthly is a sensible default for most businesses, weekly if you are in an active push. Answer engines update, competitors publish, and your own consensus shifts, so a single measurement is a snapshot, not a trend. The value is in watching the line move.

About the author. Scott Tischler is the Founder & Chairman of Alrecommend.ai and a practitioner-authority on AI search and Answer Engine Optimization. With 20+ years in marketing technology — including American Express, MetLife, and UBS — and executive and professional study at Wharton, Harvard, and Oxford, he helps businesses become the ones AI recommends.