

seer
interactive



THE SIGNAL

The AI *Stuff* Marketers Can't Ignore

(and what to do about it)



Alisa Scharf
Chief AI Officer



Nick Haigler
R&D Lead, AI Innovation

SEPTEMBER 2026

WHAT'S HAPPENING?

How reviews impact AI search visibility.

- 01 Trustpilot research review:
How reviews shape AI responses
- 02 Prompt Frequency Pilot:
**How many runs does it take before
AI responses stabilize?**
- 03 Ask us anything

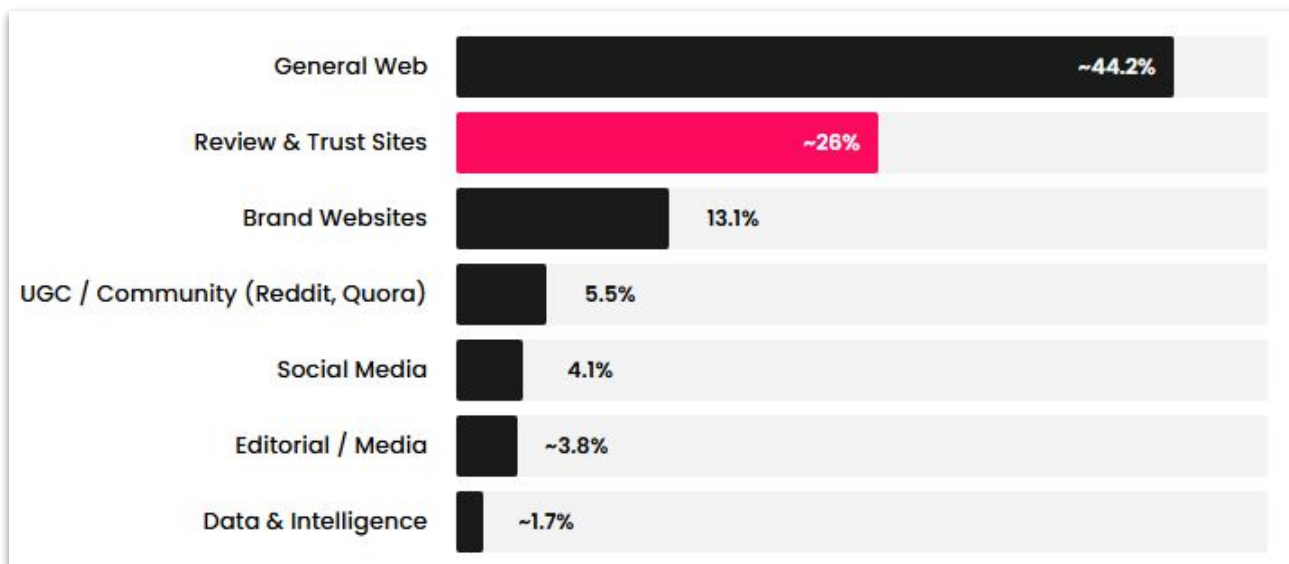


Alisa's Soapbox of the Month

TRUSTPILOT x SEER RESEARCH

How reviews shape AI responses

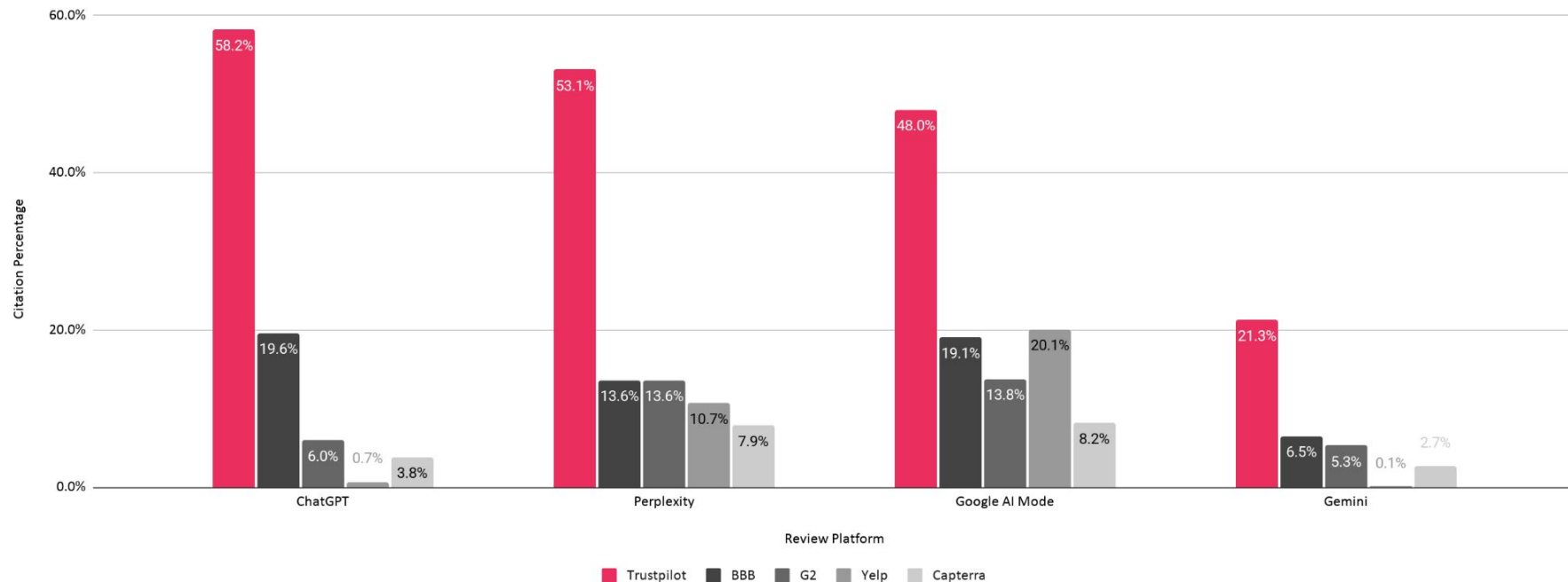
Review and trust sites are the #2 citation source in AI responses at ~24.5% of all citations



Review sites grow from 3.0% of citations at Awareness to 23.5% at Decision.

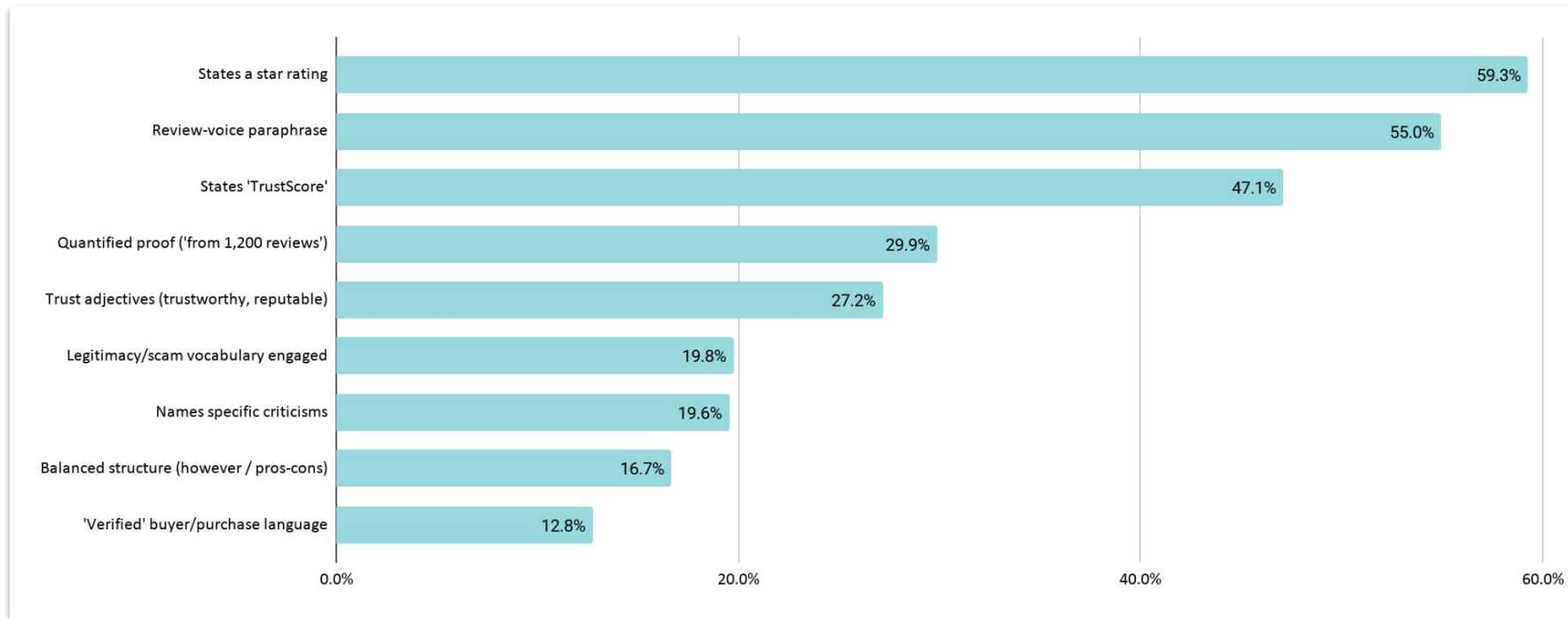
An ~8x expansion in share as purchase intent increases.

Trustpilot leads review platform citations on all four AI platforms tested



Trust signals are most featured in responses

Quantifiable metrics are heavy hitters



Review content amplifies sentiment in AI

Positive amplification

"On Trustpilot, the company has a 4.5/5 rating from tens of thousands of reviews, with many buyers specifically complimenting the build quality and value."

- ChatGPT

Negative amplification

"Trustpilot pages for the [Brand] location show several negative reviews (examples include claims of unmet vehicle promises, deposit or billing confusion, and delays)."

- Perplexity

Comparative amplification

"Trustpilot shows a high overall customer rating (4.8/5 from many thousands of reviews), while BBB customer reviews and WalletHub include substantially more criticism, particularly about collections and loan costs."

- ChatGPT

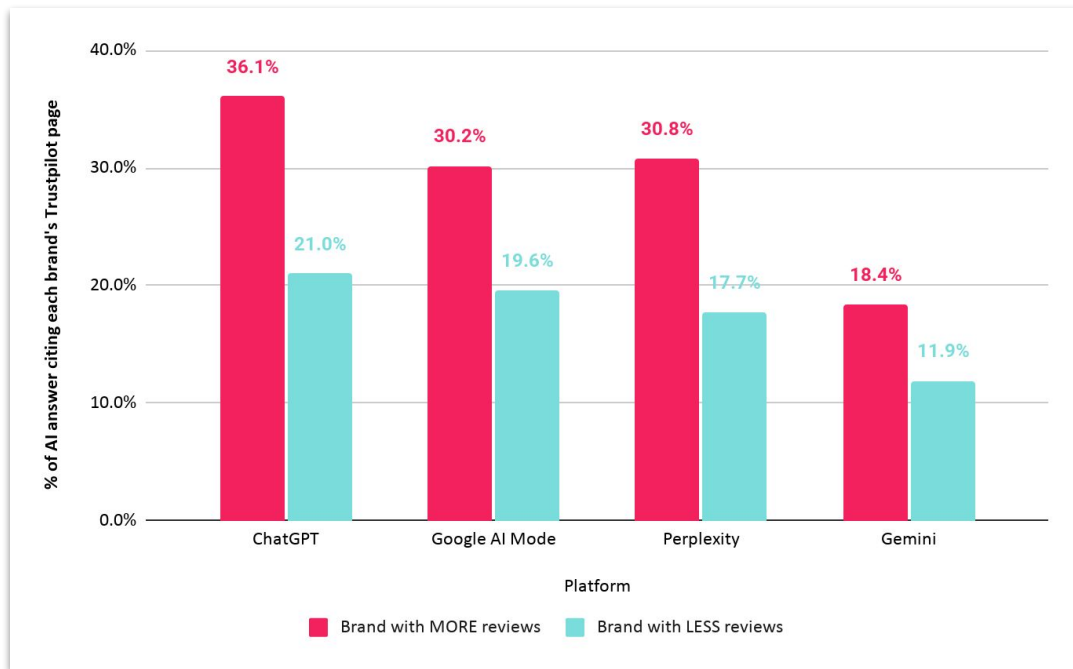
Absence signal

"There is a distinct lack of extensive, verified third-party reviews or mainstream consumer ratings for [Brand] on trusted platforms like Trustpilot or the Better Business Bureau."

- Gemini

“Brand A vs Brand B”

**The brand with
more reviews
gets its
evidence pulled
into AI
responses
up to 1.7× more**



Prompt:

"Which is better, [Brand A] or [Brand B] in Miami?"

"[Brand A] appears to have a stronger public reputation based on a large number of recent candidate reviews (around 4.9/5 overall on Trustpilot, with consistently positive feedback for its Miami office)."

"[Brand B] has much less publicly available employee feedback, making it harder to evaluate as a workplace. Most of the available information comes from the company's own website rather than independent reviews."

PROMPT FREQUENCY PILOT

**How many times do you need to run
a prompt before you can trust the
answer?**

One run shows less than half the field: 5.8 brands per run vs 12.7 across 10 runs

5.8

brands named in a
single run

12.7

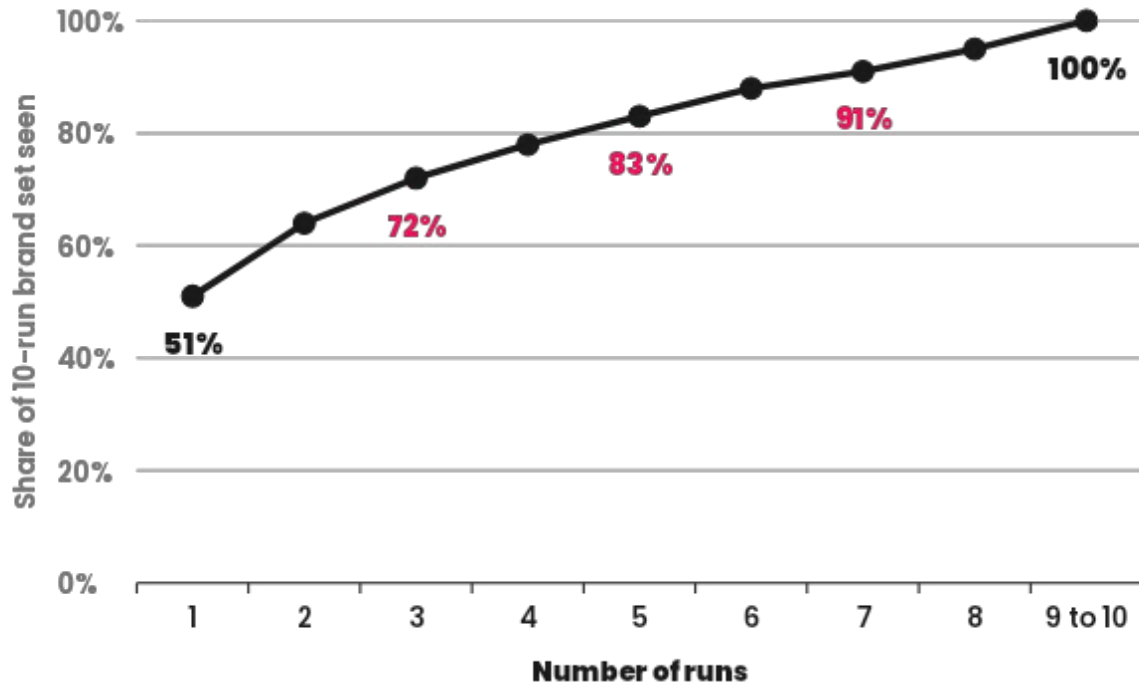
distinct brands named
across 10 runs

29%

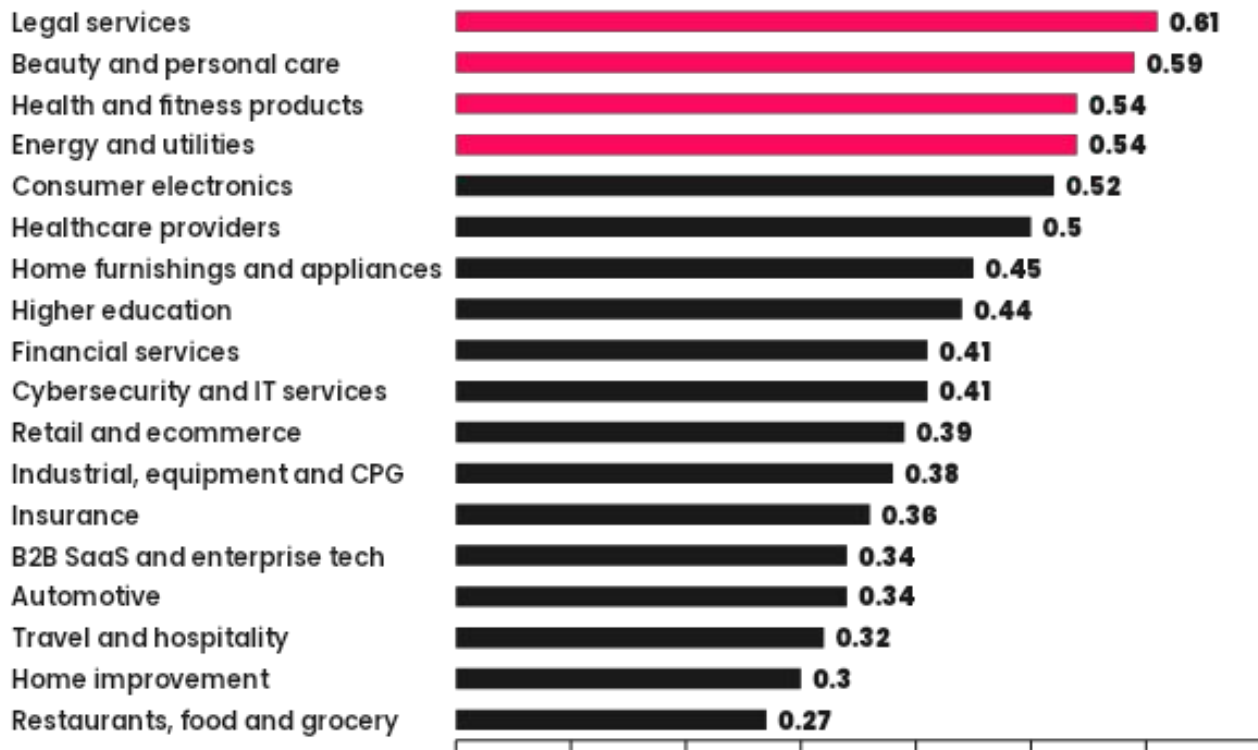
of every brand that ever
appeared showed up in
exactly **one** run out of ten.

If you checked your AI visibility once, you saw a sample, not the field. The brands you missed were still getting recommended.

Three runs surface 72% of the brands ChatGPT will ever name. Seven get you 90%.



Legal and Beauty reshuffle 2x Restaurants and Travel



The floor is 3 runs per prompt. If you need one number, use 5.

3 to 4

RUNS

Home improvement, retail, restaurants, travel, B2B SaaS, insurance, automotive, healthcare providers, energy, cybersecurity

5

RUNS

Industrial and CPG, financial services, home furnishings and appliances

6 to 7

RUNS

Higher education, consumer electronics, legal services, health and fitness products, beauty and personal care

Every industry had a prompt that needed 8 or 9 runs. Treat these as starting points, then let a baseline burst reassign individual prompts. Add roughly 2 runs per tier for 90% coverage.

The #1 pick is stickier than the list. The default held while the roster churned.

75%

of runs repeated the same lead
recommendation

8 of 90 prompts named zero brands, run after run

- “Best HVAC company for a new AC system”
- “Best personal injury lawyer for a car accident”
- “Best law firms for medical malpractice”
- “Best electricity supplier for fixed rates”
- “Best energy supplier for a small business”
- “Best probiotics for gut health”
- “Best kitchen remodeling company”
- “Best shampoo for thinning hair”

The pattern: hyper-local services and medically sensitive products.

The model asks where you live, or says see a doctor, instead of naming a brand.

Appearing at all is the question. That makes localized and persona prompts mandatory for these categories.

Be Seen. Be Believed. Be Chosen.

BE SEEN

THE QUESTION

Do the answers include
you?

BE BELIEVED

THE QUESTION

Are the answers that
include you accurate?

BE CHOSEN

THE QUESTION

Does your audience act
on it?

... SO, NOW WHAT?

THREE things to do this month

- ☐ Analyze what review platforms are getting cited for your brand's evaluation prompts.
- ☐ Claim your review profile and monitor the themes and sentiment. Those are likely being pulled through into AI responses.
- ☐ Run a 10x baseline burst on 20 to 30 prompts. At about two cents a run, it costs less than lunch. Classify each prompt as locked, mixed or volatile.

Ask us anything

Let's talk about what this means for your team.