

What 5,055 ChatGPT conversations reveal about AI visibility

A structured analysis of 131 B2B SaaS brands

July 2026

Powered By Viziquo

AI visibility is not one score

A brand can be known to the model, rarely discovered organically, weakly cited, and cautiously framed—all at once.

12.2%

Organic discovery

The tested brand appeared in only 364 of 2,995 organic conversations.

43.9%

Competitor-owned citations

In organic answers where the brand appeared, competitor sites supplied the largest citation bucket.

$r = -0.027$

Visibility vs positive tone

Higher organic visibility did not reliably predict more favorable overall framing.

A broad, structured test across the B2B SaaS market

The goal was to observe patterns across question types—not to crown a winner.

131

B2B SaaS brands

5,055

Conversations

12

Business functions

20 days

July 8–27, 2026

CONVERSATION CONTEXT

2,995 organic • 59.2%

1,065 prompted

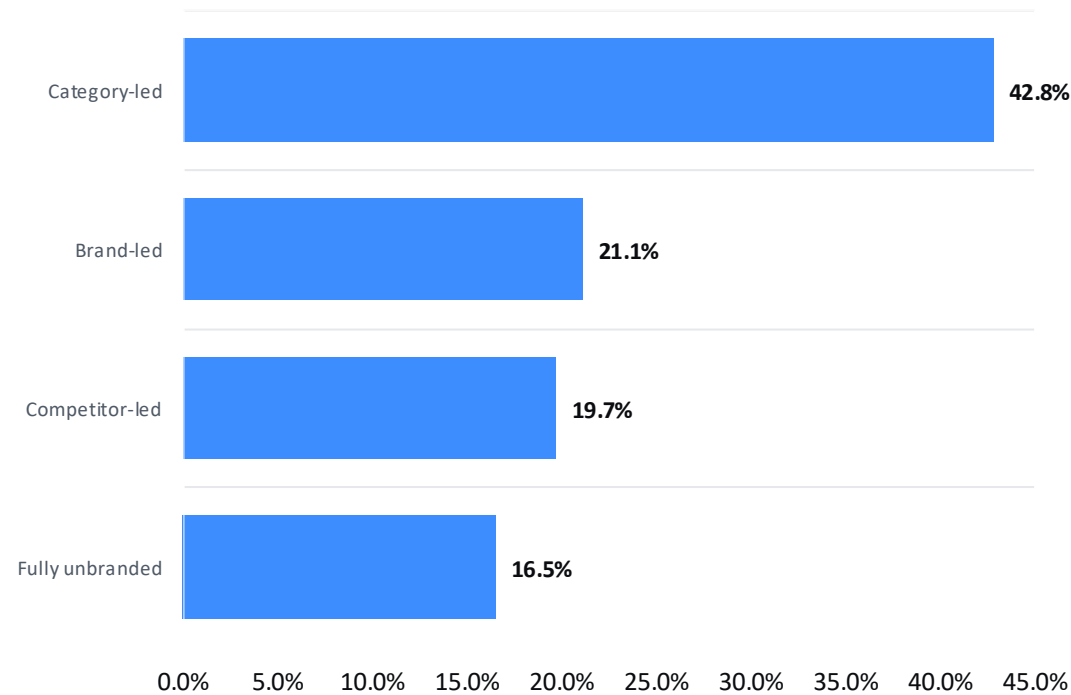
995 competitor-led

77.3% single-turn • 22.7% multi-turn • Synthetic, structured tests—not real customer logs

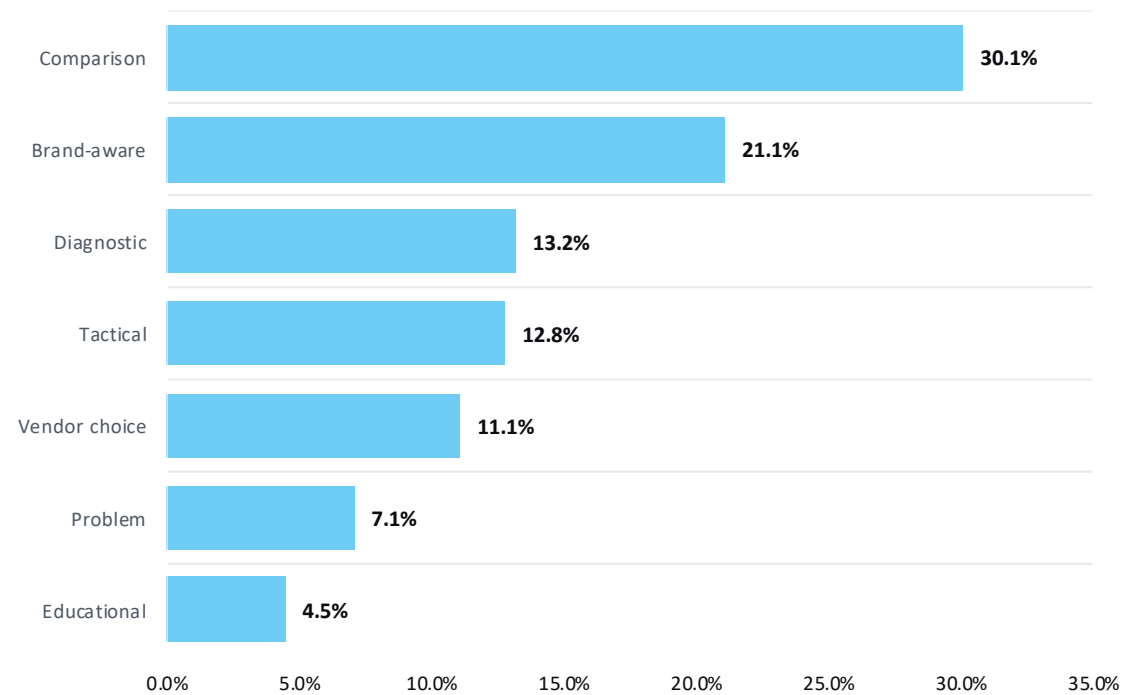
The test set covered both discovery and consideration

Most questions were category-led or comparative, but the corpus also included fully unbranded and educational prompts.

QUESTION CONTEXT / BIAS

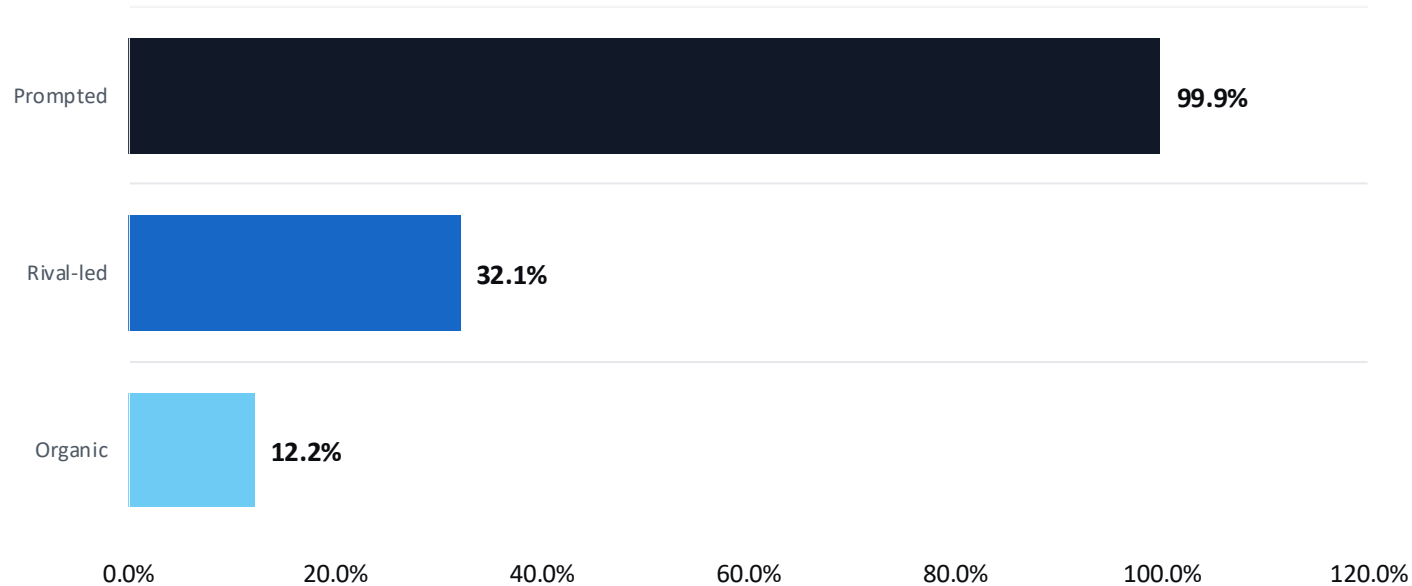


BUYER INTENT



Being known is not the same as being discovered

Explicitly naming the brand almost guaranteed appearance. Organic prompts told a very different story.



29 / 131

brands had zero organic appearances

78 / 131

had organic visibility $\leq 10\%$

91.7 pp

median prompted–organic gap

Category language acts as a discovery bridge

A prompt did not need to name a brand—but it usually needed to give the model enough category context.

5.4%

Fully unbranded organic

45 of 833 conversations

14.8%

Category-led organic

319 of 2,162 conversations

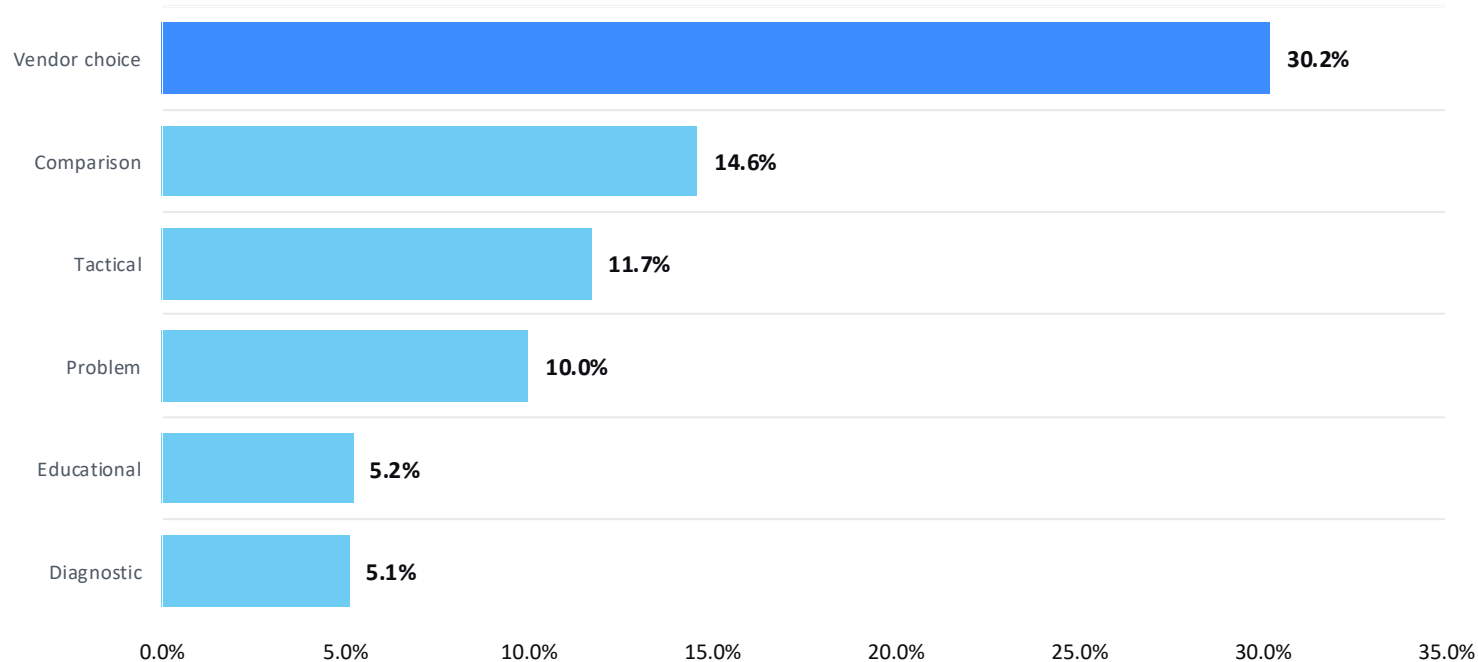
2.7×

higher appearance rate

Marketing implication: category clarity is part of AI discoverability. Broad problem language often leaves the model with no reason to introduce a vendor.

Visibility rises when the buyer is closer to a decision

Organic brand appearance varied sharply by question intent.



30.2%

Vendor-selection prompts

The highest organic appearance rate: 91 of 301 conversations.

5.1%

Diagnostic prompts

Brands were rarely introduced when the user asked for diagnosis rather than options.

Recommendation prompts create the strongest opening

Question mode mattered even after separating organic prompts from brand-prompted tests.

21.5%

Recommendation-seeking

278 of 1,294 organic conversations

5.1%

Informational

86 of 1,701 organic conversations

4.2× higher

Appearance does not mean ownership of the narrative

Among the 364 organic answers where the tested brand appeared, competitor-owned sources supplied the largest share of citations.

43.9%

Competitor-owned

771 of 1,757 citations. Competitor pages often helped define the category and comparison frame.

38.6%

Third-party

679 citations from media, communities, review sites, directories, and other independent sources.

17.5%

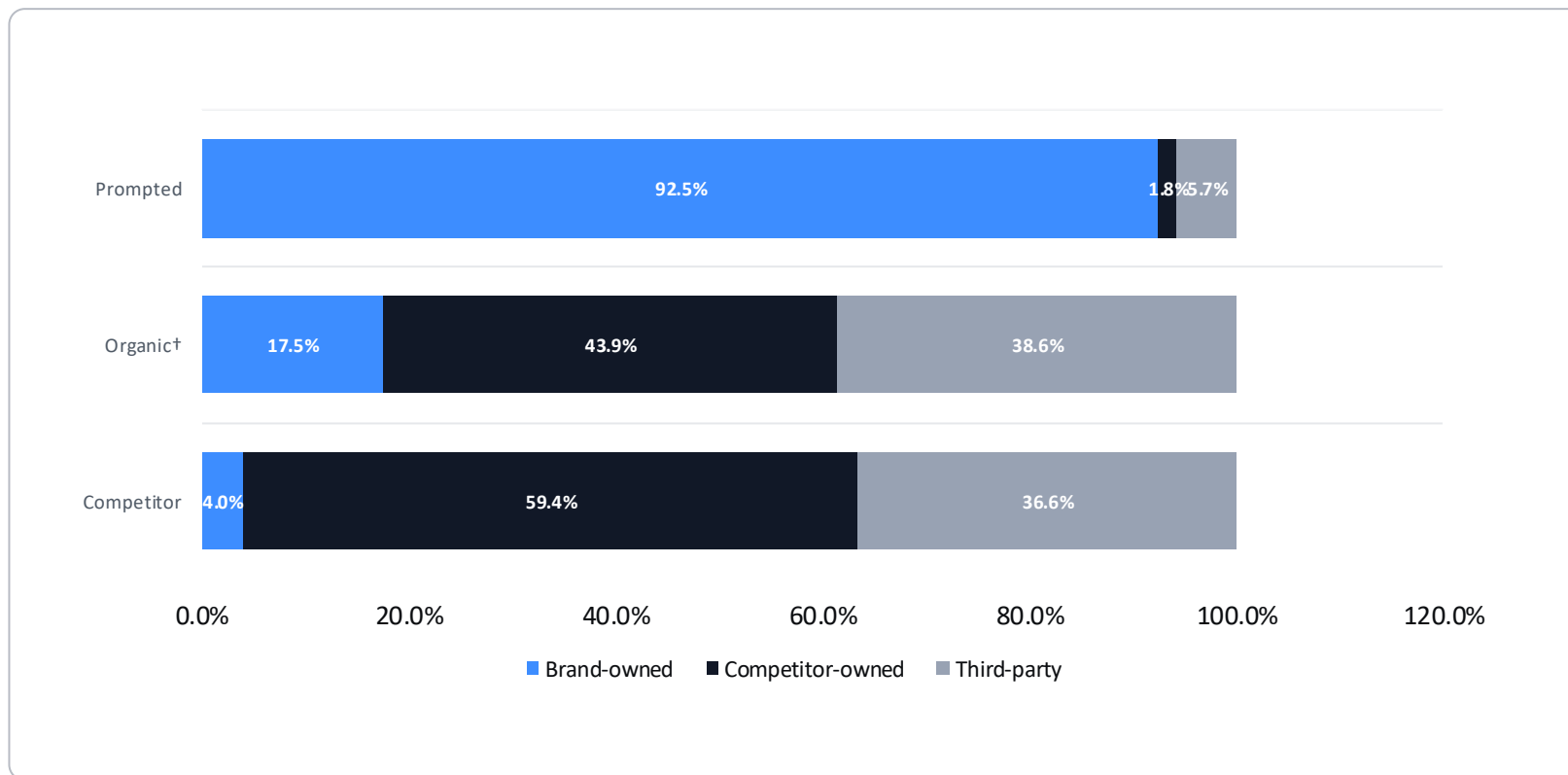
Brand-owned

307 citations from the tested brands' own domains—the smallest citation bucket.

27 of 102 organically visible brands received no brand-owned citation at all.

Prompt context changes which sources the model trusts

Citation mix changes dramatically depending on whether the user names a brand, stays organic, or starts with a competitor.



PROMPTED

92.5% brand-owned

ORGANIC, BRAND PRESENT

43.9% competitor-owned

COMPETITOR-LED

59.4% competitor-owned

Same model. Different prompt context.
Different evidence base.

Most missed visibility was not a competitor win

The tested brand was absent from 2,631 organic conversations. In most, the model still did not introduce a competing vendor.

2,631

organic conversations without the tested brand

38.2% competitor present

61.8% no competitor entity

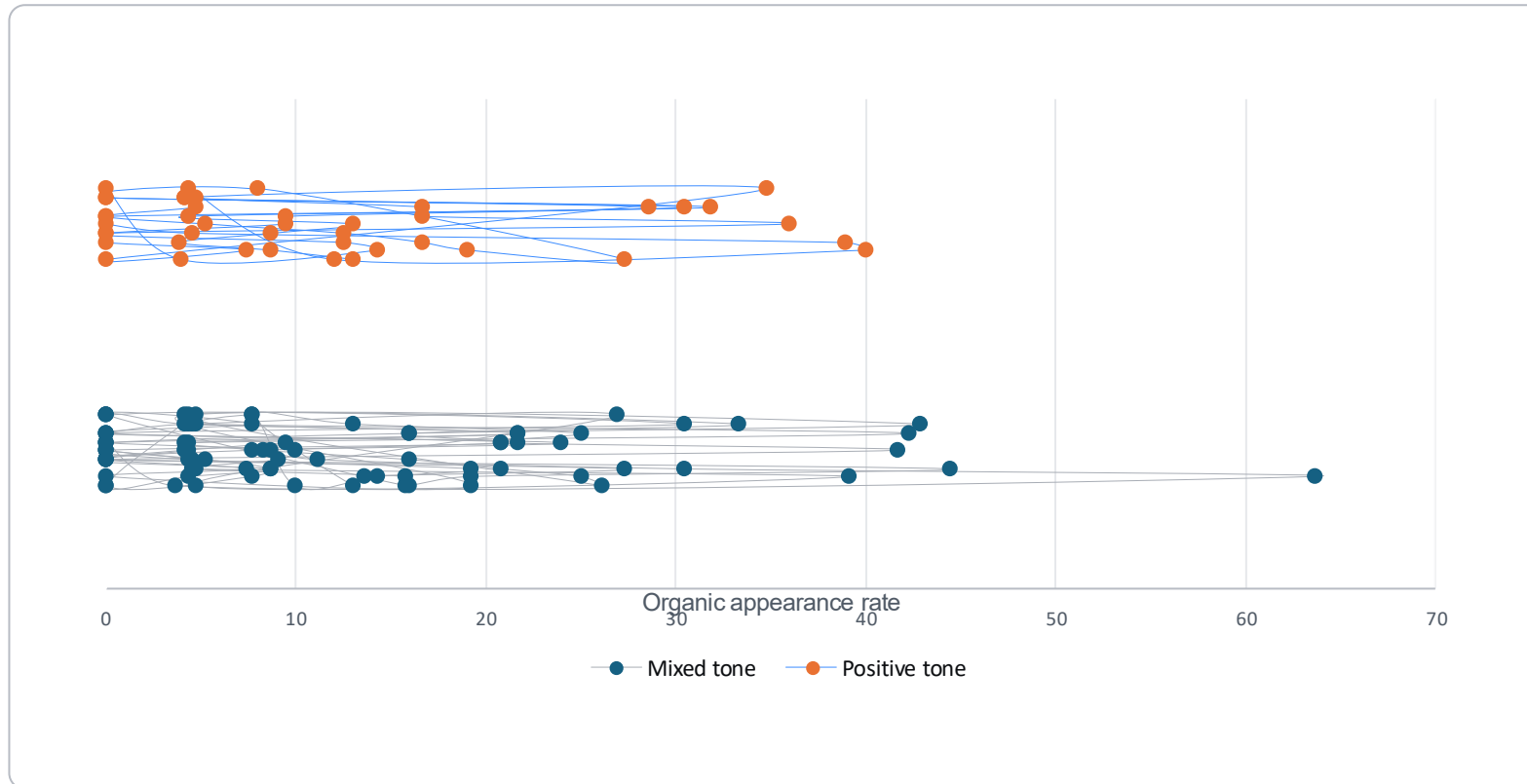
The larger opportunity

In many missed prompts, the model stayed at the level of general advice. The job is not only to beat a named competitor—it is to make the category and use case legible enough for the model to introduce any vendor.

30.6% cited a competitor

Higher visibility did not guarantee favorable framing

Brands with positive and mixed overall tone appeared across the visibility spectrum.



$$r = -0.027$$

organic visibility vs positive overall tone

Top visibility quartile

24 mixed • 9 positive

Recommendation pattern correlation: $\rho = 0.135$ (weak)

The core patterns were not driven by one SaaS category

We checked the direction of the findings across 12 business functions and within a more balanced 88-brand cohort.

Pattern checked	Direction held	Scope / note
Recommendation-seeking > informational	12 of 12	Every business function
Competitor-owned > brand-owned citations	12 of 12	Brand-present organic answers
Category-led > fully unbranded	11 of 12	Exception: AI Agents (4 brands)

Balanced cohort

88 brands × 40 conversations: 11.9% organic appearance • 5.6% unbranded • 14.3% category-led • 20.9% recommendation-seeking • 5.1% informational

Multi-turn helps consideration more than discovery

The first answer did most of the brand-discovery work. Later turns added evidence and citation depth.

92

Brand first appeared on turn 1

Among 718 organic multi-turn conversations, most discovery happened immediately.

6 / 626

Brand appeared only later

Just 1.0% of conversations that lacked the brand on turn 1 introduced it later. Only two later introductions were recommendations.

355 / 718

New citations appeared later

49.4% of organic multi-turn conversations added at least one new citation after the first answer.

DISCOVERY

REFINEMENT

EVIDENCE & FRAMING

Measure AI visibility as four separate questions

A single visibility score hides where the real problem sits.

01 DISCOVERY

Are we introduced organically?

Track fully unbranded and category-led prompts separately.

02 RELEVANCE

For which intents and personas?

Segment by decision stage, use case, persona, and turn.

03 AUTHORITY

Whose sources shape the answer?

Separate brand-owned, competitor-owned, and third-party citations.

04 FRAMING

What does the model say about us?

Audit tone, recommendation strength, caveats, and repeated narratives.

How we built the test set

A four-step workflow turned each brand's website into buyer-style conversations.

COHORT

131 B2B SaaS brands across five vertical groups: Sales & Revenue; Marketing & Growth; Demo, Content & Visual Communication; AI Engineering, Evaluation & Security; and Customer Success, Support & Retention.

01

Brand profiling

Built a detailed profile for each brand from 15–20 pages on its website.

02

User profiles (ICPs)

Created 3–5 profiles per brand, capturing the user's role—not job title—their problems, and the language they would or would not use with ChatGPT.

03

Question generation

Generated 20–40 buyer-style conversational questions per brand, with a healthy mix of organic and brand-led prompts. We avoided keyword-first questions.

WE CHANGED THE PATTERN

“Best software in [category]”



“Can you recommend software for [user problem]?”

“[Brand] vs [competitor]”



“Between [brand] and [competitor], what would you recommend for [user problem]?”

1,150

Multi-turn conversations

22.7% of the dataset • adaptive follow-up • maximum three turns

Every conversation began with a buyer-style question. Most ended after the first answer; follow-ups were added only when the response warranted deeper exploration.