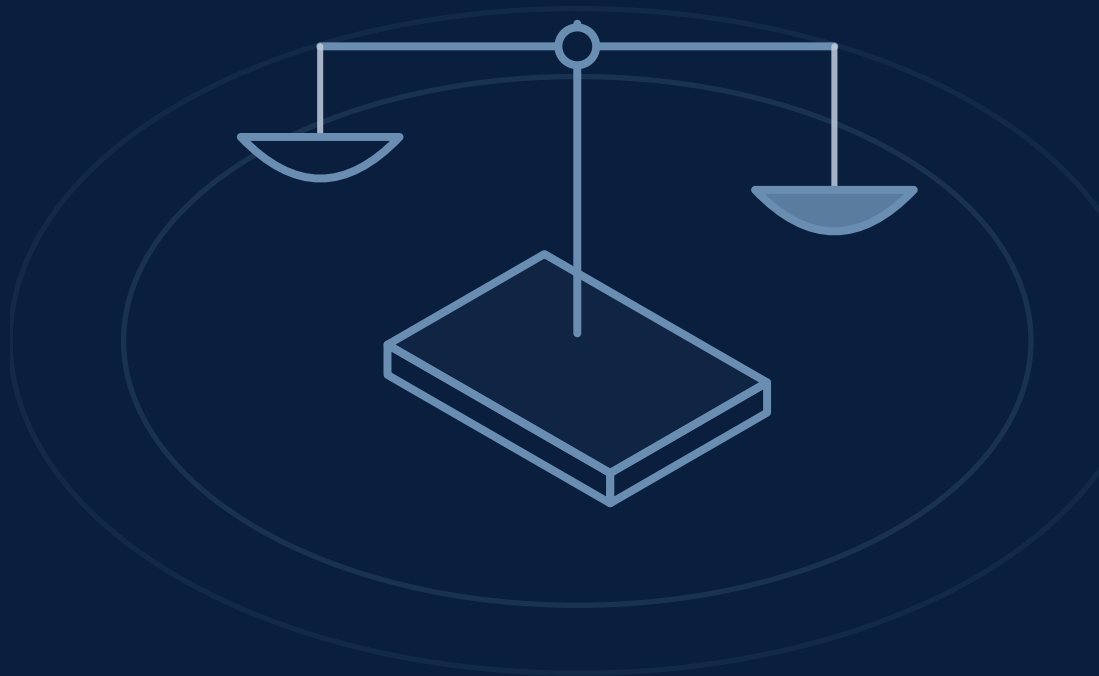


Resonance

RESEARCH · 2026

What AI gets wrong about B2B tech brands.

And where the mistakes come from.



We checked almost 50,000 AI answers against the facts.

Between April and September 2026 we asked ChatGPT, Gemini, Perplexity and a few others the kind of questions B2B buyers ask about 116 technology companies. Then we checked every answer against a pack of verified facts about each company.

We wanted to know how often the assistants get a company wrong, and where the mistakes come from when they do. If the source is your own website, you can fix it this week. If it's someone else's page, or no page at all, you need a longer plan.

THE STORY IN TEN HEADINGS

01	Most answers get something wrong, usually something small	3
02	The most common mistake is confident detail nobody published	4
03	Errors about company facts and pricing are rarer but more often serious	5
04	AI leaves out more than it gets wrong	6
05	Most errors don't come from any page	7
06	When a page is to blame, it's more often your own, but third-party errors do more harm	7
07	Some assistants, and some kinds of question, are worse than others	8
08	It's getting harder to pass	9
09	What we'd fix first	10
10	How we checked	11

Most answers get something wrong, usually something small.

Two in three answers we checked had at least one mistake about the company. Most were minor, like a feature described too broadly or a figure nobody can trace.

65%

of answers had at least one error about the company.

15%

had an error our judge rated serious. That's roughly one answer in seven.

4.5

verified facts left out of the average answer, against 1.7 errors.

AVERAGE SCORE OUT OF 100

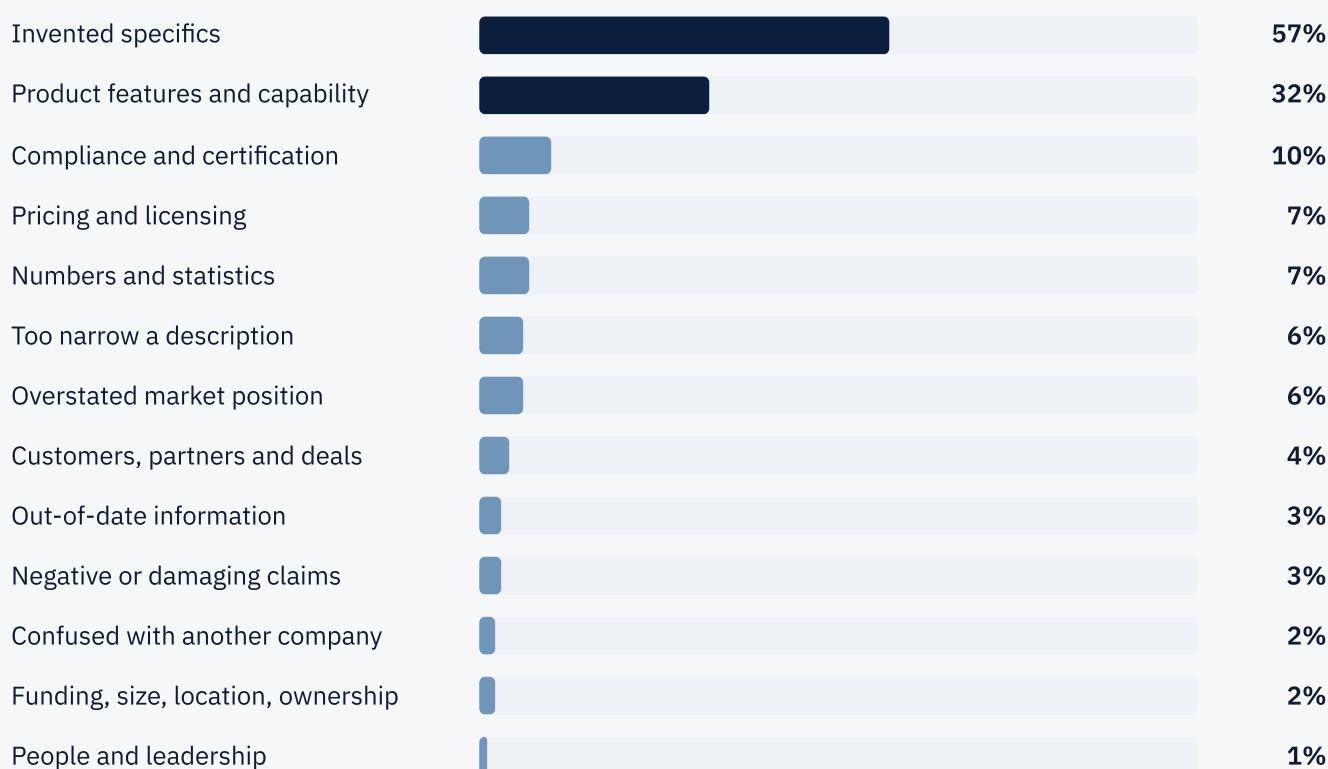


The assistants are generally warm about the companies they describe, and sentiment scored 82 out of 100. They slip on the detail. Facts, completeness and recency all came in at around 70. In practice that's an answer that sounds confident and is mostly right, but often describes the company as it was a year or two ago.

The most common mistake is confident detail nobody published.

We sorted more than 80,000 errors by what they were about. In more than half, the assistant had added specifics the company has never put in public.

SHARE OF ERRORS BY TOPIC (ONE ERROR CAN COVER MORE THAN ONE)



One answer credited an automation platform with cutting a customer's audit time by 75%. Another said an AI infrastructure company guarantees 98% service quality. Neither company has ever published those numbers.

It happens because buyers ask precise questions, like whether a product works with their stack or is certified for their sector. When the company hasn't answered that in public, the assistant writes a plausible answer anyway.

Errors about company facts and pricing are rarer but more often serious.

Only about one error in fifty is about funding, size, location or ownership. When an assistant gets one of those wrong, though, our judge rated it serious almost a third of the time.

SHARE OF EACH KIND OF ERROR RATED SERIOUS



Across all errors, 13% were rated serious.

Assistants get positioning wrong in both directions.

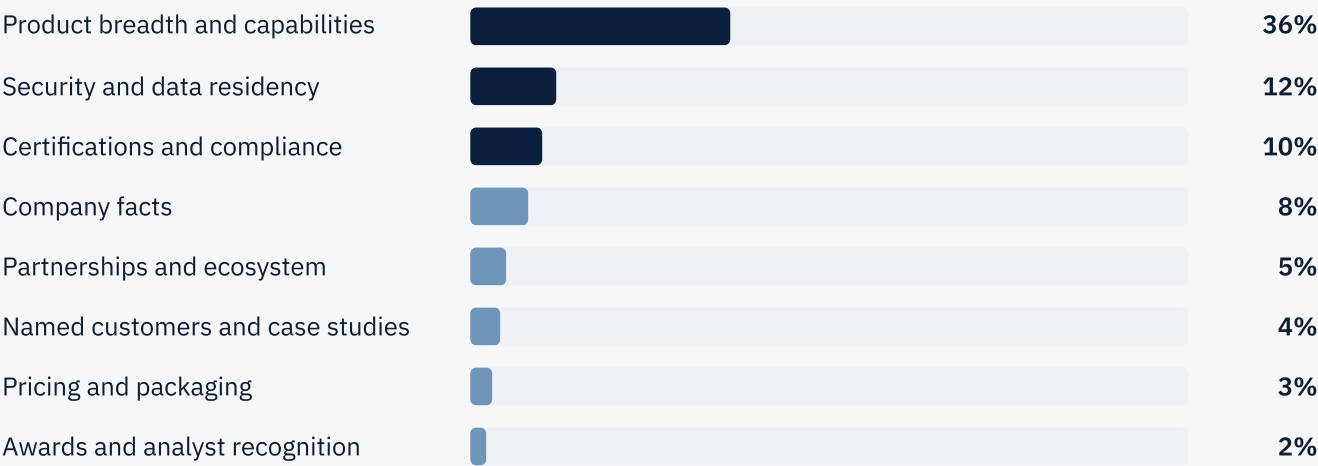
About a fifth of errors are about how a company is positioned. Some shrink it. We saw a broad automation platform described as mainly an RPA tool, a label it moved beyond years ago. Others inflate it, like the fund data provider called the most widely used network in Europe.

The flattering ones are easy to shrug off. A buyer who checks the claim and finds nothing behind it will trust the rest of the answer less, and that includes the parts about you that were right.

AI leaves out more than it gets wrong.

We logged more than 200,000 verified facts that answers should have included and didn't. That's four or five per answer, more than twice the number of errors. Nobody reading the answer can spot what isn't there.

SHARE OF MISSED FACTS BY TOPIC



The rest cover a wide mix of topics that don't fit one group.

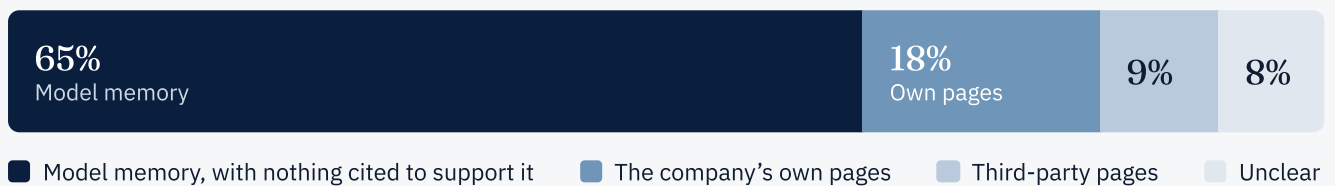
The biggest gap is breadth. About a third of missed facts are about the rest of the product. The assistant describes what a company was known for and skips what it has added since, which is where a lot of the pigeonholing on the previous page comes from.

Security, data residency and certifications make up another fifth. For buyers in regulated sectors those can decide a deal, and they're some of the easiest facts for a company to publish clearly.

Most errors don't come from any page.

For each error, our judge looked at what the answer cited and recorded where the claim most likely came from. It only blamed a page if that page plainly supported the claim.

WHERE ERRORS CAME FROM (ABOUT 59,000 ERRORS, JUNE TO SEPTEMBER)



About two thirds of errors had nothing behind them. The model knew the company vaguely and filled in the rest from memory. You can't correct a page that doesn't exist, so the only fix is for the web to tell a consistent story for long enough that the models pick it up.

When a page is to blame, it's more often your own, but third-party errors do more harm.

A company's own pages were behind 18% of errors, about twice as many as third-party pages. Most of those were the assistant stretching something the company did say, such as applying a figure more widely than it was meant. Only 7% were serious.

Errors from third-party pages were rarer and worse, with 15% rated serious, the highest of any source. The usual culprit was an old review, directory listing or aggregator page that nobody had updated.

SHARE OF ERRORS RATED SERIOUS, BY SOURCE



Some assistants, and some kinds of question, are worse than others.

We checked more than 15,000 answers each from ChatGPT, Perplexity and Gemini, which is enough to compare them fairly. The gap was wider than we expected.

ASSISTANT	FACTUAL SCORE	ANSWERS WITH AN ERROR	WITH A SERIOUS ERROR
ChatGPT	79	49%	5%
Perplexity	71	63%	16%
Gemini	62	83%	27%

Gemini got something wrong in more than eight answers out of ten, and in about a quarter the mistake was serious. ChatGPT was wrong in about half. The sources differ too. ChatGPT leans hardest on companies’ own sites, so more of its mistakes are misreadings of what you published, while Perplexity draws on a wider spread of third-party pages.

WHERE EACH ASSISTANT’S ERRORS CAME FROM



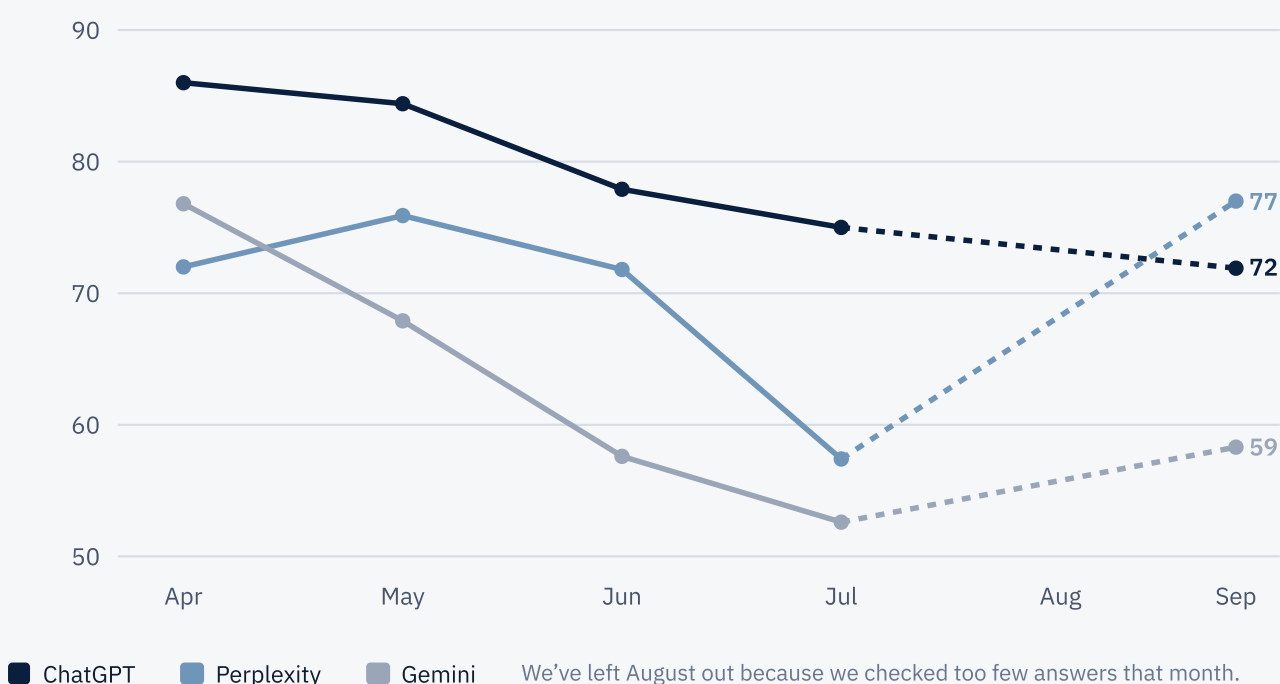
FACTUAL SCORE BY KIND OF QUESTION, THE LOWEST THREE AND THE HIGHEST



It's getting harder to pass.

Factual scores fell for all three assistants between spring and summer. ChatGPT slipped from 86 in April to 72 in September. Gemini dropped sharply in July and has partly recovered, and Perplexity was back up to 77 by September.

AVERAGE FACTUAL SCORE BY MONTH, APRIL TO SEPTEMBER 2026



Some of the fall is down to us. The fact packs we check against got more detailed over the six months, so there was more to catch. The order never changed, though, with ChatGPT ahead and Gemini behind all the way through.

Recency scored lowest of everything we measured, and about one error in ten turned up again in later answers. Once something wrong gets into an assistant's picture of you, expect it to stay there until the sources change.

What we'd fix first.

Most errors fill a gap, so the work is closing gaps in public.

We'd start with your own site because it's quickest, then move on to the third-party pages behind the worst errors.

01 Publish the details buyers ask about.

Your pricing approach, certifications, integrations, data residency and customer numbers. Put them in plain text on a dated page so an assistant can quote you instead of guessing.

02 Read your own pages the way an assistant would.

Look for loose claims that could be stretched, old product names and figures with no context. This is the quickest win in the report.

03 Keep a facts page and an up-to-date newsroom.

One page that says what you do, who uses it and what you're certified for, and a newsroom that shows what has changed and when.

04 Correct the third-party pages with the most reach.

Review sites, software directories and reference pages cause the most serious errors. Update your profiles and ask for corrections where you can.

05 Describe yourself the same way everywhere.

If assistants pigeonhole you, it's because parts of the web still do. Use your current category in press releases, analyst briefings, bylines and partner pages, and look for trade coverage that states the facts.

06 Check again every month.

Ask the questions your buyers ask, in the assistants they use, and keep a list of the errors that come back. Memory moves slowly, so judge progress over a quarter.

How we checked.

The answers

Agency, our AI visibility platform, collected almost 50,000 answers about 116 B2B technology companies between April and September 2026. The questions are the ones buyers ask, from early research through to comparing vendors. Most answers came from ChatGPT, Gemini and Perplexity, with smaller samples from Claude, Copilot and Google's AI features.

The facts

Each company has a pack of verified facts covering what it does, its products, certifications, customers and partners, and how it describes itself. We build the packs from the company's own material and people check them.

The judge

An AI judge compared every answer with the fact pack. It logged each wrong claim with a severity and noted whether it was about fact or positioning and whether it had come up before. It also logged the verified facts each answer missed.

Where errors came from

From June, the judge also recorded the most likely source of each wrong claim, based on what the answer cited. We told it to be conservative, and to say "unclear" rather than guess.

The topics

We grouped errors and missed facts by the words used in each claim and correction. One error can cover more than one topic, so the shares add up to more than 100%.

What to bear in mind

An error here means the fact pack doesn't support the claim. Some are plainly false and some may happen to be true but can't be checked. The companies lean towards enterprise software, AI infrastructure, fintech, cybersecurity, sustainability software and data. The fact packs got more detailed over the period. We're reporting patterns, and we don't name any company. The examples are composites.

Resonance

ABOUT THIS REPORT

Find out what AI says about you before your buyers do.

Resonance is a strategic communications consultancy that began in B2B technology PR. We help technology companies get their story straight across media, analysts and the AI assistants buyers now ask first.

Agency, the PR operations platform Resonance built, checks what AI assistants say about brands every week, and where they get it from.

hello@resonancecrowd.com

resonancecrowd.com/research

agency.com.pr

London · Manchester

Resonance