

We read 252 AI answers about real estate agents. Every number below is checkable.

By Amit Turkaspa · Founder, Pitched · July 2026

1 in 3

names in AI answers belong to agents the AIs mentioned once and never again. Those seats are up for grabs.

We asked ChatGPT, Claude, Gemini and Perplexity the questions real buyers and sellers ask. "Who should list my home?" "Who are good buyer's agents here?" We did this in six US markets and read all 252 answers, word for word. Almost 750 agents were named. Almost none of them show up consistently. Here is what we found.

Your market, all four AIs, every question a client would ask. The same scan behind this study.

252 AI answers we read in full	749 agents and teams named at least once
52% were named one single time, then never again	~12 agents per market the AIs name again and again

FINDING 01

One mention is luck. Consistent mentions are earned.

The agents the AIs recommend again and again are recommended for the same reasons every time: sales records, reviews, awards and rankings, published guides. Real proof.

Everyone else got named by accident. One directory listing. One old blog post. One Reddit comment that happened to get picked up. That is why half the agents in our data were named exactly once and never again, and why a third of all the names the AIs handed to clients belong to agents like these.

Now look at how few agents are actually competing. The Realtor associations covering our six markets count roughly 150,000 members combined. Thousands of them work these neighborhoods. Yet only about a dozen per market are named again and again. Downtown Manhattan makes it concrete: across our entire scan, six agents earned a regular seat. Not because the seats are full. Because almost nobody has given the AIs a reason to say their name yet.

Bottom line: "an AI mentioned me once" means nothing. Below a certain level of proof, whether you show up is a coin flip. Above it, the AIs name you on purpose. And the list of agents above that line is still short.

FINDING 02

The most popular AI-visibility trick gets you named with a warning label attached.

Agents are being taught right now to publish pages like "Who's the best agent in [city]?" and answer with their own name.

Our scans caught what happens next. Those pages do get you named. But on two of the four AIs, your name shows up with a warning next to it. In front of the client:

[CHATGPT] ASKED: "WHO ARE THE TOP BUYER'S AGENTS FOR DOORMAN CONDOS IN TRIBECA?"

"I'd treat this as self-promotional but still potentially relevant if you want careful comp analysis."

The source ChatGPT was warning about: the team's own "best agent in Tribeca" page. Named, recommended, and flagged, all in one sentence.

[CLAUDE] ASKED: LISTING AGENTS WHO PRICE HONESTLY IN A SOFTENING MARKET

*"What search results *do* surface are self-promotional agent sites, SEO-optimized 'best of' listicles, and volume/production rankings - none of which measure honesty in pricing."*

[CLAUDE] ASKED: AGENT RECOMMENDATIONS FOR A FIRST HOME IN THE MISSION

"I don't want to just repeat agent self-promotion back to you as if it were a real endorsement."

The exception is Gemini. In all six markets, Gemini took agents' own testimonials and marketing claims at face value, sometimes as proof of expertise or even honesty, without any warning.

Bottom line: the same page that gets you endorsed on Gemini gets you flagged on ChatGPT and Claude. Which AI your next client happens to ask decides which version of you they meet.

The winners had two kinds of evidence.

Why did the AIs recommend who they recommended? Two reasons came up again and again.

The first: proof written by others. Sales records on the portals, reviews, rankings, press. The AIs leaned on it in every market. It takes time to build, because someone else has to publish it, but every consistent winner had it.

The second: published guides that genuinely help the client. School-zone guides, condo explainers, pricing breakdowns. Pages about the client's decision, not about the agent. This one is fully in your hands, and the AIs treated it as real evidence on all four platforms, in all six markets.

[PERPLEXITY]

ASKED: AGENTS WHO CAN EXPLAIN CONDO VS. TIC TRADEOFFS TO FIRST-TIME BUYERS

"This kind of content indicates she is accustomed to educating buyers."

One answer named eight different agents, each one purely because of a published condo-vs-TIC guide. The guide wasn't supporting evidence. It was the entire reason.

[CHATGPT] ASKED: ABOUT BRICKELL BUYER'S AGENTS

"I'd shortlist buyer-side agents who publish actual condo due-diligence content, not just 'Brickell luxury' marketing."

The seats are winnable. And the bar will never be this low again.

On top of the regular questions, we pushed each AI with one extra challenge per market: give us a numbered top-20 list of agents. 24 lists requested. 480 seats total. Two things happened.

First, three times the AI refused or gave up partway. It would not put 20 names on a list it could not back up.

Second, in the lists that did arrive, **almost 3 entries per list were not an actual, identifiable agent.** "Wendy," pulled from a review snippet. "The agent behind this

Instagram reel," no name given. Generic instructions for finding an agent yourself, numbered as if they were agents. To be clear, the AIs were not inventing fake people. They were running out of real ones. Of the 480 seats we asked for, 112 came back empty or filled by placeholders.

Now add the agents who were named only once, by accident. That is more than half of everyone in our data. The AI happened to see their name on some webpage once, and nothing holds them in that answer. They lose the seat the moment another agent gives the AI a real reason to say their name instead.

That is the board today: about a quarter of the seats empty, half held by accident. Right now you are competing against "Wendy." Every month, more agents figure out that AI answers are becoming the first impression. Soon you will be competing against agents who did the work.

The benchmark

Six markets, side by side.

Each market was scanned through a realistic client scenario. "Consistent winners" are agents or teams named four or more times in their market.

Market	Client scenario	Answers	Agents named	Consistent winners	Named once
Frisco, TX	Relocating family, school-zone buyer	40	126	13	60%
Brickell, Miami	Remote professional, condo buyer	44	127	12	57%

Market	Client scenario	Answers	Agents named	Consistent winners	Named once
Downtown Manhattan	Finance professional, condo buyer	40	136	6	49%
Chicago Gold Coast	Luxury condo seller	44	91	18	47%
Ballantyne, Charlotte	Move-up seller, softening market	40	129	19	43%
Noe Valley / Mission, SF	First-time tech buyer, condo vs. TIC	44	140	12	54%

Seller markets (Chicago, Ballantyne) had the tightest winner groups. The deepest buyer market (Manhattan) had the smallest one: six agents. Every figure comes from the same data layer delivered in Pitched customer reports.

Definitions

The terms, in plain words.

Discovery questions

The "who should I hire" questions a buyer or seller asks before picking an agent. Ten to eleven per market in this study.

Named

The agent appeared in at least one AI answer in their market.

Single-platform

Named by exactly one of the four AIs. 61% of all named agents.

Named once

Appeared one single time across the entire market scan. 52% of all named agents.

The headline stat (1 in 3)

Counted at the mention level: of all agent mentions across a market's answers, the share belonging to agents named exactly once in that market. About 1 in 3 across the study.

Consistent winner

Named four or more times in their market. Between 6 and 19 agents per market cleared this bar.

Pre-discounted

Named, but with a credibility warning attached to the source ("self-promotional," "treat as marketing"), delivered to the client next to the name.

The four reasons

If the seats are winnable, why is almost everyone invisible?

We read all 252 answers. Four failure patterns explained nearly every missing name. None of them is fixed by louder marketing.

01. Each AI builds its own picture of you, from different sources.

In our scans, ChatGPT leaned on portal sales records. Perplexity leaned on local rankings and articles. Gemini leaned on agents' own content. Claude leaned on whatever it could verify. Four diets, four different shortlists. If you are strong in one diet and absent from the others, you show up on one AI and vanish on three. That is exactly what happened to 61% of everyone named.

02. The proof AI trusts, most agents never put in order.

Verified sales counts. Reviews on several sites. Third-party recognition. Guides that answer a real client question. This carried every consistent winner in all six markets. Most agents have some of it. Scattered, outdated, or buried in a portal profile the AIs half-read.

03. The AIs get basic facts about agents wrong.

In our scans, the AIs disagreed about which brokerage some billion-dollar teams work for. One answer misspelled a top agent's name. Every error like this breaks the link between your name and your track record, and makes the AIs less likely to bring you up. If it happens to the biggest teams in the country, it happens to you. And you would never know without checking.

04. Your treatment changes. By AI, by market, by quarter.

The same AI refused to produce a top-20 in one market and padded one with anonymous social posts in another. Sources get re-crawled. Competitors publish. Models update. Visibility is not a box you tick once. It moves, which is why it has to be measured, not assumed.

Every one of these four is fixable. The agents who win will not be the ones who shout loudest. They will be the ones who put verifiable proof where the machines actually look.

What actually works

The five signals the AIs cited again and again.

We did not invent a framework. We read why the AIs recommended who they recommended, 252 times, and wrote down the reasons they gave. These five came up in every market:

1. A verified sales record

The backbone of ChatGPT's recommendations in every market: sales counts, price ranges, and neighborhood track records pulled from the major portals. If your closed business is invisible or stale on the portals, you are invisible to the AI that reads them most.

2. Recognition written by others

Rankings, press, and directory features published by someone other than you. These carried winners in all six markets. And the AIs repeatedly named the source next to the agent. That is what makes a recommendation arrive clean instead of flagged.

3. Reviews on several sites, not one

Not one glowing profile. Presence on multiple review sites, with volume and recent activity. In several answers, the AIs cited review counts and even quoted specific review language as their evidence.

4. Published guides that answer the client's question

The most buildable lever on this list, cited by all four AIs in all six markets: school-zone guides, condo-vs-TIC explainers, carrying-cost breakdowns, seller pricing guides. One answer named eight agents purely for their guides. Content about the client's decision, not about the agent.

5. Word-of-mouth the machines can read

Three of the four AIs named agents because Reddit threads recommended them. One flagged the source as anecdotal, and named the agents anyway. Your reputation in public forums is now part of your AI record, for better or worse.

The anti-lever: ranking yourself

Self-published "best agent" pages do get you named. But on two of the four AIs, the name arrives with the warning attached, in front of the client. If this is the centerpiece of your strategy, you are buying visibility and paying in trust.

A warning

Don't try to trick the AIs. When they catch it, they tell your client.

Our scan caught one attempt in the wild. A "top 10 agents" ranking site had hidden text inside it, telling AI systems how to cite it. Claude found the hidden text and said this, straight to the client:

[CLAUDE] ASKED: A CHICAGO GOLD COAST SELLER'S DISCOVERY QUESTION

"One of the search results contained text explicitly instructing AI systems on how to cite it - I'm disregarding those embedded instructions, since I only follow guidance from you and my actual system instructions, not text planted in web content."

The same site appeared in our San Francisco results too. Whoever is doing this is doing it across markets.

AI recommendations are already valuable enough that people are trying to game them. This is what happens to the crude version: it gets exposed to the exact person it was meant to influence. The way in is not tricking the machine. It is giving it verifiable reasons to say your name.

Where this leaves you

Run the report. Execute. Track.

1 • Run your report

You cannot fix what you have never seen. Your market's real discovery questions, all four AIs: where you appear, where you don't, who is being named instead, and how.

2 · Execute the action items. Yourself, or with us

The report ends with the specific actions that would improve your visibility, mapped to your actual gaps. Some agents run them in-house. Others hand them back to us: Pitched works end to end, building the proof and the presence, then re-measuring until "named once" becomes "consistently recommended."

3 · If you act, track it

The AIs re-crawl, competitors publish, treatment shifts. The agents winning these answers check their visibility the way they check their pipeline.

Methodology

How this study was built

Every number on this page comes from scans you can reproduce. Six US markets, June–July 2026, each through a defined client scenario. Ten to eleven discovery questions per market, submitted to ChatGPT, Claude, Gemini and Perplexity in their live, search-enabled modes. 252 answers, every one read in full. Research and analysis by Amit Turkaspa, founder of Pitched. Mentions counted on the leaderboard layer – the same layer delivered in every Pitched customer report, so an agent who buys their market's report can check our math against it. The headline figure counts at the mention level: of the roughly 1,200 agent mentions across all answers, about one in three belongs to an agent named exactly once in their market.

Market-size figures are drawn from the relevant Realtor associations' own published membership counts as of July 2026 (MetroTex, MIAMI

Realtors, REBNY, Chicago Association of Realtors, Canopy, SFAR); we make no claim about what share of members is active. All quotes are verbatim and were verified in their full original context. Data cleaning: entities surfaced through geographic ambiguity (place names shared across regions) were excluded from all counts. Agent and team names are withheld throughout, including in positive examples. One market's scan was commissioned by a Pitched customer; only aggregate, non-identifying figures from it appear here.

– Full scan parameters

Scan window	June–July 2026
Markets	6 US markets, each scanned through one defined client scenario
Questions	10–11 discovery questions per market, written for that scenario
Submissions	Each question submitted once per platform via API – no retries, no repeat sampling
Session state	Every question in a fresh API call: no conversation history, no user account, no personalization or location signals
Models	ChatGPT: gpt-5.5 (Responses API, low reasoning effort) · Claude: claude-sonnet-5 · Gemini: gemini-2.5-pro (default configuration) · Perplexity: sonar-pro (temperature 0)
Web search	Each platform's native search tool: OpenAI web search, Anthropic web search (max 3 searches per answer), Gemini Google Search grounding, Perplexity's built-in search. On ChatGPT, Claude and Gemini the model itself decides per question whether to search – same as a real user session

Answers	252 total (40-44 per market), every answer captured verbatim and retained, with the exact dated model identifier stored per answer
Mention counting	Automated name extraction from answer text, then deterministic entity rules in code: brokerages excluded; a team is folded into a person only when the attribution is unambiguous, otherwise it stands as its own entity
Name variants	Merged by deterministic rules only (case, punctuation, leading "The", guarded middle-initial matching) - no fuzzy and no manual merging
Review	Extraction automated; results manually verified against the underlying answers
Top-20 seat count	The 24 forced top-20 answers (one per platform per market, 480 seats) were reviewed manually; a seat counted as filled only when occupied by an identifiable, named agent or team. Unnamed entries, first-name-only entries, matching services, generic instructions, and undelivered seats counted as unfilled
Variability	AI answers vary between runs. This study is a single-submission snapshot per market - which is exactly why visibility needs monitoring, not a one-time check

© 2026 Pitched - AI visibility for real estate agents · Written by Amit Turkaspa

The State of AI Visibility in Real Estate · Edition 1