
AI Visibility Snapshot

Fijara LLC · Portland, Oregon

How Fijara LLC shows up when buyers ask AI — free teaser report

17%

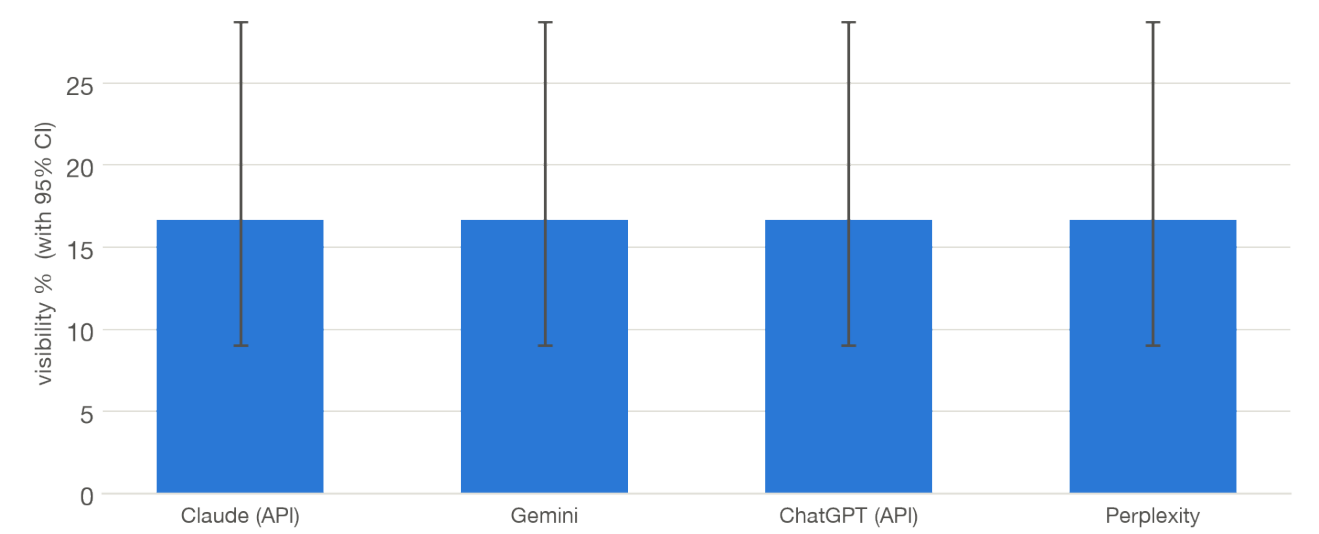
of AI answers to real buyer questions mention Fijara LLC (95% CI 12%–22%, n=216 runs).

The headline

17% VISIBILITY SCORE 95% CI 12%–22%	13% CITATION RATE (YOUR SITE LINKED) 28 of 216 runs	216 AI ANSWERS ANALYSED 3 runs per prompt	232 SOURCE DOMAINS CITED BY ENGINES your roadmap lives here
--	--	--	--

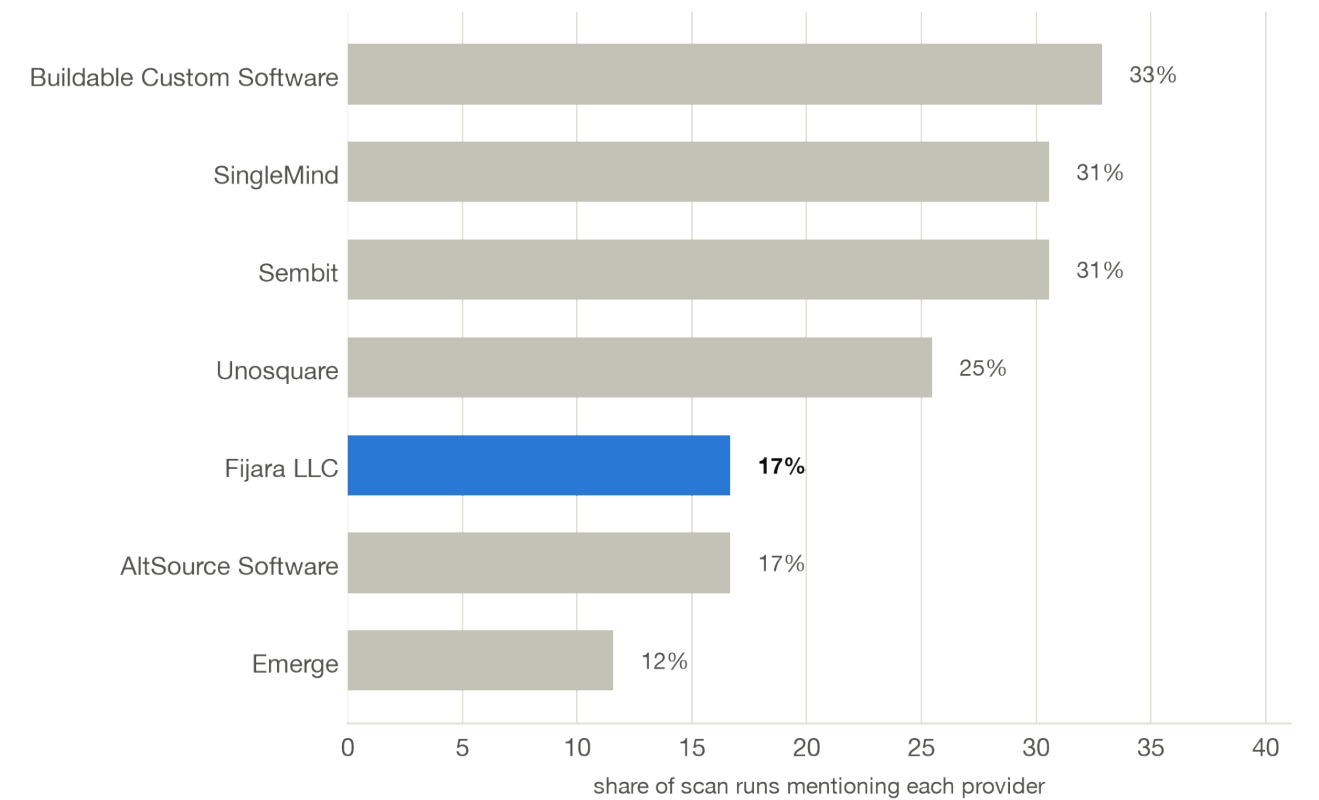
We asked Claude, Google Gemini, ChatGPT, Perplexity 18 genuine buyer questions, 3 times each — 216 AI answers in total. Fijara LLC appeared in 17% of them, typically near the top of the answer.

Per-engine visibility



Error bars are 95% confidence intervals — the honest range given 3 runs per prompt. Where bars differ sharply, buyer experience differs by engine.

Who buyers hear about instead



PROVIDER	MENTIONED IN	SHARE OF VOICE
Buildable Custom Software	33% of runs	20%
SingleMind	31% of runs	19%
Sembit	31% of runs	19%
Unosquare	25% of runs	15%
Fijara LLC (you)	17% of runs	10%
AltSource Software	17% of runs	10%
Emerge	12% of runs	7%

Share of voice = each provider's mentions as a share of all provider mentions across every answer in this scan.

Five buyer questions you're losing right now

These are verbatim prompts a real buyer would type. Frequency is how often Fijara LLC appeared across 3 runs of each.

BUYER QUESTION (VERBATIM PROMPT)	YOUR FREQUENCY	WHO SHOWS UP INSTEAD
“best custom software and app development agency for startups and growing businesses”	0%	—
“top custom software and app development agency companies in Portland, Oregon”	0%	SingleMind, Buildable Custom Software, Unosquare
“best custom software and app development agency in Portland, Oregon”	0%	Buildable Custom Software, Unosquare, Sembit
“who are the most reliable custom software and app development agencies in Portland, Oregon”	0%	Buildable Custom Software, Sembit, SingleMind
“recommended custom software and app development agency for a growing startups and growing busin”	0%	—

The full audit covers all 50 prompts across five intent tiers — category discovery, comparisons, problem-first, qualification, and searches on your own brand name.

Where the engines get their answers

AI engines don't invent recommendations — they retrieve them from a small set of sources: review sites, comparison listicles, directories, forums. These are the domains the engines cited most in this scan:

SOURCE DOMAIN	CITED IN	ENGINES CITING IT
clutch.co	55% of runs	Claude, Gemini, ChatGPT, Perplexity
designrush.com	26% of runs	Claude, Gemini, Perplexity
techreviewer.co	23% of runs	Claude, Gemini, Perplexity

Presence in the sources the engines already trust is the single most direct, evidence-backed route into these answers — not schema markup, not an llms.txt file.

The full picture

The complete audit maps **every** cited source domain (232 in this scan alone), every prompt across all five intent tiers on 5 runs each, competitor share of voice per engine, sentiment of every mention, and a prioritised 90-day plan ranked by effort and impact.

Reply to the email this came with, and it's a 20-minute conversation.

Methodology — read this before the numbers

Why every prompt runs multiple times

The same prompt returns a materially different brand list run-to-run (SparkToro, Jan 2026: the same list appears in fewer than 1 in 100 repeat runs). A single query is noise. Every number in this report is a **frequency across 3 runs per prompt**, with a 95% Wilson confidence interval.

Engines queried

ENGINE	SURFACE	RUNS ANALYSED	STATUS
Claude (Anthropic API)	Anthropic Messages API + web search	54	ok
Google Gemini	Gemini API + Google Search grounding	54	ok
ChatGPT (OpenAI API)	OpenAI Responses API + web search	54	ok
Perplexity	Perplexity sonar (search-grounded)	54	ok

Date of measurement

August 16, 2026 (run ID 20260816-172100-free). AI answers change constantly; treat this as a snapshot, and trends across repeat scans as the signal.

API responses vs. what a human sees

We query each provider's API with search/grounding enabled. Consumer apps (ChatGPT.com, gemini.google.com, etc.) add their own system prompts, retrieval, history and personalisation, so an individual user's answer can differ from what we measure. We state this openly: the API view is the consistent, comparable, repeatable measurement — not a screenshot of one lucky (or unlucky) session.

What we will never claim

No one controls AI answers. This report is descriptive — “you appear in 17% of runs” — never predictive. Anyone promising a specific ranking or visibility outcome is selling something they cannot deliver.