

AI Search Visibility in Enterprise Bot Management

A multi-engine, multi-market benchmark study of how six AI search engines — ChatGPT, Perplexity, Gemini, Microsoft Copilot, Google AI Overview and Google AI Mode — answer buyer-intent questions about enterprise bot management, bot mitigation and automated-traffic defence across the United States, Canada, India and Germany.

Table 1. Study scale by AI engine — responses recorded and sources cited

AI engine	Responses	Cited sources	Markets
ChatGPT	40,000	1,950,000	4
Microsoft Copilot	40,000	155,000	4
Gemini	40,000	166,000	4
Google AI Mode	40,000	766,000	4
Google AI Overview	40,000	469,000	4
Perplexity	40,000	400,000	4
Total	240,000	3,906,000	4

Published: August 2026

Research period: August 2026

Category: Enterprise bot management (bot management, bot mitigation, automated-traffic defence)

Scope: 10,000 buyer-intent queries · 240,000 AI engine responses · 6 AI engines · 4 markets

Evidence base: 3,906,000 cited sources · 84 distinct product entries

Prepared by: GrackerAI · AI Search Visibility Benchmark Series

Abstract

Security teams that buy automated-traffic defence now begin the evaluation inside a generated response, so the vendor set an AI engine returns forms the consideration set before any vendor is contacted. This study measures how six commercial AI search engines answer buyer-intent questions about enterprise bot management. Each engine received 10,000 enterprise buyer-intent queries in each of four markets during a single collection window in August 2026, producing 240,000 responses that contained 3,906,000 cited sources and covered 84 distinct product entries. For every response the study recorded length, cited pages and their source type, page reachability, hedging, recency anchoring, truncation, refusal, and each brand mention with its ordinal position. The engines converge on a core: 10 of 37 brand strings were named by all six engines, and the four most-mentioned — Cloudflare, Akamai, DataDome and Imperva — took 55.8% of all mentions, with 18 strings named by a single engine. Cloudflare was named most often, at 232,000 mentions, and opened the response in 5 of the six engines; DataDome opened ChatGPT, where it averaged position 1.5 against Cloudflare's 3.1. The engines cite largely separate domains, with a mean pairwise overlap of 0.17, and no two engines share a most-cited domain. 47.3% of everything they cite sits on a vendor's own site, but vendor sites are the largest class for only 1 of the six engines. Dead links are rare, at 2.0% of cited pages. Cloudflare led every market; the order below it changed by market within the same four brands.

Contents

1. Introduction

- 1.1 Background
- 1.2 Why this category
- 1.3 Contribution

2. Research Questions

3. Methodology

- 3.1 Study design
- 3.2 Engines under test
- 3.3 Markets
- 3.4 Query construction
- 3.5 Corpus and unit of analysis
- 3.6 Metric definitions
- 3.7 Data collection and processing
- 3.8 Scope exclusions

4. Results

- 4.1 Response characteristics by engine
- 4.2 Brand visibility
- 4.3 Product-level results
- 4.4 Citation volume and source type

- 4.5 Source quality and link decay

- 4.6 Cross-engine source overlap

- 4.7 First-mention position

- 4.8 Geographic variance

- 4.9 Inter-engine disagreement

5. Discussion

6. Threats to Validity

7. Limitations

8. Practical Implications

9. Conclusion

10. Reproducibility and Data Availability

Disclosure

How to cite this report

Appendix A. Brand visibility by AI engine

Appendix B. Most-cited domains and source types

Appendix C. Product entries

Appendix D. Study parameters

Appendix E. Metric computation

1. Introduction

1.1 Background

A security team evaluating bot management now begins inside a generated response. The buyer asks an AI search engine which products separate automated traffic from human traffic at the edge, protect login and checkout flows from credential stuffing and account takeover, and admit the automated agents the business wants while refusing the ones it does not. The engine returns a short ordered list of named vendors with a paragraph of justification for each. That list is not a ranking of the market. It is the output of one system reading one collection of web sources at one moment, and it arrives before the buyer has opened a single vendor site.

The structural consequence is narrow. A vendor absent from the returned list is not ranked low; it is not present. A generated response has no second page, so the buyer has no mechanism equivalent to scrolling further. Where a search results page offered ten positions and a set of related queries, a generated response typically names between three and six vendors and closes. Visibility in this setting is closer to inclusion than to ranking, and the difference between being inside and outside the returned set is discontinuous.

The measurement problem that follows is that no single engine describes the market. Engines differ in the sources they read, the length at which they answer, and the order in which they name vendors. A vendor that measures its position in one engine learns about that engine. This study measures six at once, in four markets, over one collection window, so that agreement and divergence between them can be separated from the behaviour of any one system.

1.2 Why this category

Enterprise bot management is a suitable instrument for this measurement. The vendor set is mature and well documented, so the engines have abundant published material to draw on, and 10 of the 37 brand strings recorded were named by every engine. A category with a settled core isolates engine behaviour from genuine market churn.

The category also has a stable and publicly articulated capability vocabulary. Across the corpus the engines returned the same selection factors: behavioural and machine-learning detection, device fingerprinting, low false-positive rates, real-time scoring, protection for login and checkout flows against credential stuffing and account takeover, integration with the content-delivery network or web application firewall already in place, and selective admission of verified automated agents. These criteria are specific enough that responses can be compared to each other rather than to a general description of security software.

The category also splits cleanly along one design axis the engines return repeatedly: platforms native to a content-delivery network against independent layers that sit in front of any network. That split gives the engines a genuine choice to make when a query asks for a single recommendation, and §4.9 records that they make it differently.

1.3 Contribution

This study provides four things a single-engine measurement does not. It establishes how far six engines agree on the vendor set in one category, separating the shared core from the long tail each engine holds alone. It measures the overlap between the engines' evidence bases directly, as a pairwise coefficient over cited domains, and classifies every cited page by source type, so that agreement on vendors can be set against what the engines actually read. It reports mention volume, first-mention position and mean ordinal position as three separate quantities, which allows

them to be shown moving independently. And it tests the vendor set across four markets in the same window, which distinguishes market-specific results from a set the engines return everywhere. All four rest on a single corpus of 240,000 responses collected under one design.

2. Research Questions

Six questions are addressed, each answerable from the recorded data and each answered explicitly in §5.

1. **RQ1.** Do the six engines converge on the same set of named vendors?
2. **RQ2.** How concentrated is brand visibility across the recorded vendor set?
3. **RQ3.** Do the engines draw on a shared evidence base?
4. **RQ4.** Does ordinal position track mention volume?
5. **RQ5.** Does the vendor set vary between the four markets?
6. **RQ6.** Do the engines cite vendors' own sites or third-party sources?

3. Methodology

3.1 Study design

The study is observational and cross-sectional. Six AI search engines were queried with a fixed query set in four markets during a single collection window in August 2026. No engine was manipulated, no result was solicited, and no intervention was applied between conditions. The design records what the engines returned; it does not test a hypothesis about why they returned it. The collected responses were processed under the instrumentation's current citation, source-type and entity rules.

3.2 Engines under test

Six engines were tested: ChatGPT, Microsoft Copilot, Gemini, Google AI Mode, Google AI Overview and Perplexity. All six are commercial services. None is version-pinned, none exposes a stable model identifier to the caller, and all are subject to continuous change by their operators. A result reported here describes the service as it behaved during the August 2026 window, not a fixed artefact that can be re-queried later in the same state. Three of the six — Gemini, Google AI Mode and Google AI Overview — are operated by the same parent, which is relevant to the overlap results in §4.6.

3.3 Markets

Four markets were tested: United States, Canada, India and Germany. Market variance was tested because engines localise results, and a vendor set returned in one market cannot be assumed to hold in another. The four combine two large English-language markets, one large non-English market, and one market with a substantial domestic buyer base and a separate regulatory environment. Each query was issued in every market, so market is a within-design condition rather than a separate sample.

3.4 Query construction

The query set contains 10,000 distinct buyer-intent queries. It combines core category questions of the form a buyer would ask when assembling a shortlist with paraphrase variants of those questions, so that the stability of a returned vendor set under rewording can be observed rather than assumed. Queries were written to elicit vendor recommendations rather than definitions.

Every query was issued to every engine in every market. The corpus identity follows directly: 10,000 queries in 4 markets gives 40,000 responses per engine, and 10,000 queries across 6 engines and 4 markets gives 240,000 responses in total.

3.5 Corpus and unit of analysis

The unit of analysis is one engine response to one query in one market. Metrics are computed over responses and aggregated to the engine, the brand, the product entry or the market as the metric requires. The corpus contains 240,000 responses, distributed evenly across the six engines at 40,000 responses each and across the four markets at 60,000 responses each. Every per-engine comparison in §4 therefore rests on the same denominator.

3.6 Metric definitions

Every metric used anywhere in this report, including the appendices, is defined here.

Table 2. Metric definitions — all metrics used in this report

Metric	Definition	Unit
Response	One engine output returned for one query in one market.	count
Mean response length	Arithmetic mean word count of an engine's responses.	words
Hedging rate	Share of an engine's responses containing explicit qualification of a recommendation.	% of responses
Recency-anchored rate	Share of an engine's responses that anchor a claim to an explicit calendar date.	% of responses
Truncation rate	Share of an engine's responses ending mid-sentence, inferred from the terminal characters of the response text.	% of responses
Refusal rate	Share of an engine's responses declining to answer the query.	% of responses
First-mention vendor	The brand most often named before any other brand in an engine's responses.	brand
Brand mention	One occurrence of a vendor brand name in a response. Repeat occurrences within one response are counted once. Spelling variants of one brand are merged.	count
Mean ordinal position	Arithmetic mean of the position at which a brand is named within a response, counting named vendors from one.	position
Mean sentiment	Arithmetic mean of the score the analysis model assigned to how favourably the response describes the brand, on a scale from minus one to one, under a fixed rubric.	score

Metric	Definition	Unit
Cited source	One distinct page a response cites, in its source list or as an in-text link, after removal of page assets, trackers and advertising links. Variants of one page differing only in tracking parameters or text fragments count once.	count
Distinct domains	The number of different registrable domains cited by that engine across its responses. Reported per engine and never summed across engines.	count
Source type	The class of the cited page's domain: vendor site (a domain of a vendor named in the response), review aggregator, news, forum, social, marketplace, government, Wikipedia, or other.	class
Dead-link rate	Share of an engine's cited pages that returned not-found, gone or a server error, or whose host could not be reached, when tested. A page that refused the request is not a dead link.	% of cited sources
Self-citation rate	Share of an engine's cited sources hosted on a domain operated by the engine's own parent.	% of cited sources
Jaccard coefficient	Size of the intersection of two engines' cited-domain sets divided by the size of their union. Zero is disjoint, one is identical.	ratio
Product entry	One product recorded against a brand, after plan, edition and deployment spellings of one product are merged.	count
Mean rank	Arithmetic mean of the ordered position of a product entry within the response that named it.	position

3.7 Data collection and processing

Responses were collected through each engine's public interface and normalised to plain text. Word counts, terminal-character inspection for truncation, and detection of qualifying and date-anchoring language were computed over that normalised text.

Cited sources were taken as the union of the source list each engine attached to a response and the pages the response linked to in its text. Page assets, tracking and advertising links were removed, and variants of one page that differ only in tracking parameters or highlighted-text fragments were counted once. Where an engine wrapped a citation in a redirect, the redirect was followed to the cited page. Each page was then reduced to its registrable domain for the distinct-domain, source-type and overlap measures. Source type was assigned by rule from the domain. A domain belonging to a vendor the response names is a vendor site. Forums, social and video sites, news publishers, marketplaces, government domains and review aggregators were matched against lists and host-name patterns, and the remainder is other.

Reachability was tested once per page. A page was recorded as a dead link when it returned not-found, gone or a server error, or when its host could not be reached after a retry. A page that refused the request, as sites behind bot protection do, was recorded as reachable, since the page exists. That rule matters more in this category than in most, because the vendors under study sell the bot protection that refuses such requests.

Brand mentions were detected against the vendor names each response used and recorded with the ordinal position at which the brand appeared among the named vendors, counting from one. Spelling variants of one brand, such as a vendor's name with and without a corporate suffix, were merged, and plan, edition and category-noun variants of one product were merged into one product entry. The first-mention vendor for an engine is the brand that most often held

position one. Sentiment was assigned by the analysis model to each brand and product mention, on a scale from minus one to one. The rubric scores how favourably the response describes the entity and is documented with the instrumentation.

One implementation detail is not documented by the instrumentation: the lexicon used to detect qualifying and date-anchoring language. §6.1 states the consequence for how those two rates should be read. Source age was recorded where a page declared a publication date, but fewer than a fifth of cited pages did, so it is not reported.

3.8 Scope exclusions

The study did not paraphrase, summarise or edit engine output before analysis. It did not check any claim an engine made against an external source, and it did not score responses for accuracy. It did not test, benchmark or evaluate any vendor product. No placement was paid for, solicited or arranged. The object of study is how the engines described the category, not whether their descriptions were correct.

4. Results

4.1 Response characteristics by engine

The six engines construct responses differently. Table 3 gives the response-construction measures for each engine over an identical number of responses.

Table 3. Response construction by AI engine — n = 40,000 responses per engine, 240,000 in total

AI engine	Responses	Mean words	Hedging	Recency-anchored	Truncation	Refusal	First-mention vendor
ChatGPT	40,000	584.0	13%	33%	23%	0%	DataDome
Microsoft Copilot	40,000	502.5	7%	68%	0%	0%	Cloudflare
Gemini	40,000	521.4	0%	0%	45%	0%	Cloudflare
Google AI Mode	40,000	347.9	0%	7%	0%	0%	Cloudflare
Google AI Overview	40,000	224.8	3%	20%	3%	0%	Cloudflare
Perplexity	40,000	105.5	3%	13%	42%	0%	Cloudflare

Mean response length spans a factor of 5.5, from 105.5 words for Perplexity to 584.0 for ChatGPT. Figure 1 shows the distribution. Three engines exceed 500 words — ChatGPT, Microsoft Copilot and Gemini — and Google AI Overview and Perplexity fall below 300.

Truncation was recorded in 4 of the six engines: 45% for Gemini, 42% for Perplexity, 23% for ChatGPT and 3% for Google AI Overview. Microsoft Copilot and Google AI Mode recorded none. Recency anchoring is highest for Microsoft Copilot at 68%, and Gemini recorded none. Hedging is highest for ChatGPT at 13%, followed by Microsoft Copilot at 7%; Gemini and Google AI Mode recorded none. No engine refused a query: the refusal rate is 0% for all six. Five engines record Cloudflare as the first-mention vendor and one records DataDome; §4.7 reports the position data behind that column.

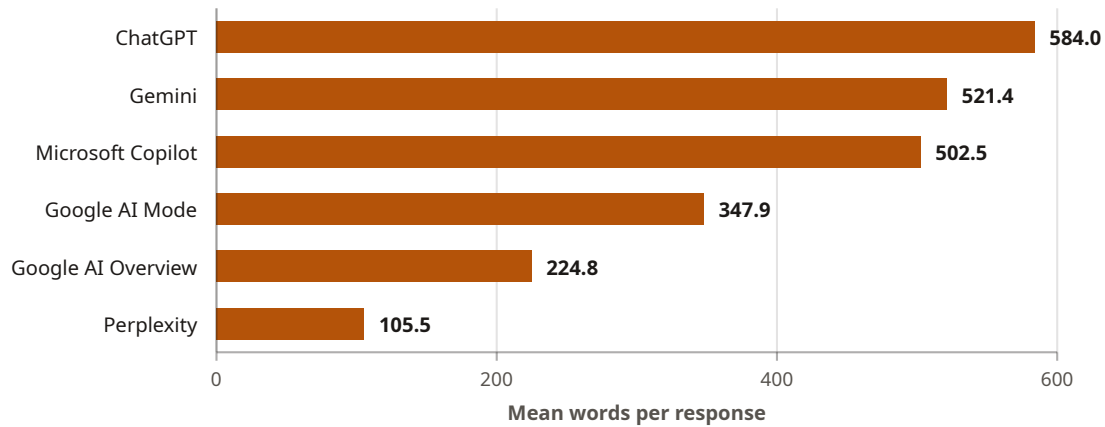


Figure 1. Mean response length by AI engine. Arithmetic mean word count per response; n = 40,000 responses per engine. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

4.2 Brand visibility

Brand visibility was recorded for 37 brand strings, listed in full in Appendix A. Table 4 gives the twelve most-mentioned, with the number of engines that named each and its mean sentiment. Figure 2 shows the same twelve.

Table 4. Brand visibility, twelve most-mentioned brands — of 37 brand strings recorded; mean sentiment where recorded, a dash otherwise; n = 240,000 responses

Brand	Total mentions	Engines recording	Mean sentiment
Cloudflare	232,000	6	0.81
Akamai	210,000	6	0.76
DataDome	203,000	6	0.81
Imperva	196,000	6	0.71
HUMAN Security	137,000	6	0.75
Radware	91,000	6	0.69
Arkose Labs	73,000	6	0.72
F5	65,000	6	0.70
Kasada	63,000	6	0.76
Netacea	45,000	4	0.72
Cequence	38,000	4	0.68
Fastly	33,000	6	0.70

Visibility is concentrated at the top and long-tailed below it. The two most-mentioned brands, Cloudflare and Akamai, account for 29.3% of all mentions recorded across the 37 brand strings, at 232,000 and 210,000 mentions, 22,000 apart. DataDome follows at 203,000 and Imperva at 196,000, 7,000 apart; the four together take 55.8%. The gap from Imperva to HUMAN Security, at 59,000, is the largest between any two adjacent brands in Table 4, and below HUMAN Security at 137,000 no brand exceeds 91,000.

10 of the 37 brand strings were named by all six engines: Cloudflare, Akamai, DataDome, Imperva, HUMAN Security, Radware, Arkose Labs, F5, Kasada and Fastly. 18 were named by exactly one engine, none of them above 11,000 mentions. The strings recorded include adjacent categories — fraud prevention, challenge and verification services, and security consultancies — each named in a small number of responses as the engines wrote them. Mean sentiment for the twelve brands in Table 4 occupies a band from 0.68 for Cequence to 0.81 for Cloudflare, a spread of 0.13.

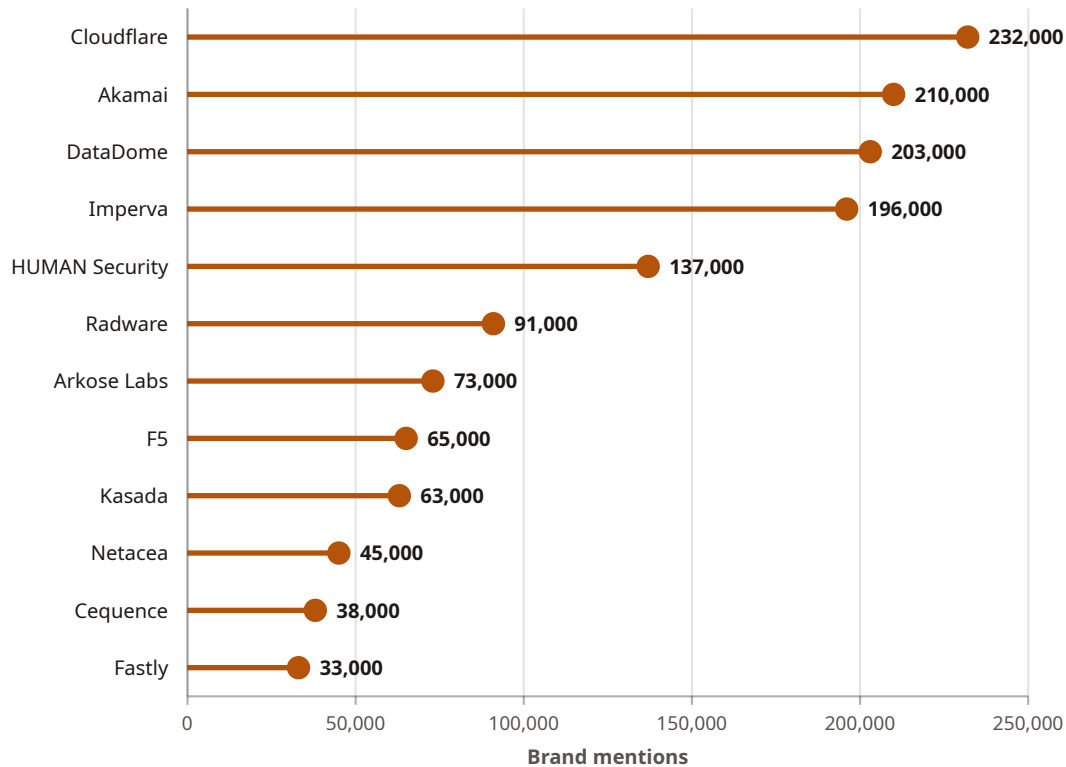


Figure 2. Brand mentions, twelve most-mentioned brands. The 12 most-mentioned of 37 brand strings recorded; n = 240,000 responses. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

4.3 Product-level results

84 distinct product entries were recorded after plan, edition and category-noun spellings of one product were merged. An entry is one product as an engine named it, so a brand with several products holds several entries: Akamai holds 8, HUMAN Security 8, and Cloudflare and Imperva 6 each. HUMAN Security is attributed under three spellings of its name in Appendix C and holds 14 entries across them. Table 5 gives the ten most-mentioned entries.

Table 5. Product entries, ten most-mentioned — of 84 entries; n = 240,000 responses

Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
Cloudflare Bot Management	Cloudflare	191,000	1.9	0.81	6	4
Akamai Bot Manager	Akamai	169,000	2.8	0.76	6	4
DataDome	DataDome	128,000	2.8	0.81	6	4
Imperva Advanced Bot Protection	Imperva	120,000	3.7	0.72	6	4
Radware Bot Manager	Radware	77,000	4.3	0.70	6	4

Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
F5 Distributed Cloud Bot Defense	F5	55,000	5.1	0.71	6	4
DataDome Bot Protect	DataDome	41,000	2.9	0.81	5	4
Kasada	Kasada	39,000	4.2	0.76	5	4
HUMAN Security	HUMAN Security	38,000	3.5	0.74	6	4
Arkose Labs	Arkose Labs	36,000	5.2	0.72	6	4

Product entries carry their own brand attributions, and Appendix C lists entries attributed to brands outside the twelve in Table 4. The distribution is long-tailed. 66 of the 84 entries record fewer than 10,000 mentions, 31 record exactly 1,000, and 43 were named by exactly one engine. 9 entries were named by all six engines, 21 were recorded in all four markets, and 35 were recorded in a single market. 7 entries carry no mean rank because the engine that named them did not place them in an ordered list. The full table is in Appendix C.

Mean rank among the ten most-mentioned entries runs from 1.9 to 5.2. The leading entry, Cloudflare Bot Management, records 191,000 mentions at a mean rank of 1.9, the earliest of the ten; the second, Akamai Bot Manager, records 169,000 at 2.8.

4.4 Citation volume and source type

The corpus contains 3,906,000 cited sources, distributed very unevenly across the six engines, as Figure 3 shows. Table 6 gives volume and domain breadth; Table 7 gives the source-type mix.

Table 6. Cited sources and domain breadth by AI engine — n = 3,906,000 cited sources

AI engine	Cited sources	Share of corpus	Distinct domains	Most-cited domain
ChatGPT	1,950,000	49.9%	118	datadome.co (211,000)
Microsoft Copilot	155,000	4.0%	24	worldmetrics.org (44,000)
Gemini	166,000	4.2%	35	radware.com (19,000)
Google AI Mode	766,000	19.6%	165	gartner.com (67,000)
Google AI Overview	469,000	12.0%	69	cequence.ai (77,000)
Perplexity	400,000	10.2%	31	indusface.com (32,000)
Total	3,906,000	—	—	—

ChatGPT alone accounts for 49.9% of all cited sources in the corpus, at 1,950,000, and Google AI Mode for a further 19.6%. Microsoft Copilot accounts for the least at 4.0%. Google AI Mode cites from the widest domain population at 165 distinct domains, followed by ChatGPT at 118; Microsoft Copilot cites from the narrowest at 24. Distinct-domain counts are reported per engine and are not summed across engines.

No two engines share a most-cited domain. ChatGPT's is datadome.co, a vendor's site; Google AI Overview's is cequence.ai and Gemini's is radware.com, both vendors' sites; Google AI Mode's is gartner.com, an analyst firm; Microsoft Copilot's is worldmetrics.org, a statistics aggregator. Appendix B gives the three most-cited domains for each engine.

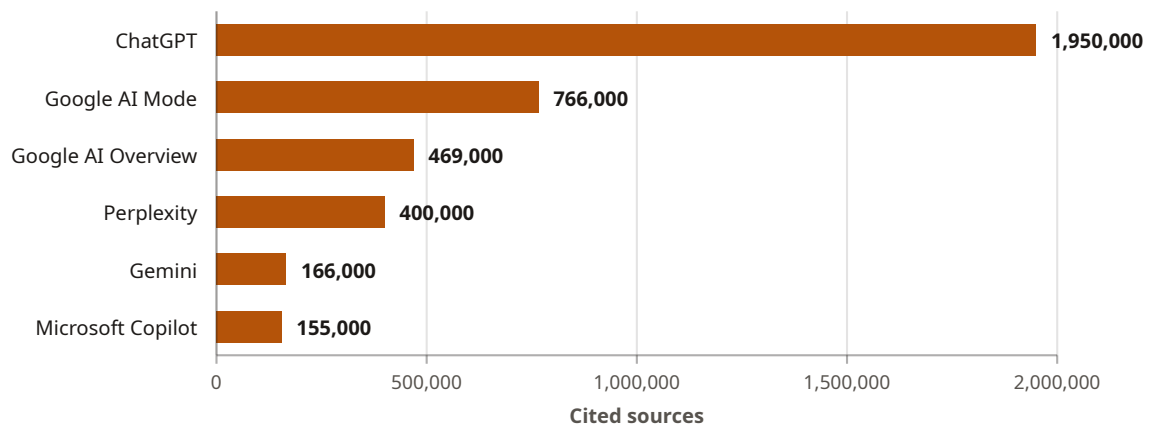


Figure 3. Cited sources by AI engine. Sources cited per engine; n = 3,906,000 cited sources. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

Table 7. Source type by AI engine — share of each engine's cited sources by the class of the cited domain; n = 3,906,000 cited sources

AI engine	Vendor site	Other	Review aggregator	Social	Forum	Marketplace	News	Government	Wikipedia
ChatGPT	65.2%	29.7%	3.8%	0.1%	0.6%	0.3%	0.1%	0.1%	0.1%
Microsoft Copilot	7.7%	88.4%	1.3%	0.0%	0.0%	1.9%	0.6%	0.0%	0.0%
Gemini	42.2%	44.0%	13.3%	0.0%	0.6%	0.0%	0.0%	0.0%	0.0%
Google AI Mode	31.9%	49.9%	12.1%	3.4%	2.1%	0.0%	0.7%	0.0%	0.0%
Google AI Overview	42.2%	50.1%	4.5%	2.6%	0.4%	0.0%	0.2%	0.0%	0.0%
Perplexity	13.5%	76.5%	10.0%	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%
All engines	47.3%	43.9%	6.5%	1.0%	0.8%	0.2%	0.2%	0.1%	0.0%

Across the corpus, 47.3% of cited sources sit on a vendor site, 1,848,000 of 3,906,000, and 43.9% fall outside every named class. Review aggregators, which here include analyst firms, account for 6.5%. The mix differs sharply by engine. Vendor sites are the largest class for 1 of the six engines, ChatGPT, which draws 65.2% of its citations from them; for the other 5 the largest class is other. Gemini and Google AI Overview each draw 42.2% from vendor sites, and Microsoft Copilot draws 7.7%, with 88.4% of its citations outside every named class. Review aggregators reach 13.3% of Gemini's citations, the highest share. No news, academic or government domain was recorded by any engine beyond a handful of citations. Figure 4 shows the vendor-site share by engine, and Appendix B gives the counts behind Table 7.

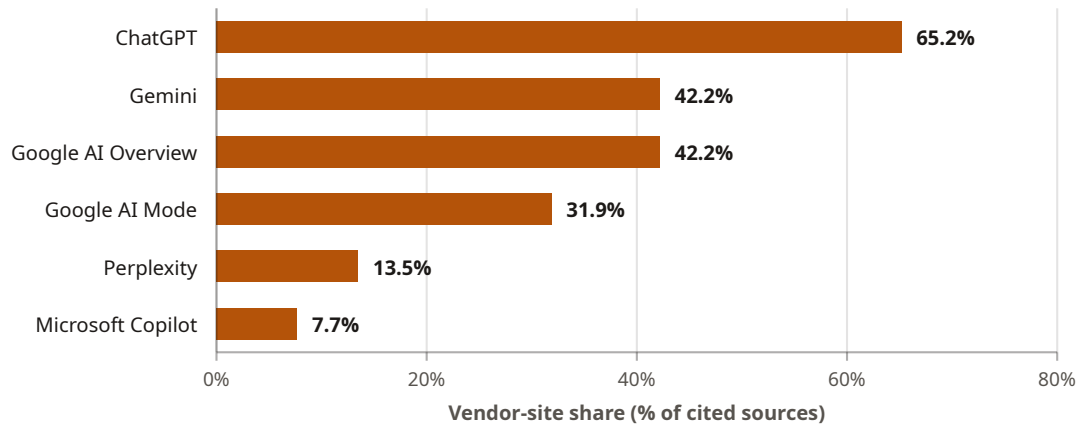


Figure 4. Share of cited sources on vendor sites by AI engine. Cited sources whose domain belongs to a vendor named in the response, as a share of that engine's cited sources; n = 3,906,000 cited sources. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

4.5 Source quality and link decay

A small share of the pages the engines cited did not resolve when tested. Table 8 gives the rate for each engine.

Table 8. Dead links and self-citation by AI engine — n = 3,906,000 cited sources

AI engine	Cited sources	Dead-link rate	Dead links	Self-citation rate
ChatGPT	1,950,000	2.3%	44,800	0%
Microsoft Copilot	155,000	0%	0	0%
Gemini	166,000	0%	0	0%
Google AI Mode	766,000	1%	7,660	0%
Google AI Overview	469,000	3%	14,100	0%
Perplexity	400,000	2.5%	10,000	0%
Total	3,906,000	2.0%	76,560	—

The overall dead-link rate is 2.0%, 76,560 of 3,906,000 cited pages. Two engines — Microsoft Copilot and Gemini — record no dead links at all, and no engine exceeds 3%, the rate for Google AI Overview. ChatGPT accounts for 44,800 of the 76,560 dead links, 58.5% of the total, in proportion to its share of citations.

The self-citation rate is 0% for all six engines.

4.6 Cross-engine source overlap

The Jaccard coefficient over each pair of engines' cited-domain sets measures how far the engines read the same web. Table 9 gives the full matrix.

Table 9. Cross-engine cited-domain overlap — Jaccard coefficient over cited-domain sets; 0 is disjoint, 1 is identical; n = 3,906,000 cited sources

	ChatGPT	Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
ChatGPT	1.00	0.10	0.19	0.18	0.22	0.17

	ChatGPT	Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
Copilot	0.10	1.00	0.07	0.03	0.08	0.25
Gemini	0.19	0.07	1.00	0.16	0.30	0.32
Google AI Mode	0.18	0.03	0.16	1.00	0.24	0.11
Google AI Overview	0.22	0.08	0.30	0.24	1.00	0.20
Perplexity	0.17	0.25	0.32	0.11	0.20	1.00

Shading increases with the coefficient: < 0.10 0.10–0.17 0.18–0.25 0.26–0.33 ≥ 0.34

The mean coefficient across the 15 distinct engine pairs is 0.17. The highest value is 0.32, between Gemini and Perplexity, two engines operated by different parents. The lowest is 0.03, for Microsoft Copilot with Google AI Mode. The three engines operated by the same parent — Gemini, Google AI Mode and Google AI Overview — record coefficients of 0.16, 0.30 and 0.24 with each other; the same parent does not produce the highest overlap in this category.

Six of the 15 pairs reach 0.20. Microsoft Copilot records the lowest coefficient with four of the five other engines, from 0.03 with Google AI Mode to 0.10 with ChatGPT, and its one elevated value is 0.25 with Perplexity. Perplexity in turn records its highest value, 0.32, with Gemini. The two engines that cite the least and the two that cite the most do not pair with each other.

4.7 First-mention position

Mention volume records how often a brand appears. Mean ordinal position records how early it appears among the vendors a response names. The first-mention vendor is the brand that most often opens a response. Table 10 gives all three for the three most-mentioned brands, and Figure 5 shows Cloudflare and DataDome.

Table 10. Mean ordinal position of the three leading vendors — lower is earlier in the response; n = 240,000 responses

AI engine	Cloudflare	Akamai	DataDome	First-mention vendor
ChatGPT	3.1	3.3	1.5	DataDome
Microsoft Copilot	1.2	2.6	3.5	Cloudflare
Gemini	2.5	2.7	2.9	Cloudflare
Google AI Mode	1.7	2.5	3.5	Cloudflare
Google AI Overview	1.7	2.9	2.4	Cloudflare
Perplexity	1.0	3.0	4.3	Cloudflare

Cloudflare holds the first-mention position in 5 of the 6 engines, on 232,000 mentions, the highest total in the corpus. DataDome holds it in 1, ChatGPT, on 203,000, the third-highest total. Akamai, second on mentions at 210,000, opens no engine.

Mean ordinal position follows the first-mention split more closely here than in other categories in this series, but not completely. Cloudflare records the earlier mean position than Akamai in all 6 engines, from 1.0 in Perplexity to 3.1 in ChatGPT. DataDome records the earlier position than Cloudflare only in ChatGPT, where it averages 1.5 against

Cloudflare's 3.1, the widest gap in Table 10; in Perplexity it averages 4.3, the latest position in the table. DataDome also records the earlier position than Akamai in ChatGPT and Google AI Overview, so the brand that opens one engine sits third or later in most of the others.

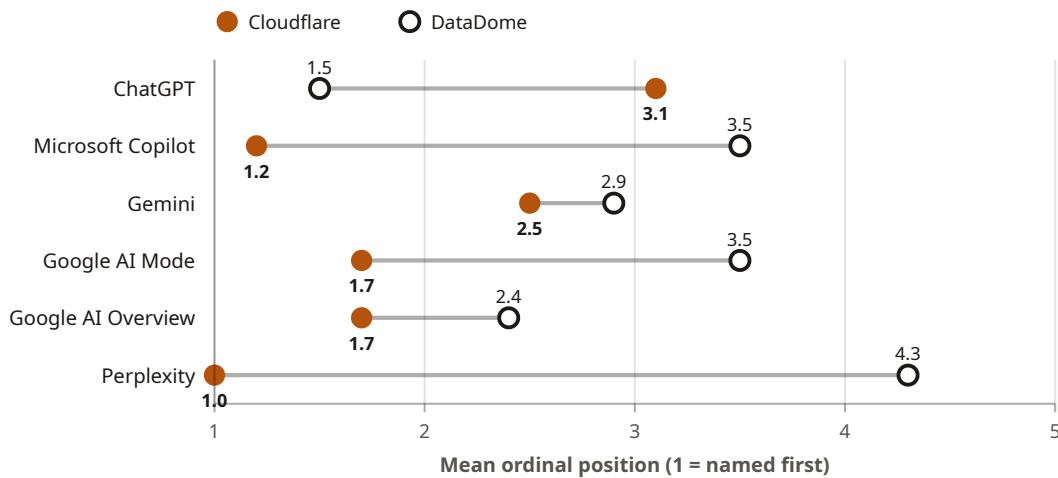


Figure 5. Mean ordinal position of Cloudflare and DataDome by AI engine. Lower is earlier in the response; n = 240,000 responses. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

4.8 Geographic variance

Table 11 gives the three leading brands in each market, ranked by mentions within that market.

Table 11. Leading brands by market — ranked by brand mentions within the market; n = 60,000 responses per market

Market	First	Second	Third
United States	Cloudflare	DataDome	Akamai
Canada	Cloudflare	DataDome	Akamai
India	Cloudflare	Akamai	Imperva
Germany	Cloudflare	Akamai	DataDome

Cloudflare is first in all four markets. The second position differs by market: DataDome in the United States and Canada and Akamai in India and Germany. The third position also differs: Akamai in the United States and Canada, Imperva in India and DataDome in Germany. The same four brands — Cloudflare, Akamai, DataDome and Imperva — occupy the four leading positions in every market, and no market returns a vendor absent from the other three in those positions. Mean sentiment by market runs from 0.72 in Canada to 0.76 in the United States.

At product level, 21 of the 84 product entries were recorded in all four markets, and every entry in Table 5 was recorded in all four.

4.9 Inter-engine disagreement

The engines diverge on three specific points in the category.

The first concerns the choice between a platform native to a content-delivery network and an independent layer. ChatGPT repeatedly ranks DataDome first, citing deployment across several delivery networks. Gemini more often leads with Cloudflare or Akamai, citing detection at the network edge. The engines therefore return different

products for the same question, and the difference matches the first-mention split in §4.7.

The second concerns claims that appear in one engine and nowhere else. ChatGPT alone states that DataDome is the only vendor supporting cryptographic authentication of automated agents and trust scoring for verified AI crawlers. Microsoft Copilot alone places Forter and Kroll in its account-takeover rankings. The study records both as observed and does not assess either.

The third concerns vendors that only one engine associates with the category. Microsoft Copilot alone returns 12 product entries, among them Forter and AWS WAF. Perplexity alone returns 7, among them Imperva Bot Defense and Fastly Bot Defense. Google AI Mode alone returns 9, Gemini 7, ChatGPT 6 and Google AI Overview 2. The divergences occur on use-case recommendations and at the edge of the vendor set rather than on its core.

5. Discussion

5.1 RQ1 — Do the six engines converge on the same set of named vendors?

They converge on a core larger than in other categories in this series, and diverge in a long tail. §4.2 shows 10 of the 37 recorded brand strings named by all six engines, and §4.3 shows 9 of 84 product entries reaching the same coverage. The difference between 10 brands and 9 products is the answer: the engines agree about vendors and disagree about which of a vendor's products to name, and the vendors in this category ship several.

The tail is where the engines part. 18 of the 37 brand strings and 43 of the 84 product entries were named by exactly one engine, and §4.9 shows each engine holding a tail the others do not. A plausible explanation is that the core set is reachable from almost any part of the category's web presence. The tail depends on which sources a particular engine happens to read and on where that engine draws the category boundary. §4.6 is consistent with that: the engines reach the same core from evidence bases that overlap by a mean Jaccard coefficient of 0.17.

5.2 RQ2 — How concentrated is brand visibility across the recorded vendor set?

Concentrated into four brands at the top and fragmented below them. §4.2 shows the two leading brands taking 29.3% of all mentions and the four leading brands 55.8%, against 37 brand strings recorded, with a gap of 59,000 mentions between the fourth brand and the fifth. §4.3 shows 66 of 84 product entries below 10,000 mentions and 31 at exactly 1,000.

The shape is a core of four vendors the engines reuse, a second tier of six the engines still all name, then a tail of dozens each engine names once or twice. The data is consistent with a threshold effect rather than a gradient: a vendor is either inside the group the engines reuse across queries and markets, or it is in a tail where mention counts fall away quickly. The measurement does not establish what places a vendor on one side of that boundary.

5.3 RQ3 — Do the engines draw on a shared evidence base?

No. §4.6 gives a mean pairwise Jaccard coefficient of 0.17 over cited domains, with a maximum of 0.32 and a minimum of 0.03. Only 6 of the 15 pairs reach 0.20, and Microsoft Copilot records the lowest coefficient with four of the five other engines. §4.4 adds that no two engines share a most-cited domain. Six engines returned substantially the same core vendors while reading substantially different sets of domains.

Volume compounds the divergence. ChatGPT contributes 49.9% of all cited sources on its own, and with Google AI Mode the two reach 69.5%, so a source list that serves either describes little of what the other four read. A practical consequence follows for measurement rather than for the vendors: a visibility programme built against one engine's source list transfers poorly to the others.

5.4 RQ4 — Does ordinal position track mention volume?

Mostly, with one instructive exception. §4.7 shows Cloudflare recording the highest mention total, the first-mention position in 5 engines, and the earlier mean position than Akamai in all six. The exception is ChatGPT, where DataDome, third on mentions across the corpus, opens the response and averages 1.5 while Cloudflare averages 3.1. Akamai, second on mentions, opens no engine and never records the earliest position of the three.

The three measures answer different questions. Mention volume records how often a vendor enters the response; first-mention position records which vendor opens it; mean ordinal position records where a vendor lands on average once named. In this category the first two agree in five engines and disagree in one, and the third shows DataDome placed third or later in 5 engines despite opening the sixth. §6.1 sets out why ordinal position should be read as a proxy for prominence rather than as a measure of it. A visibility measurement that reports only one of the three is incomplete.

5.5 RQ5 — Does the vendor set vary between the four markets?

The set does not; the order below the leader does. §4.8 shows Cloudflare first in all four markets and the same four brands in the four leading positions everywhere, with the second position held by DataDome in the United States and Canada and by Akamai in India and Germany. 21 of 84 product entries were recorded in all four markets, and every entry in Table 5 was.

The result is a negative one for membership and is useful as such. It removes market localisation as an explanation for a vendor's absence from the returned set: a vendor missing in one market is, in this corpus, missing in all of them. The study covered four markets in one window and cannot establish that the same holds for markets or languages it did not test.

5.6 RQ6 — Do the engines cite vendors' own sites or third-party sources?

Split, and split by engine. §4.4 shows 47.3% of all cited sources on a vendor site, but that figure is carried by ChatGPT, which supplies 49.9% of the corpus and draws 65.2% of its own citations from vendor domains. For the other five engines the largest class is other, and Microsoft Copilot draws only 7.7% from vendor sites. Analyst and review sites are a smaller class than in other categories in this series, at 6.5% of the corpus and 13.3% of Gemini's citations.

The data is consistent with two distinct evidence strategies behind one vendor list. One engine reads the vendors' own documentation and product pages in quantity; the others read a broader and largely unclassified web, of which vendor pages are a minority. That the six arrive at the same core set from such different mixes is the strongest indication in the study that the core set is not an artefact of any one kind of source. What the measurement does not establish is whether a vendor's own site earns citations because the vendor is already in the set, or the reverse.

5.7 What these results do not resolve

The study measures outputs, not mechanisms. It cannot say why a vendor is inside the core set or what would move one into it, because it observes no vendor changing position. It records that the engines cite largely separate domains but not whether a citation influenced the vendor list, since citation and recommendation were recorded from the same response and cannot be ordered causally. It classifies cited domains but does not read the cited pages, so it cannot say what on a vendor site was cited.

Two further limits apply to specific metrics. Truncation is inferred from the terminal characters of a response and cannot distinguish a response that ended mid-sentence from one whose final sentence was recorded without terminal punctuation; the 45% for Gemini and 42% for Perplexity should be read with that in mind. Sentiment is a model's reading of the response under a rubric, and the 0.13 spread across the twelve leading brands is too narrow to support a distinction between any two of them.

6. Threats to Validity

6.1 Construct validity

Mean ordinal position is used as a proxy for prominence. It records where a brand name falls in the sequence of vendors a response names, which is not the same as how prominently a reader perceives it. A vendor named first in a list and then qualified, and a vendor named third and then recommended, receive positions that do not reflect that difference. The metric is defensible for comparison between engines over the same corpus and should not be read as a measure of persuasive weight. The first-mention vendor is a mode, not a mean.

Sentiment is assigned by the analysis model under a documented rubric (§3.7). The values in §4.2 occupy a narrow band, which is consistent either with the engines describing these vendors in similar terms or with the rubric having limited resolution over this kind of text. The study cannot distinguish the two, and no claim in this report rests on a difference between two sentiment values.

Source type is assigned from the domain, not from the page. A vendor's blog and its product documentation both count as vendor site; an analyst firm's page counts as a review aggregator. The classes are coarse by design, and the residual other class is the largest class for five of the six engines in this category. Truncation is an inference from the terminal characters of a response, not an observation of a process. Hedging and recency anchoring rest on an undocumented lexicon (§3.7) and are subject to the same limit. The same detectors were applied to all six engines, so the values are comparable within this corpus. They are not comparable to figures produced by another detector.

6.2 Internal validity

All responses were collected in one window in August 2026. The vendor sets and the cited domains describe that window; reachability describes the date on which the pages were tested. A collection window that included a major vendor acquisition or an analyst publication would be expected to produce different results, and the study has no second window against which to test that.

The engines are non-deterministic. The same query issued twice to the same engine can return different text, a different vendor ordering, and a different citation list. The design mitigates this through corpus size and through paraphrase variants within the query set, which spread the measurement over many samples of each engine's behaviour rather than relying on single responses. It does not eliminate it.

Personalisation and session state were not controlled beyond issuing each query in a defined market. Where an engine adapts to prior interaction, that adaptation is not observable from the response text, and its contribution to the results cannot be separated. No engine was version-pinned, because none of the six exposes a version identifier to the caller (§3.2), so a change made by an operator during the window would be indistinguishable from variation in the responses.

6.3 External validity

The study covers one category, four markets and one collection window. Enterprise bot management has a mature core and a long tail (§1.2, §4.2). The core is what makes it a usable instrument; the tail is what limits generalisation, since a category with a different vendor population would not be expected to produce the same shape.

Query sampling is purposive, not random. The query set was written to elicit buyer-intent vendor recommendations, so it samples a deliberately chosen region of the space of things a buyer might ask rather than a representative draw from it. Combined with the non-determinism in §6.2 and the absence of version pinning, this precludes statistical inference from these results. All figures in this report are descriptive statistics over the recorded corpus, and no significance claim, interval or test is reported anywhere in it.

6.4 Reliability

A repeat run in a later window would be expected to reproduce the structural results. Those are the concentration of visibility into a core of four with a long tail, the leading brand opening most engines while one engine opens with a different vendor, the low overlap between the engines' evidence bases, the dependence of the vendor-site share on a single engine, and the stability of the leading group across markets. All rest on large differences that ordinary variation would not close.

Point values would not be expected to reproduce. Individual mention counts, per-engine cited-source totals, the composition of the tail, and the specific most-cited domain for an engine depend on what each engine read during the window. Dead-link rates would be expected to move, because reachability was tested once (§3.7) and the reachability of the cited web changes independently of the engines.

7. Limitations

- One category. The results describe enterprise bot management and do not transfer to other security categories without measurement.
- Four markets, all tested in one window in August 2026. Markets and languages outside the four were not tested.
- No claim an engine made was adjudicated. Where two engines returned incompatible statements (§4.9), the study records the disagreement and does not determine which is correct.
- Source type was assigned from the domain and the residual other class is the largest class for five engines; the study says which kinds of site the engines cite, not what on those sites was cited.
- Source age is not reported, because fewer than a fifth of cited pages declared a publication date.
- One brand, HUMAN Security, is attributed under three spellings in the product table (Appendix C); its brand-level total in Table 4 merges them, its product entries do not.

8. Practical Implications

This section is derived from the results rather than measured by them.

Measure who opens the response, not only who is in it. §4.7 shows the most-mentioned brand opening five engines and the third-placed brand opening the sixth, where it also averages the earliest position. Track mention volume, first-mention share and mean position as three numbers, and decide from them whether the problem is inclusion, opening position or average placement. The three do not respond to the same work.

Treat the vendor's own site as the primary citation surface for ChatGPT and as one surface among several for the rest. §4.4 shows vendor sites as the largest class for one engine only, and the other class as the largest for five. In this category the pages the engines read outside vendor sites — statistics aggregators, comparison sites and the long tail of unclassified pages — carry most of the citations for five of the six engines.

Instrument every engine separately, and at least three from different vendor families. §4.6 gives a mean pairwise cited-domain overlap of 0.17, only 6 pairs at or above 0.20, and no two engines sharing a most-cited domain. A source list that serves ChatGPT describes little of what Microsoft Copilot or Perplexity read.

Name products consistently, and expect several entries per vendor regardless. §4.3 records 84 product entries, with Akamai holding 8 and HUMAN Security 14 across three spellings of its own name. Some of that is real product breadth and some is naming drift; only the vendor can tell which, and only consistent naming across published material lets a measurement consolidate it.

Do not build a market-by-market visibility programme for these four markets. §4.8 shows the same brand leading all four and the same four brands in the leading positions everywhere, with only the order below the leader changing. Effort that would go into market-specific work is better spent on the inclusion problem, which is the same problem in every market tested.

Do not treat link decay as a lever in this category. §4.5 records an overall dead-link rate of 2.0% and no engine above 3%. Do not treat sentiment as one either: the twelve leading brands sit within a 0.13 band (§4.2).

9. Conclusion

This study put 10,000 enterprise buyer-intent queries to six AI search engines in four markets, producing 240,000 responses containing 3,906,000 cited sources and covering 84 distinct product entries in enterprise bot management.

The engines agree about a core and disagree about evidence. 10 of 37 brand strings were named by all six engines, led by Cloudflare at 232,000 and Akamai at 210,000, with 18 strings named by one engine only. Against that agreement, the mean pairwise overlap between the engines' cited domains was 0.17, the highest value anywhere in the matrix was 0.32, and no two engines shared a most-cited domain. Vendor sites carry 47.3% of everything the engines cite, but that share rests on one engine; for the other five, most citations fall outside every class the study tracks.

Two further results qualify that list. Cloudflare opens 5 of the six engines and is read earlier than Akamai in all six, yet DataDome opens ChatGPT, where it averages 1.5 against Cloudflare's 3.1. Source reliability is high: 2.0% of cited pages were unreachable, and no engine exceeded 3%. Cloudflare led every market, and the same four brands held the leading positions in all of them.

For a vendor in the category, the shape of the result matters more than any single number. Visibility here is concentrated into a core the engines reuse across markets, phrasings and engines, reached through six largely separate evidence bases. Entering that core is a different problem from opening the response once inside it, and both have to be solved once per engine.

10. Reproducibility and Data Availability

Responses were collected through each engine's public interface using the study's own instrumentation and retained as normalised plain text. Retained alongside them are the extracted citation lists with source type and reachability results, and the brand and product mention records with their ordinal positions. The dataset for this report was generated in August 2026 and processed under the instrumentation's current citation and entity rules.

Reproducing the study would require the query set, the market configuration, the vendor name list used for mention detection and the reachability test, applied to the same six engines within a comparable collection window. Exact reproduction is not achievable, because the engines are commercial services under continuous change and none is version-pinned; §6.4 states what a repeat run would and would not be expected to reproduce.

Requests for the dataset supporting the tables in this report may be directed to GrackerAI.

Disclosure

GrackerAI produces this research and also sells AI search visibility measurement for the category it covers. The study measured how six AI search engines describe vendors in enterprise bot management. It did not test, evaluate or verify any vendor's product, and no statement in this report is an assessment of product quality. No vendor named in this report paid for, requested, reviewed or was given advance sight of any placement, ordering or finding in it.

How to cite this report

GrackerAI. (August 2026). *AI Search Visibility in Enterprise Bot Management: A Multi-Engine, Multi-Market Benchmark Study*. GrackerAI AI Search Visibility Benchmark Series. <https://gracker.ai/data-and-research-reports>

Appendix A. Brand visibility by AI engine

Table A1. Brand mentions and mean ordinal position by AI engine — all 37 brand strings recorded; mean ordinal position in parentheses, a dash where the engine recorded no mention or no position; n = 240,000 responses

Brand	Total	ChatGPT	Microsoft Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
Cloudflare	232,000	38,000 (3.1)	40,000 (1.2)	41,000 (2.5)	38,000 (1.7)	37,000 (1.7)	38,000 (1.0)
Akamai	210,000	38,000 (3.3)	33,000 (2.6)	40,000 (2.7)	38,000 (2.5)	30,000 (2.9)	31,000 (3.0)
DataDome	203,000	39,000 (1.5)	35,000 (3.5)	32,000 (2.9)	30,000 (3.5)	34,000 (2.4)	33,000 (4.3)

Brand	Total	ChatGPT	Microsoft Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
Imperva	196,000	30,000 (4.9)	36,000 (3.5)	32,000 (4.3)	35,000 (3.4)	33,000 (3.2)	30,000 (2.9)
HUMAN Security	137,000	37,000 (3.0)	14,000 (5.7)	25,000 (3.6)	26,000 (4.0)	21,000 (3.3)	14,000 (4.7)
Radware	91,000	6,000 (5.3)	11,000 (5.3)	21,000 (3.6)	16,000 (5.5)	11,000 (4.1)	26,000 (3.5)
Arkose Labs	73,000	17,000 (4.2)	14,000 (6.2)	4,000 (5.0)	8,000 (5.2)	15,000 (4.2)	15,000 (6.7)
F5	65,000	10,000 (4.3)	4,000 (4.7)	11,000 (5.1)	22,000 (5.1)	14,000 (5.1)	4,000 (8.0)
Kasada	63,000	20,000 (3.3)	11,000 (5.0)	15,000 (4.3)	9,000 (5.0)	3,000 (5.0)	5,000
Netacea	45,000	8,000 (5.8)	27,000 (5.1)	—	—	1,000 (6.0)	9,000 (5.7)
Cequence	38,000	—	—	7,000 (5.0)	4,000 (5.5)	15,000 (5.3)	12,000 (3.0)
Fastly	33,000	2,000 (6.0)	2,000 (4.0)	6,000 (4.0)	5,000 (4.4)	4,000 (4.7)	14,000 (3.0)
CHEQ	22,000	2,000	6,000	—	—	—	14,000 (7.3)
AWS	18,000	—	10,000 (4.0)	3,000 (5.0)	2,000 (8.0)	—	3,000
Kroll	11,000	—	11,000 (5.3)	—	—	—	—
Google	8,000	2,000	3,000 (4.0)	3,000 (4.0)	—	—	—
Sift	7,000	—	5,000 (7.0)	—	1,000 (8.0)	1,000 (5.0)	—
Forter	6,000	—	6,000 (2.0)	—	—	—	—
Indusface	6,000	—	—	6,000 (6.0)	—	—	—
GeeTest	5,000	—	1,000 (6.0)	2,000 (7.0)	2,000 (7.5)	—	—
Booz Allen Hamilton	4,000	—	4,000 (3.0)	—	—	—	—
Fingerprint	4,000	—	—	—	3,000 (8.0)	1,000 (6.0)	—
Kount	4,000	—	4,000	—	—	—	—
NCC Group	4,000	—	4,000 (2.0)	—	—	—	—
Barracuda	3,000	—	—	—	—	—	3,000
Fortinet	3,000	—	—	—	—	—	3,000 (5.0)
Nardello & Co. / Kroll	3,000	—	3,000	—	—	—	—
Clearout	2,000	—	—	1,000	1,000 (7.0)	—	—
Thales	2,000	—	—	2,000	—	—	—
Azion	1,000	—	—	—	—	—	1,000
DataVisor	1,000	—	1,000	—	—	—	—
hCaptcha	1,000	1,000	—	—	—	—	—

Brand	Total	ChatGPT	Microsoft Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
IO River	1,000	—	—	1,000	—	—	—
NodeStack	1,000	—	—	1,000	—	—	—
Trusted Accounts	1,000	—	—	1,000	—	—	—
Webz.io	1,000	—	1,000	—	—	—	—
White Ops	1,000	—	1,000	—	—	—	—
Total	1,506,000	250,000	287,000	254,000	240,000	220,000	255,000

Appendix B. Most-cited domains and source types

Table B1. Three most-cited domains by AI engine — cited-source counts in parentheses; n = 3,906,000 cited sources

AI engine	Cited sources	Distinct domains	Most-cited domains
ChatGPT	1,950,000	118	datadome.co (211,000), developers.cloudflare.com (184,000), imperva.com (172,000)
Microsoft Copilot	155,000	24	worldmetrics.org (44,000), bot-management-solutions.com (21,000), gitnux.org (17,000)
Gemini	166,000	35	radware.com (19,000), indusface.com (17,000), gartner.com (13,000)
Google AI Mode	766,000	165	gartner.com (67,000), cequence.ai (48,000), geetest.com (42,000)
Google AI Overview	469,000	69	cequence.ai (77,000), cloudflare.com (31,000), akamai.com (29,000)
Perplexity	400,000	31	indusface.com (32,000), clearout.io (24,000), cequence.ai (24,000)

Table B2. Cited sources by source type and AI engine — counts; n = 3,906,000 cited sources

AI engine	Vendor site	Other	Review aggregator	Social	Forum	Marketplace	News	Government	Wikipedia	Total
ChatGPT	1,270,000	581,000	75,000	2,000	12,000	6,000	1,000	2,000	1,000	1,950,000
Microsoft Copilot	12,000	137,000	2,000	0	0	3,000	1,000	0	0	155,000
Gemini	70,000	73,000	22,000	0	1,000	0	0	0	0	166,000
Google AI Mode	244,000	382,000	93,000	26,000	16,000	0	5,000	0	0	766,000
Google AI Overview	198,000	235,000	21,000	12,000	2,000	0	1,000	0	0	469,000
Perplexity	54,000	306,000	40,000	0	0	0	0	0	0	400,000
Total	1,848,000	1,714,000	253,000	40,000	31,000	9,000	8,000	2,000	1,000	3,906,000

Appendix C. Product entries

Table C1. All product entries — 84 entries ordered by mentions; a dash where no mean rank was recorded; n = 240,000 responses

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
1	Cloudflare Bot Management	Cloudflare	191,000	1.9	0.81	6	4
2	Akamai Bot Manager	Akamai	169,000	2.8	0.76	6	4
3	DataDome	DataDome	128,000	2.8	0.81	6	4
4	Imperva Advanced Bot Protection	Imperva	120,000	3.7	0.72	6	4
5	Radware Bot Manager	Radware	77,000	4.3	0.70	6	4
6	F5 Distributed Cloud Bot Defense	F5	55,000	5.1	0.71	6	4
7	DataDome Bot Protect	DataDome	41,000	2.9	0.81	5	4
8	Kasada	Kasada	39,000	4.2	0.76	5	4
9	HUMAN Security	HUMAN Security	38,000	3.5	0.74	6	4
10	Arkose Labs	Arkose Labs	36,000	5.2	0.72	6	4
11	Netacea	Netacea	34,000	5.3	0.73	4	4
12	HUMAN Bot Defender	HUMAN	33,000	3.7	0.73	5	4
13	Imperva Advanced Bot Management	Imperva	32,000	3.6	0.68	6	4
14	Fastly Bot Management	Fastly	22,000	4.0	0.73	5	4
15	HUMAN Sightline	HUMAN	20,000	3.9	0.81	5	4
16	Cequence Bot Management	Cequence	20,000	4.7	0.71	4	4
17	Arkose Bot Manager	Arkose Labs	14,000	4.0	0.72	4	4
18	AWS WAF Bot Control	AWS	10,000	4.8	0.67	4	4
19	CHEQ	CHEQ	9,000	7.3	0.64	2	4
20	Cloudflare Bot Management & AI Crawl Control	Cloudflare	7,000	1.0	0.90	2	4
21	Cequence Security	Cequence	7,000	5.0	0.74	3	2
22	Human Defense Platform	HUMAN Security	6,000	3.4	0.78	4	3
23	Akamai Bot Manager (App & API Protector)	Akamai	5,000	2.0	0.84	2	3
24	HUMAN Sightline Cyberfraud Defense	HUMAN	5,000	4.3	0.80	4	3
25	Imperva	Imperva	5,000	4.4	0.68	2	4
26	Imperva Bot Defense	Imperva	4,000	2.0	0.84	1	3
27	Akamai Bot Manager / Account Protector	Akamai	4,000	2.3	0.82	3	2
28	Cloudflare	Cloudflare	4,000	2.8	0.81	3	3
29	Akamai Account Protector	Akamai	4,000	3.3	0.82	3	2

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
30	PerimeterX	Human Security	4,000	5.5	0.53	1	3
31	Fingerprint	Fingerprint	4,000	7.3	0.66	2	2
32	Forter	Forter	3,000	2.0	0.85	1	3
33	DataDome Bot Protection	DataDome	3,000	3.3	0.83	3	3
34	Kroll	Kroll	3,000	4.7	0.77	1	3
35	Google reCAPTCHA Enterprise	Google	3,000	4.0	0.40	2	2
36	Fastly Next-Gen WAF	Fastly	3,000	5.0	0.70	2	2
37	Akamai	Akamai	3,000	5.3	0.70	3	1
38	AppTrana WAAP	Indusface	3,000	6.0	0.68	1	2
39	GeeTest	GeeTest	3,000	7.3	0.72	2	2
40	Fastly Bot Defense	Fastly	2,000	—	0.80	1	2
41	DataDome Bot & Agent Trust Management	DataDome	2,000	1.5	0.85	1	1
42	Cloudflare Bot & Agent Management	Cloudflare	2,000	1.5	0.80	1	2
43	DataDome Account Protect	DataDome	2,000	2.0	0.88	2	2
44	Kasada Bot Defense & AI Agent Trust	Kasada	2,000	3.5	0.80	2	1
45	Imperva Advanced Bot Management (ABM)	Imperva	2,000	3.5	0.78	1	2
46	AWS WAF + CloudFront	AWS	2,000	4.0	0.55	1	2
47	Kasada Bot Defense	Kasada	2,000	4.5	0.72	2	2
48	Imperva Application Security	Imperva	2,000	6.0	0.70	1	2
49	Sift	Sift	2,000	6.0	0.60	2	1
50	PerimeterX Bot Defender	PerimeterX	2,000	6.0	0.55	1	2
51	AppTrana	AppTrana	2,000	7.0	0.63	2	2
52	Clearout Form Guard	Clearout	2,000	7.0	0.55	2	2
53	Cloudflare Turnstile	Cloudflare	2,000	8.0	0.68	2	2
54	HUMAN Security (Human Sightline)	HUMAN Security	1,000	1.0	0.90	1	1
55	HUMAN Bot Defender / Sightline	HUMAN	1,000	1.0	0.90	1	1
56	Azion Bot Protector	Azion	1,000	—	0.80	1	1
57	Human Security (Human Defense Platform)	Human Security	1,000	—	0.80	1	1
58	Cloudflare Bot Management / Turnstile	Cloudflare	1,000	1.0	0.80	1	1
59	HUMAN Security (Human Bot Defender)	HUMAN Security	1,000	—	0.70	1	1
60	Fortinet	Fortinet	1,000	—	0.70	1	1
61	Barracuda Bot Protection	Barracuda	1,000	—	0.70	1	1
62	AWS WAF	AWS	1,000	—	0.60	1	1

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
63	F5 Distributed Cloud Bot Defense / Account Protection	F5	1,000	2.0	0.90	1	1
64	HUMAN Financial/Cyberfraud Defense	HUMAN Security	1,000	2.0	0.85	1	1
65	Kasada AI Agent Trust	Kasada	1,000	2.0	0.85	1	1
66	Akamai Bot Manager (Premier)	Akamai	1,000	2.0	0.85	1	1
67	Akamai App & API Protector	Akamai	1,000	2.0	0.85	1	1
68	HUMAN Security Sightline / Bot Defender	HUMAN Security	1,000	2.0	0.80	1	1
69	NCC Group	NCC Group	1,000	2.0	0.60	1	1
70	Kasada Bot and Agent Protection	Kasada	1,000	3.0	0.80	1	1
71	Akamai Bot & Agent Control	Akamai	1,000	3.0	0.80	1	1
72	Booz Allen Hamilton	Booz Allen Hamilton	1,000	3.0	0.60	1	1
73	HUMAN Security (Human Bot Defender / Sightline)	HUMAN Security	1,000	4.0	0.75	1	1
74	Google Cloud Armor	Google	1,000	4.0	0.70	1	1
75	Fastly Signal Sciences	Fastly	1,000	4.0	0.50	1	1
76	Netacea Bot Management	Netacea	1,000	5.0	0.80	1	1
77	Cequence Unified API Protection	Cequence Security	1,000	5.0	0.75	1	1
78	Arkose Labs Passages/MatchKey	Arkose Labs	1,000	5.0	0.70	1	1
79	HUMAN Bot Protection	HUMAN Security	1,000	5.0	0.70	1	1
80	Fortinet/F5 Distributed Cloud Bot Defence	Fortinet	1,000	5.0	0.50	1	1
81	Arkose MatchKey	Arkose Labs	1,000	6.0	0.75	1	1
82	Kroll / NCC Group / Booz Allen Hamilton	Kroll	1,000	7.0	0.65	1	1
83	GeeTest Bot Management	GeeTest	1,000	6.0	0.40	1	1
84	Sift Digital Trust & Safety	Sift	1,000	8.0	0.60	1	1

Appendix D. Study parameters

Table D1. Study parameters — as configured for this study

Parameter	Value
Category	Enterprise bot management (bot management, bot mitigation, automated-traffic defence)
AI engines	ChatGPT, Microsoft Copilot, Gemini, Google AI Mode, Google AI Overview, Perplexity
Markets	United States, Canada, India, Germany

Parameter	Value
Query set	10,000 distinct buyer-intent queries, core questions plus paraphrase variants
Responses per engine	40,000
Responses per market	60,000
Total responses	240,000
Cited sources	3,906,000
Brand strings recorded	37
Product entries	84
Collection window	August 2026
Unit of analysis	One engine response to one query in one market
Design	Observational, cross-sectional, single collection window
Statistical treatment	Descriptive statistics only; no inference is reported

Appendix E. Metric computation

Table E1. Metric computation — the arithmetic behind each derived metric in this report

Metric	Computation
Responses per engine	Queries × markets.
Total responses	Queries × engines × markets.
Responses per market	Queries × engines.
Mean response length	Sum of response word counts for an engine, divided by that engine's responses.
Hedging, recency-anchored, truncation and refusal rates	Responses of an engine carrying the characteristic, divided by that engine's responses.
Total cited sources	Sum of the per-engine cited-source counts.
Share of corpus	That engine's cited sources divided by total cited sources.
Source-type share	That engine's cited sources of the type, divided by that engine's cited sources.
Dead links per engine	That engine's cited sources × its dead-link rate.
Total dead links	Sum of the per-engine dead-link counts.
Overall dead-link rate	Total dead links divided by total cited sources.
Total brand mentions	Sum of a brand's per-engine mention counts, summed over brands for the table total.
Engines recording	Count of engines whose mention count for that brand is greater than zero.

Metric	Computation
Mean ordinal position	Sum of the positions at which a brand was named, divided by the mentions carrying a position.
Mean sentiment	Sum of the per-mention sentiment scores for a brand, divided by that brand's mentions.
Jaccard coefficient	Domains cited by both engines, divided by domains cited by either engine.
Mean rank	Sum of a product entry's ordered positions, divided by the mentions carrying a position.
Leading brands by market	Brands ordered by their mention count within that market's responses.