

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	1,960,000	4
Microsoft Copilot	40,000	155,000	4
Gemini	40,000	361,000	4
Google AI Mode	40,000	454,000	4
Google AI Overview	40,000	400,000	4
Perplexity	40,000	400,000	4
Total	240,000	3,730,000	4

Published: September 2026

Research period: August 2026

Category: Enterprise bot management, bot mitigation and automated-traffic defence

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

Evidence base: 3,730,000 cited sources · 98 distinct product entries

Prepared by: GrackerAI · AI Search Visibility Benchmark Series

Abstract

Enterprise security buyers increasingly open an evaluation inside an AI-generated answer rather than on a search results page, so the vendor set an AI engine returns now 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, bot mitigation and automated-traffic defence. 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 3,730,000 cited sources and covered 98 distinct product entries. For every response the study recorded length, cited domains, source reachability, hedging, recency anchoring, truncation, refusal, and each brand mention with its ordinal position. The engines converge on a narrow vendor set: 10 of 16 brands were named by all six engines, and the four leading vendors took 58.7% of the mentions recorded across those brands, led by Cloudflare at 221,000. They diverge on evidence: the highest cross-engine source overlap is a Jaccard coefficient of 0.38, between Gemini and Google AI Overview, and the lowest is 0.07. Source reliability varies by more than a factor of fifteen, from a 1.9% dead-link rate to 29.7%, against an overall rate of 24.4%. Ordinal position and mention volume move independently, and the same three vendors lead every market.

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

- 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

Appendix C. Product entries

Appendix D. Study parameters

Appendix E. Metric computation

1. Introduction

1.1 Background

An enterprise buyer evaluating automated-traffic defence now begins the process inside a generated answer. The buyer asks an AI search engine which products address credential stuffing, scraping or account takeover, and 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 visited a single vendor site.

The structural consequence is narrow. A vendor absent from the returned list is not ranked low; it is not present. The buyer has no mechanism equivalent to scrolling to the second page of results, because a generated answer has no second page. Where a search results page offered ten positions and a set of related queries, a generated answer typically names between three and eight vendors and closes. Visibility in this setting is closer to inclusion than to ranking, and the gap 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 bounded: the engines in this study named a stable core of large edge and application-security vendors alongside a set of independent bot specialists, and 10 of the 16 brands recorded appeared in every engine. A category with a settled vendor population isolates engine behaviour from genuine market churn.

The evaluation criteria are also stable and technical. Across the corpus the engines returned the same selection factors. They are behavioural detection at the edge, device fingerprinting, latency and false-positive characteristics, and policy control over automated crawlers. To those they add native integration with a content delivery network or web application firewall, and compatibility with more than one such network. These criteria are specific enough that responses can be compared to each other rather than to a general description of security software.

Buying behaviour in the category is research-heavy and the cost of invisibility is high. Automated-traffic defence is bought by security and platform teams who shortlist before contacting vendors, and contracts are enterprise-scale and multi-year. A vendor omitted from the generated shortlist loses the opportunity at the earliest stage of the cycle, before any conversation in which the omission could be corrected.

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 tail each engine holds alone. It measures the overlap between the engines' evidence bases directly, as a pairwise coefficient over cited domains, rather than inferring shared sourcing from shared conclusions. It reports mention volume and ordinal position as separate quantities, which allows the two 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 separate bot controls bundled into an edge platform from independent bot-specific layers?

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.

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
Brand mention	One occurrence of a vendor brand name in a response. Repeat occurrences within one response are counted once.	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 polarity score assigned to the sentence containing a brand mention, on a zero-to-one scale where higher is more favourable.	score
Cited source	One URL attached to a response by the engine as a reference.	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
Dead-link rate	Share of an engine's cited URLs that did not resolve to a reachable page when tested.	% 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

Metric	Definition	Unit
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 distinct product string recorded against a brand. A brand may hold several entries where engines named its products differently.	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.

Citations were extracted as the set of URLs the engine attached to a response. Each URL was reduced to its registrable domain for the distinct-domain and overlap measures, and each was tested for reachability; a URL that did not resolve to a reachable page at test time was recorded as a dead link. Reachability was tested once, during the collection window, so the dead-link results describe the state of the cited web at that moment.

Brand mentions were detected against a vendor name list and recorded with the ordinal position at which the brand appeared among the named vendors in that response, counting from one. Product entries were recorded as the distinct product strings the engines used, which is why one brand may hold several entries in Appendix C. Sentiment was scored on the sentence containing each mention and reported on a zero-to-one scale.

Two implementation details are not documented by the instrumentation and are therefore not described here: the specific lexicon used to detect qualifying and date-anchoring language, and the model used to assign sentiment polarity. §6.1 states the consequence for how those two metrics should be read. Source-type classification was recorded by the instrumentation but did not separate the corpus, 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. The three engines above 500 words — ChatGPT, Gemini, Microsoft Copilot — sit close together, and the remaining three fall away steeply.

Truncation was recorded in 4 of the six engines, ranging from 3% to 45%, the latter for Gemini. Microsoft Copilot and Google AI Mode recorded none. Recency anchoring is highest for Microsoft Copilot at 68%, more than double the next value, and lowest for Gemini, which recorded none. Hedging is highest for ChatGPT at 13%. No engine refused a query: the refusal rate is 0% for all six.

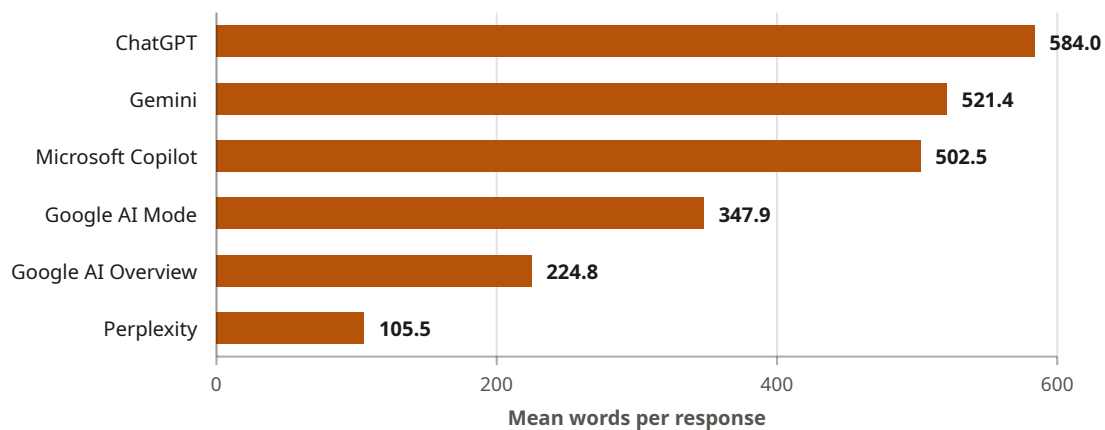


Figure 1. Mean response length by AI engine. Arithmetic mean word count per response; n = 40,000 responses per engine. Source: [GrackerAI AI Search Visibility Benchmark Series](#), September 2026.

4.2 Brand visibility

Brand visibility was recorded for 16 brands. Table 4 gives each brand's total mentions, the number of engines that named it, and its mean sentiment where one was recorded.

Table 4. Brand visibility — all 16 brands recorded; n = 240,000 responses

Brand	Total mentions	Engines recording	Mean sentiment
Cloudflare	221,000	6	0.81
Akamai	205,000	6	0.76
DataDome	203,000	6	0.81
Imperva	194,000	6	0.71

Brand	Total mentions	Engines recording	Mean sentiment
Radware	91,000	6	0.69
HUMAN Security	89,000	6	0.75
Arkose Labs	67,000	6	0.73
F5	64,000	6	0.70
Kasada	63,000	6	0.76
HUMAN	48,000	5	0.75
Netacea	45,000	4	0.72
Fastly	33,000	6	0.70
Cequence	29,000	4	—
CHEQ	22,000	3	—
AWS	17,000	4	—
Kroll	11,000	1	—
Total	1,402,000	—	—

Visibility is concentrated. The four most-mentioned brands — Cloudflare, Akamai, DataDome and Imperva — account for 58.7% of the mentions recorded across those brands, and they sit within 27,000 mentions of each other, from 221,000 for Cloudflare to 194,000 for Imperva. Figure 2 shows the distribution. Below that group the counts fall by more than half, to 91,000 for Radware.

10 of the 16 brands were named by all six engines. Kroll was named by one engine only, Microsoft Copilot, and recorded 11,000 mentions. Mean sentiment, where recorded, occupies a narrow band from 0.69 for Radware to 0.81, the value recorded for both Cloudflare and DataDome.

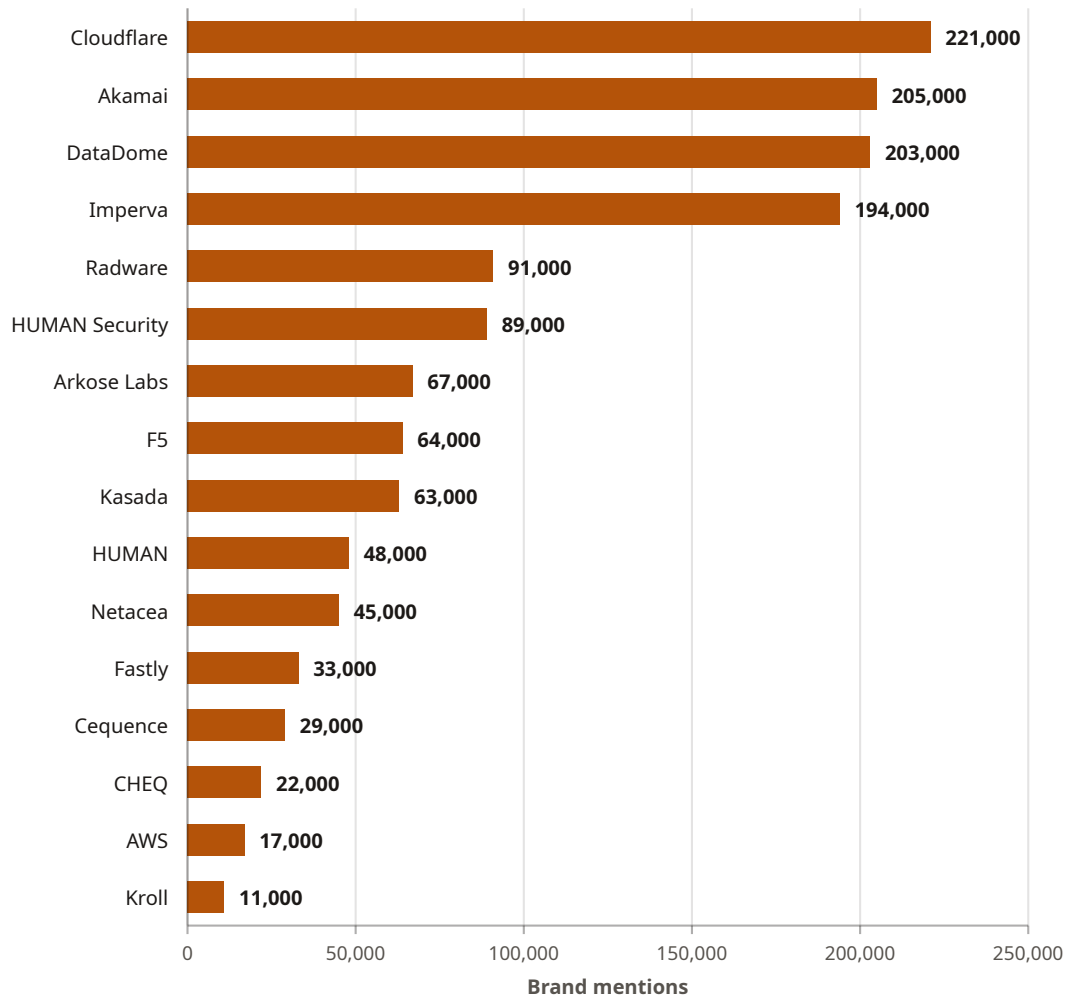


Figure 2. Brand mentions by vendor. All 16 brands recorded; n = 240,000 responses. Source: [GrackerAI AI Search Visibility Benchmark Series](#), September 2026.

4.3 Product-level results

98 distinct product entries were recorded. An entry is one product string as an engine wrote it, so a brand holds several entries where engines named its products differently; Akamai holds 10, the most of any brand. Table 5 gives the ten most-mentioned entries.

Table 5. Product entries, ten most-mentioned — of 98 entries; n = 240,000 responses

Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
Cloudflare Bot Management	Cloudflare	187,000	2.0	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
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

Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
Kasada	Kasada	39,000	4.2	0.76	5	4
Arkose Labs	Arkose Labs	36,000	5.2	0.72	6	4
Netacea	Netacea	34,000	5.3	0.73	4	4

Product entries carry their own brand attributions, and Appendix C lists entries attributed to brands outside the 16 in Table 4. The distribution is long-tailed. 82 of the 98 entries record fewer than 10,000 mentions, and 53 were named by exactly one engine. 8 entries were named by all six engines and 22 were recorded in all four markets. The full table is in Appendix C.

Mean rank among the ten most-mentioned entries runs from 2.0 to 5.3. The leading entry, Cloudflare Bot Management, records 187,000 mentions at a mean rank of 2.0, and the second and third entries, Akamai Bot Manager and DataDome, share a mean rank of 2.8 on 169,000 and 128,000 mentions respectively.

4.4 Citation volume

The corpus contains 3,730,000 cited sources, distributed unevenly across the six engines, as Figure 3 shows.

Table 6. Cited sources and domain breadth by AI engine — n = 3,730,000 cited sources

AI engine	Cited sources	Distinct domains	Most-cited domain
ChatGPT	1,960,000	118	datadome.co (210,000)
Microsoft Copilot	155,000	24	worldmetrics.org (44,000)
Gemini	361,000	36	radware.com (51,000)
Google AI Mode	454,000	72	gartner.com (52,000)
Google AI Overview	400,000	52	cequence.ai (74,000)
Perplexity	400,000	31	indusface.com (32,000)
Total	3,730,000	—	—

ChatGPT accounts for 52.5% of all cited sources in the corpus, 1,960,000 of 3,730,000, and cites from the widest domain population at 118 distinct domains. Microsoft Copilot cites from the narrowest at 24. Distinct-domain counts are reported per engine and are not summed across engines.

The most-cited domain differs for every engine, and no engine's most-cited domain appears as another engine's most-cited domain. Appendix B gives the three most-cited domains for each engine.

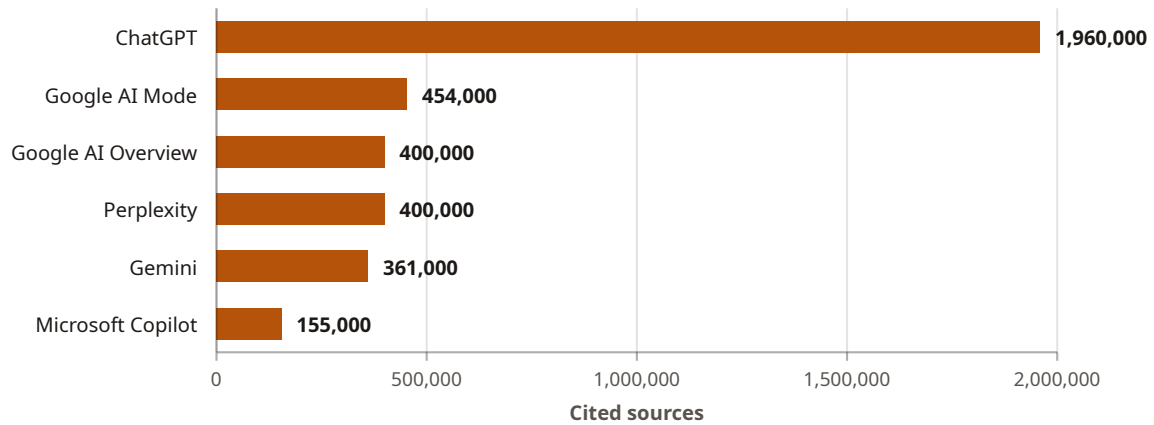


Figure 3. Cited sources by AI engine. URLs attached to responses; n = 3,730,000 cited sources. Source: [GrackerAI AI Search Visibility Benchmark Series](#), September 2026.

4.5 Source quality and link decay

A share of the URLs the engines cited did not resolve to a reachable page when tested. Figure 4 gives the rate for each engine.

Table 7. Dead links and self-citation by AI engine — n = 3,730,000 cited sources

AI engine	Cited sources	Dead-link rate	Dead links	Self-citation rate
ChatGPT	1,960,000	29.7%	582,000	0%
Microsoft Copilot	155,000	1.9%	2,950	0%
Gemini	361,000	17.7%	63,900	0%
Google AI Mode	454,000	26.2%	119,000	0%
Google AI Overview	400,000	21.8%	87,200	0%
Perplexity	400,000	14%	56,000	0%
Total	3,730,000	24.4%	911,050	—

The overall dead-link rate is 24.4%, 911,050 of 3,730,000 cited URLs. Per-engine rates span more than a factor of fifteen, from 1.9% for Microsoft Copilot to 29.7% for ChatGPT. ChatGPT contributes 1,960,000 of the 3,730,000 cited sources in the corpus and 582,000 of the 911,050 dead links.

The self-citation rate is 0% for all six engines, including the three operated by the same parent as one of the search properties they could cite.

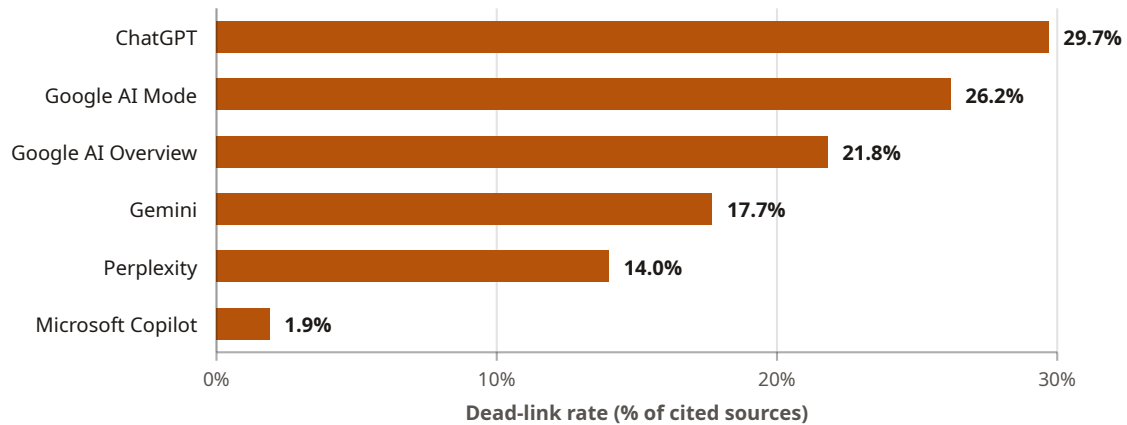


Figure 4. Dead-link rate by AI engine. Share of cited URLs that did not resolve when tested; n = 3,730,000 cited sources. Source: [GrackerAI AI Search Visibility Benchmark Series](#), September 2026.

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 8 gives the full matrix.

Table 8. Cross-engine cited-domain overlap — Jaccard coefficient over cited-domain sets; 0 is disjoint, 1 is identical; n = 3,730,000 cited sources

	ChatGPT	Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
ChatGPT	1.00	0.10	0.18	0.22	0.22	0.17
Copilot	0.10	1.00	0.07	0.07	0.09	0.25
Gemini	0.18	0.07	1.00	0.33	0.38	0.31
Google AI Mode	0.22	0.07	0.33	1.00	0.33	0.20
Google AI Overview	0.22	0.09	0.38	0.33	1.00	0.24
Perplexity	0.17	0.25	0.31	0.20	0.24	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 0.21. The highest value is 0.38, between Gemini and Google AI Overview. The lowest is 0.07, recorded twice, for Microsoft Copilot with Gemini and with Google AI Mode. The three engines operated by the same parent — Gemini, Google AI Mode and Google AI Overview — record coefficients of 0.33, 0.38 and 0.33 with each other, the three highest values in the matrix.

Microsoft Copilot records the lowest coefficients with four of the other five engines, between 0.07 and 0.10, and its highest value anywhere in the matrix is 0.25, with Perplexity. No pair of engines outside the same-parent group exceeds 0.31.

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. Table 9 gives both for the two brands that took a first-mention position, together with Akamai, the third most-mentioned brand, which took none.

Table 9. Mean ordinal position of the leading vendors — lower is earlier in the response; n = 240,000 responses

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

Cloudflare holds the first-mention position in 5 of the 6 engines and records the highest mention total in the corpus at 221,000. DataDome holds it in 1, ChatGPT, on 203,000 mentions.

The two measures move independently. In ChatGPT DataDome records a mean ordinal position of 1.5 against 3.1 for Cloudflare, so the brand named earlier is not the brand named most often. In Perplexity the ordering reverses and widens: Cloudflare records 1.0, the earliest mean position recorded anywhere in Table 9, while DataDome records 4.3, the latest. Akamai holds no first-mention position and its mean position varies least across the engines, from 2.5 to 3.3.

4.8 Geographic variance

Table 10 gives the three leading vendors in each market.

Table 10. Leading products by market — n = 60,000 responses per market

Market	First	Second	Third
United States	Cloudflare Bot Management	DataDome	Akamai Bot Manager
Canada	Cloudflare Bot Management	Akamai Bot Manager	DataDome
India	DataDome	Cloudflare Bot Management	Akamai Bot Manager
Germany	Cloudflare Bot Management	Akamai Bot Manager	DataDome

The same three vendors lead all four markets, in three different orders. Cloudflare leads three of them; India names DataDome first. United States places DataDome second and Akamai third, and Canada and Germany reverse that pair. No vendor outside the three appears in any market's leading positions, and no market returns a vendor absent from the other three.

At product level the same stability holds: 22 of the 98 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 control over automated crawlers. ChatGPT returns explicit crawler-trust features in DataDome and Cloudflare as a selection criterion. Perplexity rarely returns crawler controls at all. The two engines therefore describe a different set of decision criteria for the same category, while returning substantially the same vendors.

The second concerns analyst positioning. ChatGPT is the only engine to return the claim that DataDome is a leader in the 2026 Forrester Wave for bot and agent trust management. No other engine returns that claim in any market.

The third concerns detection architecture. Gemini is the only engine to single out Cequence for network-based detection that requires no client-side code. Cequence records 29,000 mentions across the corpus, against 221,000 for Cloudflare. Microsoft Copilot separately returns vendors that no other engine associates with the category, including Forter at 3,000 mentions and Kroll at 11,000, both recorded in Microsoft Copilot alone. The divergences occur on capability claims and on 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 the core and diverge in the tail. §4.2 shows 10 of the 16 recorded brands named by all six engines, and §4.3 shows only 8 of 98 product entries reaching the same coverage. The difference between those two figures is the answer: the engines agree about companies and disagree about product names. A brand is recognised across all six; the string an engine uses for that brand's product is not.

The tail is where the engines part. 53 of the 98 entries were named by exactly one engine, and §4.9 shows Microsoft Copilot returning vendors the other five do not associate with the category. A plausible explanation is that the core set is reachable from almost any part of the category's web presence, while the tail depends on which sources a particular engine happens to read. §4.6 is consistent with that: the engines reach the same core from evidence bases that overlap by a mean Jaccard coefficient of 0.21.

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

Highly concentrated at the top and thin below it. §4.2 shows the four leading brands taking 58.7% of the mentions recorded across those brands, and the next brand recording less than half the fourth's total. §4.3 shows 82 of 98 product entries below 10,000 mentions.

The shape matters more than the ordering within it. The four leading brands are separated by 27,000 mentions in total, which is small against the 103,000 that separates the fourth from the fifth. 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 0.21 over cited domains, with a maximum of 0.38 and a minimum of 0.07. The three highest values in the matrix all fall between engines operated by the same parent, and no pair outside that group exceeds 0.31. Six engines returned substantially the same vendors while reading substantially different parts of the web.

Volume compounds the divergence. §4.4 shows ChatGPT contributing 52.5% of the corpus's cited sources and citing from 118 distinct domains against 24 for Microsoft Copilot. The engines differ not only in which sources they read but by how much they 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, and there is no single source list that serves all six.

5.4 RQ4 — Does ordinal position track mention volume?

No. §4.7 shows Cloudflare recording the highest mention total in the corpus while recording a mean ordinal position of 3.1 in ChatGPT, against 1.5 for DataDome, the brand ChatGPT names first. The same two brands reverse in Perplexity, and the gap there is wider than in any other engine.

The two measures answer different questions. Mention volume records how often a vendor enters the conversation; ordinal position records where in the ordered list it lands once it has. A vendor can be present in most responses and consistently placed third, or present in fewer and placed first. §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 frequency is incomplete, and one that reports only position is equally so.

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

Not materially. §4.8 shows the same three vendors leading all four markets. Their order varies: India is the only market to change the leading vendor, and United States ranks the second and third differently from Canada and Germany. The set does not change. 22 of 98 product entries were recorded in all four markets, and every entry in Table 5 was.

The result is a negative one 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 separate bot controls bundled into an edge platform from independent bot-specific layers?

The engines name both kinds of vendor and do not order them consistently. §4.2 places Cloudflare and Akamai, whose controls sit inside an edge platform, first and second by mention volume, with DataDome, an independent layer, third at 203,000 mentions against 205,000. §4.7 shows the ordering between the two kinds changing by engine rather than holding.

The corpus records the distinction as a stated trade-off rather than as a ranking. Responses associate edge-bundled controls with lower latency and fewer moving parts, and independent layers with compatibility across more than one delivery network and with bot-specific detection depth. §4.9 shows the criteria themselves varying between engines: ChatGPT returns crawler-trust features as a selection factor and Perplexity rarely returns them. The data is consistent with the engines reproducing a genuine architectural split in the category, and it does not establish that either side of that split is favoured.

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.

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. Sentiment was scored by an undocumented model (§3.7), so the narrow band it produces across brands should not be read as evidence that the engines describe those brands with comparable warmth.

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.

Sentiment is scored on the sentence containing each mention by a model the instrumentation does not document (§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 scorer 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.

Truncation is an inference from the terminal characters of a response, not an observation of a process. It cannot separate a response that stopped mid-sentence from one whose last sentence lacked terminal punctuation, and it is reported here as a characteristic of the recorded text rather than as a property of the engine. Hedging and recency anchoring rest on undocumented lexicons (§3.7) and are subject to the same limit. The same detector was 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, the cited domains and the reachability results describe that window. A collection window that included a major analyst publication or a vendor acquisition 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 and bounded vendor set (§1.2). That is what makes it a usable instrument, and it is also what limits generalisation. A category with a larger or faster-moving vendor population would not be expected to produce the same concentration.

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. The concentration of visibility into a small leading group, the separation between mention volume and ordinal position, the elevated overlap between same-parent engines, and the stability of the vendor set across markets. All four 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 in particular, because reachability was tested once at collection time (§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.
- Two processing steps are undocumented by the instrumentation: the lexicon behind the hedging and recency measures, and the sentiment model. Both are stated where they are used.
- Source-type classification did not separate the corpus and is therefore not reported, so the study says nothing about the kinds of site the engines cite, only about the domains.

8. Practical Implications

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

Treat inclusion and ordering as two problems. §4.2 and §4.7 show that the most-mentioned brand is not the earliest-named one in every engine. Measure both, and decide which one is the constraint before acting: a vendor outside the group in §4.2 has an inclusion problem, and a vendor inside it with a late mean position has an ordering problem. The two do not respond to the same work.

Measure every engine separately. §4.6 gives a mean pairwise cited-domain overlap of 0.21, and the highest values in the matrix are between engines sharing a parent. A source list that serves ChatGPT is not the list that serves Microsoft Copilot, whose coefficients with the others are the lowest recorded. Budget for six measurements, not one.

Publish the specifics the engines reuse. §1.2 and §4.9 show the engines returning a consistent set of selection criteria: integration with a content delivery network or web application firewall, compatibility with more than one such network, latency and false-positive characteristics, and policy control over automated crawlers. Documentation that states these concretely gives an engine something specific to extract; a general description of the category does not.

Do not build a market-by-market visibility programme for these four markets. §4.8 shows the same three vendors leading all four, with only the first two places reversed in India. Effort that would go into market-specific work is better spent on the inclusion problem, which is the same problem in every market tested.

Maintain the pages the engines cite. §4.5 records an overall dead-link rate of 24.4%, reaching 29.7% for ChatGPT. A vendor's own retired or moved URLs are part of that population, and a citation pointing at an unreachable page carries the brand name without carrying the reader.

9. Conclusion

This study put 10,000 enterprise buyer-intent queries to six AI search engines in four markets, producing 240,000 responses containing 3,730,000 cited sources and covering 98 distinct product entries in enterprise bot management, bot mitigation and automated-traffic defence.

The engines agree about vendors and disagree about evidence. 10 of 16 brands were named by all six engines, and the four leading vendors took 58.7% of the mentions recorded across those brands, from 221,000 for Cloudflare to 194,000 for Imperva. Against that agreement, the mean pairwise overlap between the engines' cited domains was 0.21, and the highest value recorded anywhere in the matrix was 0.38, between two engines operated by the same parent. Six systems reading substantially different parts of the web arrived at substantially the same list of names.

Two further results qualify that list. Ordinal position does not track mention volume: Cloudflare takes the first-mention position in 5 of the six engines, yet in the engine where DataDome is named first its mean ordinal position is 3.1 against 1.5. Source reliability varies by more than a factor of fifteen across the engines, from 1.9% to 29.7% of cited URLs unreachable, against an overall rate of 24.4%. The vendor set did not vary materially across the four markets.

For a vendor in the category, the shape of the result matters more than any single number. Visibility here is concentrated into a small group of names the engines reuse across markets, phrasings and engines, and it is reached through six largely separate evidence bases. Entering that group is a different problem from improving a position within it, and it has 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, the brand and product mention records with their ordinal positions, and the URL reachability results. The dataset for this report was generated in August 2026.

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

One of the domains listed among Gemini's most-cited sources in Appendix B is operated by a co-founder of GrackerAI. It was recorded by the same detection applied to every other domain in the study, and neither it nor its position was selected, solicited or adjusted.

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. (September 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 16 brands 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	221,000	37,000 (3.1)	34,000 (1.2)	39,000 (2.5)	38,000 (1.7)	35,000 (1.7)	38,000 (1.0)
Akamai	205,000	35,000 (3.3)	32,000 (2.6)	39,000 (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)
Imperva	194,000	30,000 (4.9)	35,000 (3.5)	31,000 (4.3)	35,000 (3.4)	33,000 (3.2)	30,000 (2.9)
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)
HUMAN Security	89,000	21,000 (3.0)	14,000 (5.7)	18,000 (2.7)	15,000 (3.8)	12,000 (2.8)	9,000 (5.0)
Arkose Labs	67,000	17,000 (4.2)	14,000 (6.2)	2,000 (5.0)	7,000 (5.0)	13,000 (4.1)	14,000 (6.7)
F5	64,000	10,000 (4.3)	4,000 (4.7)	10,000 (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
HUMAN	48,000	16,000 (3.0)	—	7,000 (5.3)	11,000 (4.1)	9,000 (3.9)	5,000 (4.0)
Netacea	45,000	8,000 (5.8)	27,000 (5.1)	—	—	1,000 (6.0)	9,000 (5.7)
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)
Cequence	29,000	—	—	5,000 (5.0)	4,000 (5.5)	8,000 (5.6)	12,000 (3.0)
CHEQ	22,000	2,000	6,000	—	—	—	14,000 (7.3)
AWS	17,000	—	10,000 (4.0)	2,000 (5.0)	2,000 (8.0)	—	3,000
Kroll	11,000	—	11,000 (5.3)	—	—	—	—
Total	1,402,000	243,000	246,000	227,000	232,000	207,000	247,000

Appendix B. Most-cited domains

Table B1. Three most-cited domains by AI engine — cited-source counts in parentheses; n = 3,730,000 cited sources

AI engine	Cited sources	Distinct domains	Most-cited domains
ChatGPT	1,960,000	118	datadome.co (210,000), developers.cloudflare.com (182,000), imperva.com (168,000)
Microsoft Copilot	155,000	24	worldmetrics.org (44,000), bot-management-solutions.com (21,000), gitnux.org (17,000)
Gemini	361,000	36	radware.com (51,000), cequence.ai (43,000), guptadeepak.com (38,000)
Google AI Mode	454,000	72	gartner.com (52,000), cequence.ai (36,000), geetest.com (33,000)
Google AI Overview	400,000	52	cequence.ai (74,000), cloudflare.com (29,000), akamai.com (28,000)
Perplexity	400,000	31	indusface.com (32,000), worldmetrics.org (24,000), cequence.ai (24,000)

Appendix C. Product entries

Table C1. All product entries — 98 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	187,000	2.0	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	Arkose Labs	Arkose Labs	36,000	5.2	0.72	6	4
10	Netacea	Netacea	34,000	5.3	0.73	4	4
11	HUMAN Bot Defender	HUMAN	33,000	3.7	0.73	5	4
12	Imperva Advanced Bot Management	Imperva	32,000	3.6	0.68	6	4
13	HUMAN Security	HUMAN Security	29,000	3.5	0.76	5	4
14	Fastly Bot Management	Fastly	22,000	4.0	0.73	5	4

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
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	9,000	3.5	0.76	4	4
18	CHEQ	CHEQ	9,000	7.3	0.64	2	4
19	AWS WAF Bot Control	AWS	8,000	5.0	0.66	3	4
20	Cloudflare Bot Management & AI Crawl Control	Cloudflare	6,000	1.0	0.90	2	4
21	Human Defense Platform	HUMAN Security	5,000	3.3	0.79	3	3
22	HUMAN	HUMAN	5,000	3.5	0.78	3	4
23	HUMAN Sightline Cyberfraud Defense	HUMAN	5,000	4.3	0.80	4	3
24	Imperva	Imperva	5,000	4.4	0.68	2	4
25	Arkose Bot Manager	Arkose	5,000	5.0	0.66	3	3
26	Cloudflare Bot Manager	Cloudflare	4,000	1.0	0.85	1	3
27	Imperva Bot Defense	Imperva	4,000	2.0	0.84	1	3
28	Cloudflare	Cloudflare	4,000	2.8	0.81	3	3
29	Akamai Account Protector	Akamai	4,000	3.3	0.82	3	2
30	HUMAN Security	HUMAN	4,000	3.3	0.56	2	2
31	PerimeterX	Human Security	4,000	5.5	0.53	1	3
32	Fingerprint	Fingerprint	4,000	7.3	0.66	2	2
33	Forter	Forter	3,000	2.0	0.85	1	3
34	Akamai Bot Manager (App & API Protector)	Akamai	3,000	2.0	0.83	2	2
35	Akamai Bot Manager / Account Protector	Akamai	3,000	2.7	0.80	2	1
36	DataDome Bot Protection	DataDome	3,000	3.3	0.83	3	3
37	Kroll	Kroll	3,000	4.7	0.77	1	3
38	Google reCAPTCHA Enterprise	Google	3,000	4.0	0.40	2	2
39	Fastly Next-Gen WAF	Fastly	3,000	5.0	0.70	2	2
40	Akamai	Akamai	3,000	5.3	0.70	3	1
41	Cequence	Cequence	3,000	5.5	0.70	2	2
42	AppTrana WAAP	Indusface	3,000	6.0	0.68	1	2
43	GeeTest	GeeTest	3,000	7.3	0.72	2	2
44	Fastly Bot Defense	Fastly	2,000	—	0.80	1	2
45	DataDome Bot & Agent Trust Management	DataDome	2,000	1.5	0.85	1	1
46	Cloudflare Bot & Agent Management	Cloudflare	2,000	1.5	0.80	1	2

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
47	DataDome Account Protect	DataDome	2,000	2.0	0.88	2	2
48	Akamai Bot Manager & App & API Protector	Akamai	2,000	2.0	0.85	2	1
49	Kasada Bot Defense & AI Agent Trust	Kasada	2,000	3.5	0.80	2	1
50	Imperva Advanced Bot Management (ABM)	Imperva	2,000	3.5	0.78	1	2
51	Cequence Security	Cequence Security	2,000	4.5	0.80	1	2
52	Kasada Bot Defense	Kasada	2,000	4.5	0.72	2	2
53	Cequence Security	Cequence	2,000	5.0	0.75	1	1
54	Sift	Sift	2,000	6.0	0.60	2	1
55	PerimeterX Bot Defender	PerimeterX	2,000	6.0	0.55	1	2
56	AppTrana	AppTrana	2,000	7.0	0.63	2	2
57	Clearout Form Guard	Clearout	2,000	7.0	0.55	2	2
58	Cloudflare Turnstile	Cloudflare	2,000	8.0	0.68	2	2
59	Cloudflare Bot Management / AI Crawl Control	Cloudflare	1,000	1.0	0.90	1	1
60	Akamai Bot Manager + Account Protector	Akamai	1,000	1.0	0.90	1	1
61	HUMAN Security (Human Sightline)	HUMAN Security	1,000	1.0	0.90	1	1
62	HUMAN Bot Defender / Sightline	HUMAN	1,000	1.0	0.90	1	1
63	Azion Bot Protector	Azion	1,000	—	0.80	1	1
64	Human Security (Human Defense Platform)	Human Security	1,000	—	0.80	1	1
65	Cloudflare Bot Management / Turnstile	Cloudflare	1,000	1.0	0.80	1	1
66	AWS WAF with Bot Control	AWS	1,000	—	0.70	1	1
67	HUMAN Security (Human Bot Defender)	HUMAN Security	1,000	—	0.70	1	1
68	Fortinet	Fortinet	1,000	—	0.70	1	1
69	Barracuda Bot Protection	Barracuda	1,000	—	0.70	1	1
70	AWS WAF	AWS	1,000	—	0.60	1	1
71	F5 Distributed Cloud Bot Defense / Account Protection	F5	1,000	2.0	0.90	1	1
72	HUMAN Financial/Cyberfraud Defense	HUMAN Security	1,000	2.0	0.85	1	1
73	Kasada AI Agent Trust	Kasada	1,000	2.0	0.85	1	1
74	Akamai Bot Manager (Premier)	Akamai	1,000	2.0	0.85	1	1
75	Akamai App & API Protector	Akamai	1,000	2.0	0.85	1	1
76	HUMAN Security Sightline / Bot Defender	HUMAN Security	1,000	2.0	0.80	1	1
77	NCC Group	NCC Group	1,000	2.0	0.60	1	1

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
78	Kasada Bot and Agent Protection	Kasada	1,000	3.0	0.80	1	1
79	Akamai Bot & Agent Control	Akamai	1,000	3.0	0.80	1	1
80	AWS WAF + Bot Control	AWS	1,000	3.0	0.70	1	1
81	AWS WAF + CloudFront	AWS	1,000	3.0	0.60	1	1
82	Booz Allen Hamilton	Booz Allen Hamilton	1,000	3.0	0.60	1	1
83	HUMAN Security (Human Bot Defender / Sightline)	HUMAN Security	1,000	4.0	0.75	1	1
84	HUMAN Defense	HUMAN	1,000	4.0	0.75	1	1
85	Google Cloud Armor	Google	1,000	4.0	0.70	1	1
86	Fastly Signal Sciences	Fastly	1,000	4.0	0.50	1	1
87	Netacea Bot Management	Netacea	1,000	5.0	0.80	1	1
88	Cequence Unified API Protection	Cequence Security	1,000	5.0	0.75	1	1
89	Arkose Labs Passages/MatchKey	Arkose Labs	1,000	5.0	0.70	1	1
90	HUMAN Bot Protection	HUMAN Security	1,000	5.0	0.70	1	1
91	Imperva Application Security	Imperva	1,000	5.0	0.70	1	1
92	AWS WAF with CloudFront	AWS	1,000	5.0	0.50	1	1
93	Fortinet/F5 Distributed Cloud Bot Defence	Fortinet	1,000	5.0	0.50	1	1
94	Arkose MatchKey	Arkose Labs	1,000	6.0	0.75	1	1
95	Imperva Application Security Platform	Imperva	1,000	7.0	0.70	1	1
96	Kroll / NCC Group / Booz Allen Hamilton	Kroll	1,000	7.0	0.65	1	1
97	GeeTest Bot Management	GeeTest	1,000	6.0	0.40	1	1
98	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 mitigation and automated-traffic defence
AI engines	ChatGPT, Microsoft Copilot, Gemini, Google AI Mode, Google AI Overview, Perplexity
Markets	United States, Canada, India, Germany
Query set	10,000 distinct buyer-intent queries, core questions plus paraphrase variants
Responses per engine	40,000

Parameter	Value
Responses per market	60,000
Total responses	240,000
Cited sources	3,730,000
Brands recorded	16
Product entries	98
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.
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.
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 sentence-level polarity scores for a brand's mentions, divided by that brand's mentions.
Jaccard coefficient	Domains cited by both engines, divided by domains cited by either engine.

Metric	Computation
Mean rank	Sum of a product entry's ordered positions, divided by the mentions carrying a position.