

AI Search Visibility in Enterprise Attack Surface 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 attack surface management, external attack surface management and cyber asset attack surface management 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	357,000	4
Microsoft Copilot	40,000	156,000	4
Gemini	40,000	112,000	4
Google AI Mode	40,000	175,000	4
Google AI Overview	40,000	351,000	4
Perplexity	40,000	21,000	4
Total	240,000	1,172,000	4

Published: August 2026

Research period: August 2026

Category: Enterprise attack surface management (ASM, EASM, CAASM)

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

Evidence base: 1,172,000 cited sources · 124 distinct product entries

Prepared by: GrackerAI · AI Search Visibility Benchmark Series

Abstract

Enterprise security buyers now begin an attack surface management 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 attack surface 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 1,172,000 cited sources and covered 124 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 narrow core: 7 of 87 brand strings were named by all six engines, led by CyCognito at 160,000 mentions and Palo Alto Networks at 157,000, 3,000 apart. Visibility below that core is thin, with 47 brand strings named by a single engine. The brand named most often is not the brand named first in most engines: Palo Alto Networks holds the earlier mean ordinal position in 4 of the six. The engines cite largely separate domains, with a mean pairwise overlap of 0.16, and no two engines share a most-cited domain, yet 51.3% of everything they cite sits on a vendor's own site. Dead links are rare, at 0.4% of cited pages. The same two vendors lead every market, in an order that changes in one of the four.

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

An enterprise buyer evaluating attack surface management now begins inside a generated response. The buyer asks an AI search engine which products discover unknown internet-facing assets, attribute them to the organisation and its subsidiaries, and feed the result into the tools the security team already runs. 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 attack surface 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 7 of the 87 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, organised around external attack surface management, cyber asset attack surface management and continuous threat exposure management. Across the corpus the engines returned the same selection factors: discovery of unknown and subsidiary assets, seedless or low-seed attribution, internet-scale scanning, and integration with ticketing, workflow and analytics tools already in place. 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-intensive and procurement cycles are long, so the cost of omission from an early shortlist is high and difficult to reverse later in the cycle. A vendor omitted from the generated shortlist loses the opportunity at the earliest stage, 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 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. This edition reprocesses the same collected responses under the instrumentation's current citation, source-type and entity rules; the responses themselves are unchanged.

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 and the first-mention results in §4.7.

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, government, 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 and review aggregators were matched against lists and host-name patterns, and the remainder is other.

Reachability was tested once per page during reprocessing. 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.

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	647.4	10%	15%	97%	0%	Palo Alto Networks
Microsoft Copilot	40,000	507.4	17%	60%	0%	0%	Palo Alto Networks
Gemini	40,000	376.2	0%	3%	23%	0%	CyCognito
Google AI Mode	40,000	284.2	0%	13%	0%	0%	CyCognito
Google AI Overview	40,000	232.3	0%	7%	0%	0%	CyCognito
Perplexity	40,000	524.0	33%	0%	33%	0%	Microsoft

Mean response length spans a factor of 2.8, from 232.3 words for Google AI Overview to 647.4 for ChatGPT. Figure 1 shows the distribution. Three engines exceed 500 words — ChatGPT, Perplexity and Microsoft Copilot — and the two Google search surfaces fall below 300.

Truncation was recorded in 3 of the six engines: 97% for ChatGPT, the highest value in the study, 33% for Perplexity and 23% for Gemini. Microsoft Copilot and the two Google search surfaces recorded none. Recency anchoring is highest for Microsoft Copilot at 60%, and Perplexity recorded none. Hedging is highest for Perplexity at 33%, followed by Microsoft Copilot at 17% and ChatGPT at 10%; the three engines operated by the same parent recorded none. No engine refused a query: the refusal rate is 0% for all six. Three engines record CyCognito as the first-mention vendor, two record Palo Alto Networks and one records Microsoft; §4.7 reports the position data behind that column.

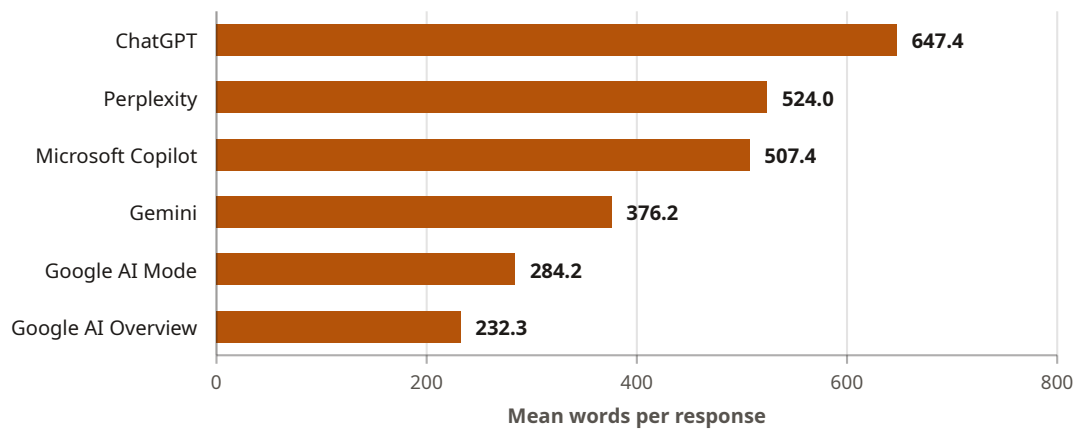


Figure 1. Mean response length by AI engine. Arithmetic mean word count per response; n = 40,000 responses per engine. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

4.2 Brand visibility

Brand visibility was recorded for 87 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 87 brand strings recorded; mean sentiment where recorded, a dash otherwise; n = 240,000 responses

Brand	Total mentions	Engines recording	Mean sentiment
CyCognito	160,000	6	0.81
Palo Alto Networks	157,000	6	0.84
Microsoft	115,000	6	0.77
Censys	113,000	5	0.78
Tenable	102,000	6	0.79
CrowdStrike	94,000	6	0.80
IONIX	75,000	6	0.81
Rapid7	42,000	5	0.69
Bitsight	35,000	5	0.80
Qualys	33,000	5	0.71
Wiz	28,000	5	0.80
UpGuard	25,000	4	—

Visibility is concentrated at the top and long-tailed below it. The two most-mentioned brands, CyCognito and Palo Alto Networks, account for 25.5% of all mentions recorded across the 87 brand strings, at 160,000 and 157,000 mentions, 3,000 apart. Microsoft follows at 115,000 and Censys at 113,000, 2,000 apart; the four together take 43.8%. Below Tenable at 102,000 no brand exceeds 94,000.

7 of the 87 brand strings were named by all six engines: CyCognito, Palo Alto Networks, Microsoft, Tenable, CrowdStrike, IONIX and Axonius. 47 were named by exactly one engine, none of them above 9,000 mentions. The strings recorded include adjacent product categories — vulnerability management, threat intelligence, ticketing and analytics — and one analyst firm, 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 0.69 for Rapid7 to 0.84 for Palo Alto Networks, a spread of 0.15.

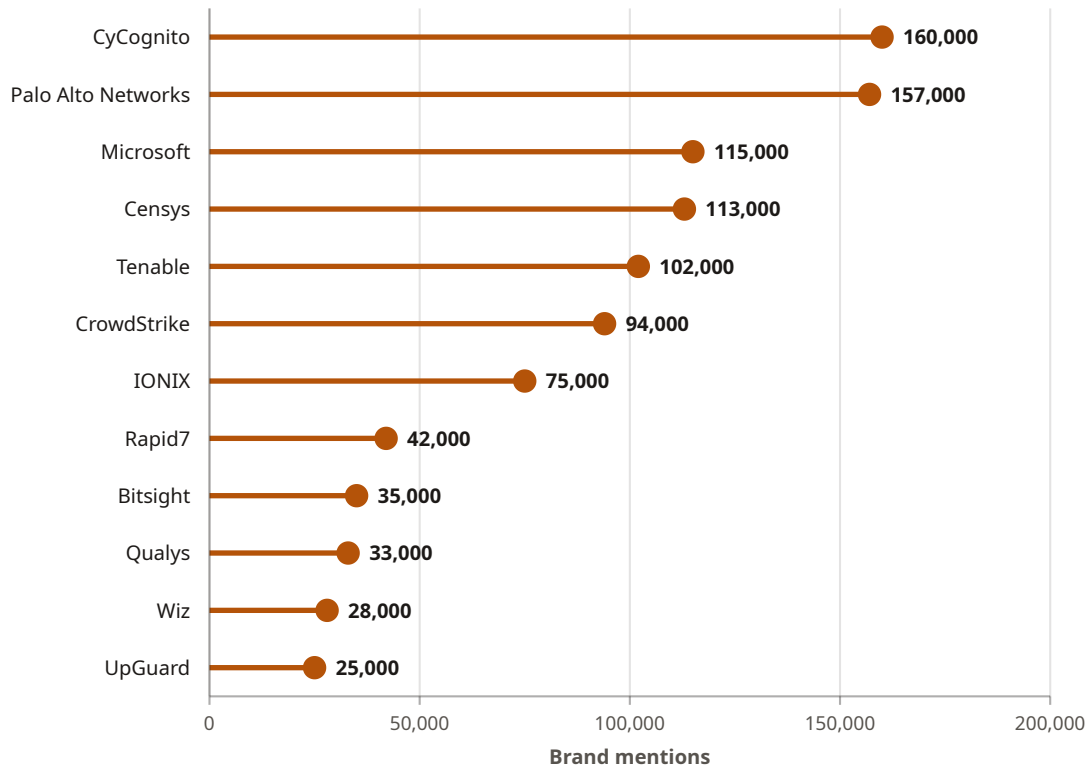


Figure 2. Brand mentions, twelve most-mentioned brands. The 12 most-mentioned of 87 brand strings recorded; n = 240,000 responses. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

4.3 Product-level results

124 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: Rapid7 holds 10, Palo Alto Networks 9 and Tenable 8. Table 5 gives the ten most-mentioned entries.

Table 5. Product entries, ten most-mentioned — of 124 entries; n = 240,000 responses

Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
CyCognito	CyCognito	160,000	2.6	0.82	6	4
Cortex Xpanse	Palo Alto Networks	149,000	1.8	0.86	6	4
Microsoft Defender EASM	Microsoft	112,000	3.4	0.78	6	4
Censys	Censys	107,000	3.5	0.81	5	4
IONIX	IONIX	73,000	3.9	0.82	6	4

Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
CrowdStrike Falcon Exposure Management	CrowdStrike	52,000	3.1	0.81	6	4
Tenable One	Tenable	51,000	3.1	0.80	5	4
CrowdStrike Falcon Surface	CrowdStrike	39,000	4.6	0.81	5	4
Tenable Attack Surface Management	Tenable	30,000	4.7	0.81	6	4
Mandiant ASM	Google Cloud	28,000	4.1	0.78	5	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. 102 of the 124 entries record fewer than 10,000 mentions, 69 record exactly 1,000, and 81 were named by exactly one engine. 6 entries were named by all six engines, 29 were recorded in all four markets, and 73 were recorded in a single market. The full table is in Appendix C.

Mean rank among the ten most-mentioned entries runs from 1.8 to 4.7. The leading entry, CyCognito, records 160,000 mentions at a mean rank of 2.6, and the second, Cortex Xpanse, records 149,000 at 1.8, the earlier of the two.

4.4 Citation volume and source type

The corpus contains 1,172,000 cited sources, distributed 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 = 1,172,000 cited sources

AI engine	Cited sources	Share of corpus	Distinct domains	Most-cited domain
ChatGPT	357,000	30.5%	56	paloaltonetworks.com (59,000)
Microsoft Copilot	156,000	13.3%	33	uinat.com (24,000)
Gemini	112,000	9.6%	36	synack.com (17,000)
Google AI Mode	175,000	14.9%	47	gartner.com (33,000)
Google AI Overview	351,000	29.9%	73	cycognito.com (53,000)
Perplexity	21,000	1.8%	19	branddefense.io (2,000)
Total	1,172,000	—	—	—

ChatGPT and Google AI Overview together account for 60.4% of all cited sources in the corpus, at 357,000 and 351,000. Google AI Overview cites from the widest domain population at 73 distinct domains, followed by ChatGPT at 56. Perplexity cites from the narrowest at 19. Distinct-domain counts are reported per engine and are not summed across engines.

No two engines share a most-cited domain. ChatGPT's is paloaltonetworks.com, Google AI Overview's is cycognito.com, Google AI Mode's is gartner.com, an analyst firm, and Microsoft Copilot's is uinat.com. Appendix B gives the three most-cited domains for each engine.

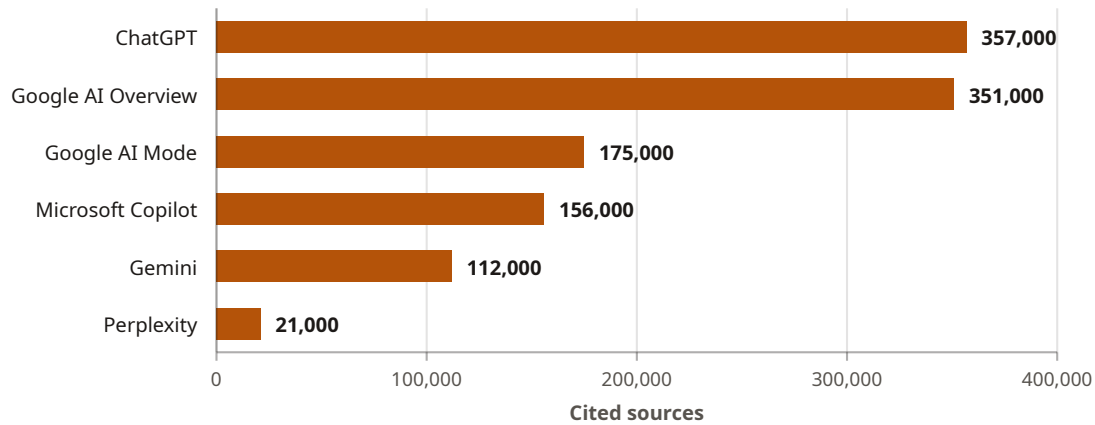


Figure 3. Cited sources by AI engine. Sources cited per engine; n = 1,172,000 cited sources. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

Table 7. Source type by AI engine — share of each engine's cited sources by the class of the cited domain; n = 1,172,000 cited sources

AI engine	Vendor site	Other	Review aggregator	Social	Forum	Parent ecosystem	Marketplace	Wikipedia
ChatGPT	69.7%	12.9%	16.5%	0.3%	0.0%	0.0%	0.3%	0.3%
Microsoft Copilot	23.1%	71.8%	5.1%	0.0%	0.0%	0.0%	0.0%	0.0%
Gemini	56.2%	36.6%	7.1%	0.0%	0.0%	0.0%	0.0%	0.0%
Google AI Mode	36.6%	35.4%	20.0%	6.9%	0.6%	0.6%	0.0%	0.0%
Google AI Overview	49.0%	39.9%	5.4%	4.3%	0.9%	0.3%	0.3%	0.0%
Perplexity	81.0%	14.3%	4.8%	0.0%	0.0%	0.0%	0.0%	0.0%
All engines	51.3%	34.5%	11.1%	2.4%	0.3%	0.2%	0.2%	0.1%

Across the corpus, 51.3% of cited sources sit on a vendor site, 601,000 of 1,172,000, and 34.5% fall outside every named class. Review aggregators, which here include analyst firms, account for 11.1%. The mix differs by engine. Vendor sites are the largest class for 5 of the six engines; Microsoft Copilot is the exception, with 71.8% of its citations in the other class and 23.1% on vendor sites. ChatGPT draws 69.7% of its citations from vendor sites and Gemini 56.2%. Review aggregators reach 20.0% of Google AI Mode's citations, the highest share. No news, academic or government domain was recorded by any engine beyond single citations. Figure 4 shows the vendor-site share by engine, and Appendix B gives the counts behind Table 7.

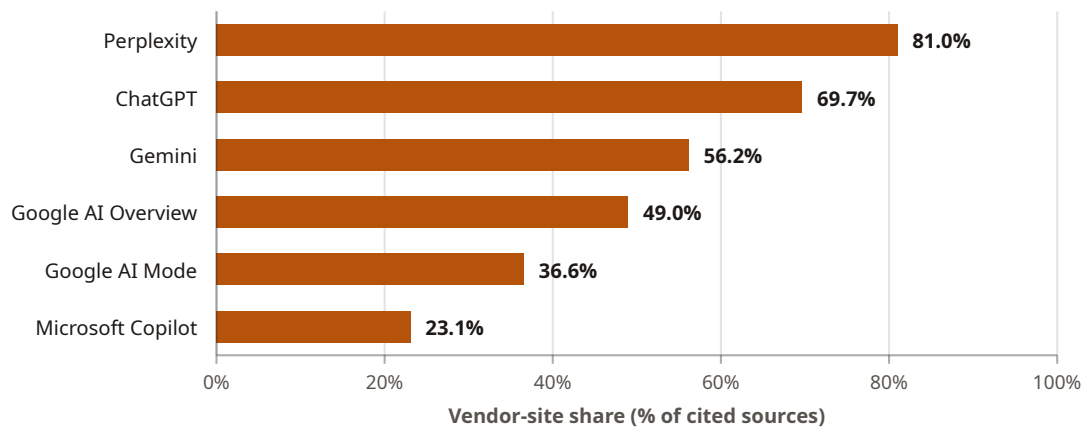


Figure 4. Share of cited sources on vendor sites by AI engine. Cited sources whose domain belongs to a vendor named in the response, as a share of that engine's cited sources; n = 1,172,000 cited sources. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

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 = 1,172,000 cited sources

AI engine	Cited sources	Dead-link rate	Dead links	Self-citation rate
ChatGPT	357,000	0.6%	2,140	0%
Microsoft Copilot	156,000	0%	0	0%
Gemini	112,000	0%	0	0%
Google AI Mode	175,000	0.6%	1,050	1%
Google AI Overview	351,000	0.3%	1,050	0%
Perplexity	21,000	0%	0	0%
Total	1,172,000	0.4%	4,240	—

The overall dead-link rate is 0.4%, 4,240 of 1,172,000 cited pages. 3 engines — Microsoft Copilot, Gemini and Perplexity — record no dead links at all, and no engine exceeds 0.6%, the rate for ChatGPT. ChatGPT accounts for 2,140 of the 4,240 dead links, 50.5% of the total.

The self-citation rate is 1% for Google AI Mode and 0% for the other five 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 = 1,172,000 cited sources

	ChatGPT	Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
ChatGPT	1.00	0.07	0.12	0.20	0.18	0.12

	ChatGPT	Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
Copilot	0.07	1.00	0.04	0.11	0.09	0.04
Gemini	0.12	0.04	1.00	0.28	0.27	0.20
Google AI Mode	0.20	0.11	0.28	1.00	0.38	0.14
Google AI Overview	0.18	0.09	0.27	0.38	1.00	0.15
Perplexity	0.12	0.04	0.20	0.14	0.15	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 0.16. The highest value is 0.38, between Google AI Mode and Google AI Overview. The lowest is 0.04, recorded twice, for Microsoft Copilot with Gemini and with Perplexity. The three engines operated by the same parent — Gemini, Google AI Mode and Google AI Overview — record coefficients of 0.28, 0.27 and 0.38 with each other, the three highest values in the matrix.

Microsoft Copilot records the lowest coefficient with every other engine: its five values run from 0.04 to 0.11, the latter with Google AI Mode. No pair of engines outside the same-parent group exceeds 0.20.

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 CyCognito and Palo Alto Networks.

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

AI engine	CyCognito	Palo Alto Networks	Microsoft	First-mention vendor
ChatGPT	4.5	1.4	3.3	Palo Alto Networks
Microsoft Copilot	2.9	1.7	3.3	Palo Alto Networks
Gemini	2.2	3.2	3.9	CyCognito
Google AI Mode	2.2	1.8	3.0	CyCognito
Google AI Overview	1.8	2.0	3.8	CyCognito
Perplexity	3.3	2.0	2.0	Microsoft

CyCognito holds the first-mention position in 3 of the 6 engines, the three operated by the same parent, on 160,000 mentions, the highest total in the corpus. Palo Alto Networks holds it in 2, ChatGPT and Microsoft Copilot, on 157,000. Microsoft holds it in 1, Perplexity.

Mean ordinal position does not follow the first-mention split. Palo Alto Networks records the earlier mean position in 4 engines, including Google AI Mode, where CyCognito is the first-mention vendor but Palo Alto Networks averages 1.8 against 2.2. CyCognito records the earlier position only in Gemini and Google AI Overview. In ChatGPT the gap is widest: Palo Alto Networks averages 1.4, the earliest mean position in Table 10, while CyCognito averages 4.5, the latest. Microsoft's mean position runs from 2.0 to 3.9.

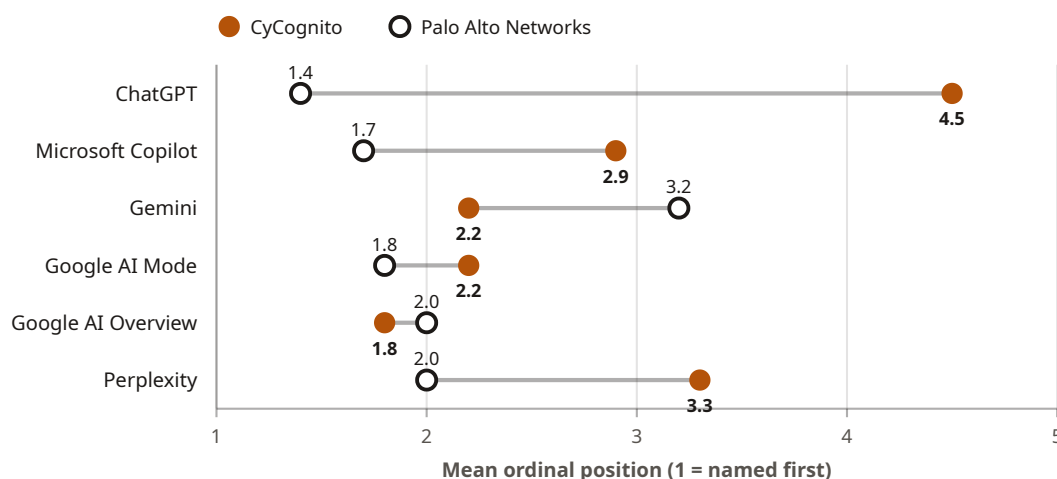


Figure 5. Mean ordinal position of CyCognito and Palo Alto Networks by AI engine. Lower is earlier in the response; n = 240,000 responses. Source: [GrackerAI AI Search Visibility Benchmark Series](#), August 2026.

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	CyCognito	Palo Alto Networks	Censys
Canada	CyCognito	Palo Alto Networks	Microsoft
India	CyCognito	Palo Alto Networks	Microsoft
Germany	Palo Alto Networks	CyCognito	Censys

The same two brands lead all four markets. CyCognito is first in the United States, India and Canada, with Palo Alto Networks second; Germany reverses the pair. The third position differs by market: Censys in the United States and Germany and Microsoft in India and Canada. No market returns a vendor absent from the other three in its leading positions.

At product level, 29 of the 124 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 discovery of unknown assets after a merger or acquisition. ChatGPT returns Cortex Xpanse as the product for that use case. Gemini and the two Google search surfaces more often lead with CyCognito, citing seedless attribution of subsidiary assets. The engines therefore return different products for the same question.

The second concerns integration. ChatGPT returns Tenable and Bitsight as the vendors to name for connectivity with ticketing, workflow and analytics tools. Google AI Overview returns IONIX and CyCognito for the same criterion.

The third concerns vendors that only one engine associates with the category. Microsoft Copilot alone returns 15 product entries, among them Qualys VMDR and ServiceNow Vulnerability Response. Perplexity alone returns 7, among them Bugcrowd Savant Vista and CATAAM. Google AI Mode alone returns 14 and Gemini alone 14. 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 small core and diverge in a very long tail. §4.2 shows 7 of the 87 recorded brand strings named by all six engines, and §4.3 shows 6 of 124 product entries reaching the same coverage. The difference between 7 brands and 6 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. 47 of the 87 brand strings and 81 of the 124 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 0.16.

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

Less concentrated at the top than in other categories in this series, and far more fragmented below it. §4.2 shows the two leading brands taking 25.5% of all mentions and the four leading brands 43.8%, against 87 brand strings recorded. §4.3 shows 102 of 124 product entries below 10,000 mentions and 69 at exactly 1,000.

The shape is a core of six or seven vendors the engines reuse, 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 0.16 over cited domains, with a maximum of 0.38 and a minimum of 0.04. The three highest values are the three same-parent pairs, no pair outside that group exceeds 0.20, and Microsoft Copilot records the lowest coefficient with every other engine. §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 and Google AI Overview together contribute 60.4% of all cited sources, so a source list that serves one of them 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?

No. §4.7 shows CyCognito recording the highest mention total and the first-mention position in three engines, while Palo Alto Networks records the earlier mean ordinal position in 4 of the six. In ChatGPT the brand named most often across the corpus averages 4.5, the latest position in Table 10, while Palo Alto Networks averages 1.4, the earliest.

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. Google AI Mode shows the second and third disagreeing within one engine, opening with CyCognito most often while placing Palo Alto Networks earlier on average. A plausible explanation is that Google AI Mode names CyCognito first in a majority of responses and much later in the remainder, which a mean captures and a first-mention count does not. §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 does, in one market. §4.8 shows CyCognito and Palo Alto Networks leading all four markets, with Germany the only market to reverse them, and a third position that differs by market within the same leading group. 29 of 124 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?

Mostly vendors' own sites, with one engine as the exception. §4.4 shows 51.3% of all cited sources on a vendor site, vendor sites as the largest class for 5 of the six engines, and ChatGPT drawing 69.7% of its citations from vendor domains. Microsoft Copilot is the opposite case at 23.1%, with 71.8% of its citations on domains outside every named class. Analyst and review sites are the second class, at 11.1% of the corpus and 20.0% of Google AI Mode's citations.

The data is consistent with two distinct evidence strategies behind one vendor list. Most engines read the vendors' own documentation and product pages, supplemented by analyst and review pages; one reads almost none of it. 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; ChatGPT's 97% should be read with that in mind. Sentiment is a model's reading of the response under a rubric, and the 0.15 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, and §4.7 shows the two disagreeing in one engine.

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 large for one engine. 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 reprocessing date. 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 attack surface management has a mature core and a very 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 small core with a long tail, the separation between mention volume and ordinal position, the elevated overlap between same-parent engines, the dominance of vendor sites in most engines' evidence bases, 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 attack surface 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 large for one engine; 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.

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 the response in three engines and averaging the latest position in another, and the second-placed brand averaging earliest where it does not open. 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, and analyst and review pages as the second. §4.4 shows 51.3% of all citations on vendor sites and 11.1% on analyst and review sites, with Microsoft Copilot the one engine that reads neither in quantity. Documentation, product and integration pages on the vendor's own domain are what most engines read.

Instrument every engine separately, and at least three from different vendor families. §4.6 gives a mean pairwise cited-domain overlap of 0.16, no cross-family pair 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 the Google surfaces read.

Name products consistently, and expect several entries per vendor regardless. §4.3 records 124 product entries, with Rapid7 holding 10. 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 two brands leading all four, with only the order changing in Germany. 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 0.4% and no engine above 0.6%. Do not treat sentiment as one either: the twelve leading brands sit within a 0.15 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 1,172,000 cited sources and covering 124 distinct product entries in enterprise attack surface management.

The engines agree about a small core and disagree about evidence. 7 of 87 brand strings were named by all six engines, led by CyCognito at 160,000 and Palo Alto Networks at 157,000, with 47 strings named by one engine only. Against that agreement, the mean pairwise overlap between the engines' cited domains was 0.16, the highest value anywhere in the matrix was 0.38, between two engines operated by the same parent, and no two engines shared a most-cited domain. What they do share is a class of source: vendor sites carry 51.3% of everything the engines cite.

Two further results qualify that list. Ordinal position does not track mention volume: CyCognito takes the first-mention position in three engines, yet Palo Alto Networks records the earlier mean position in 4 of six. In ChatGPT Palo Alto Networks is read first on average and CyCognito last of the three. Source reliability is high: 0.4% of cited pages were unreachable, and no engine exceeded 0.6%. The same two brands led every market, in an order that changed only in Germany.

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 that share a class of source rather than any one site. 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 reprocessed in August 2026 under the instrumentation's current citation and entity rules; the figures in this edition supersede those in the August 2026 edition.

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 attack surface 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 Attack Surface 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 87 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
CyCognito	160,000	26,000 (4.5)	32,000 (2.9)	31,000 (2.2)	32,000 (2.2)	36,000 (1.8)	3,000 (3.3)
Palo Alto Networks	157,000	34,000 (1.4)	29,000 (1.7)	35,000 (3.2)	32,000 (1.8)	26,000 (2.0)	1,000 (2.0)
Microsoft	115,000	29,000 (3.3)	27,000 (3.3)	26,000 (3.9)	17,000 (3.0)	14,000 (3.8)	2,000 (2.0)
Censys	113,000	30,000 (3.8)	16,000 (4.2)	28,000 (3.0)	23,000 (3.0)	16,000 (3.7)	—

Brand	Total	ChatGPT	Microsoft Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
Tenable	102,000	27,000 (2.8)	16,000 (3.5)	22,000 (5.0)	19,000 (3.9)	17,000 (4.2)	1,000 (3.0)
CrowdStrike	94,000	14,000 (4.2)	29,000 (2.8)	12,000 (4.3)	22,000 (4.0)	16,000 (4.5)	1,000 (1.0)
IONIX	75,000	12,000 (6.1)	11,000 (3.9)	11,000 (3.1)	14,000 (3.6)	26,000 (3.2)	1,000 (4.0)
Rapid7	42,000	13,000 (5.0)	10,000 (4.6)	3,000 (7.0)	8,000 (4.4)	8,000 (4.5)	—
Bitsight	35,000	6,000 (4.3)	—	19,000 (3.2)	8,000 (5.5)	1,000 (5.0)	1,000 (8.0)
Qualys	33,000	7,000 (7.1)	8,000 (3.4)	4,000 (4.3)	6,000 (5.2)	8,000 (4.4)	—
Wiz	28,000	4,000 (5.8)	9,000 (5.9)	10,000 (5.6)	—	4,000 (5.0)	1,000 (10.0)
UpGuard	25,000	—	10,000 (6.1)	3,000 (5.7)	7,000 (5.6)	5,000 (4.2)	—
Axonius	23,000	4,000 (7.3)	2,000 (8.0)	9,000 (3.3)	6,000 (3.7)	1,000 (1.0)	1,000 (5.0)
Detectify	20,000	4,000 (5.7)	14,000 (5.5)	—	1,000 (4.0)	1,000	—
runZero	16,000	11,000 (8.4)	1,000	1,000 (5.0)	—	2,000 (7.5)	1,000 (6.0)
Synack	15,000	—	—	9,000 (4.7)	—	5,000 (3.3)	1,000 (6.0)
Intruder	14,000	—	1,000	4,000 (3.0)	2,000 (5.0)	7,000 (3.4)	—
ServiceNow	14,000	1,000	10,000 (5.8)	1,000	2,000	—	—
JupiterOne	10,000	—	2,000 (7.0)	5,000 (5.0)	2,000 (5.0)	1,000 (5.0)	—
Shodan	9,000	—	—	9,000 (3.4)	—	—	—
Mandiant	7,000	2,000 (3.5)	3,000 (4.5)	2,000 (3.4)	—	—	—
Recorded Future	7,000	4,000 (4.5)	—	1,000 (3.0)	—	2,000 (4.0)	—
Armris	6,000	1,000 (11.0)	—	1,000 (6.0)	3,000 (4.3)	1,000 (2.0)	—
Bishop Fox	6,000	1,000 (2.0)	—	—	—	5,000 (5.0)	—
ImmuniWeb	6,000	2,000 (5.5)	3,000 (4.3)	—	—	—	1,000 (6.0)
Nucleus Security	6,000	—	—	4,000 (5.0)	—	2,000 (3.5)	—
HivePro	5,000	2,000 (11.0)	—	—	2,000 (5.5)	1,000 (2.0)	—
Splunk	5,000	1,000	1,000	1,000	2,000	—	—
XM Cyber	5,000	5,000 (4.0)	—	—	—	—	—
ThreatClaw	4,000	—	4,000 (6.5)	—	—	—	—
ZeroFox	4,000	—	—	—	—	2,000 (4.0)	2,000 (7.0)
FullHunt	3,000	3,000 (6.5)	—	—	—	—	—

Brand	Total	ChatGPT	Microsoft Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
Hadrian	3,000	—	3,000	—	—	—	—
IBM	3,000	1,000 (6.0)	2,000 (5.3)	—	—	—	—
Jira	3,000	—	—	1,000	2,000	—	—
SentinelOne	3,000	—	—	1,000	—	2,000 (5.5)	—
watchTowr	3,000	1,000	2,000	—	—	—	—
Astra Security	2,000	—	2,000 (6.5)	—	—	—	—
Brandefense	2,000	1,000	—	—	—	—	1,000 (2.0)
Bugcrowd	2,000	1,000	—	—	—	—	1,000 (1.0)
Check Point	2,000	—	—	—	2,000 (5.5)	—	—
Cyble	2,000	2,000 (7.0)	—	—	—	—	—
Exabeam	2,000	—	—	1,000	1,000	—	—
FireCompass	2,000	—	—	1,000	—	1,000 (4.0)	—
NordLayer	2,000	—	1,000 (6.0)	1,000	—	—	—
PlexTrac	2,000	—	—	1,000 (6.0)	—	1,000 (4.0)	—
Praetorian	2,000	—	—	2,000	—	—	—
SafeHill	2,000	—	2,000 (7.5)	—	—	—	—
SecurityScorecard	2,000	1,000	—	—	1,000 (8.0)	—	—
Strobes	2,000	—	—	1,000 (10.0)	—	—	1,000 (8.0)
Zafran	2,000	—	2,000 (6.0)	—	—	—	—
Atlassian	1,000	—	1,000	—	—	—	—
Bionic	1,000	—	—	1,000	—	—	—
Bitdefender	1,000	1,000 (7.0)	—	—	—	—	—
BreachLock	1,000	—	1,000 (6.0)	—	—	—	—
Bspeka	1,000	—	—	1,000	—	—	—
CATAAM	1,000	—	—	—	—	—	1,000 (3.0)
Cisco Vulnerability Management	1,000	1,000	—	—	—	—	—
CloudSEK	1,000	—	—	—	—	1,000	—
CybelAngel	1,000	—	—	1,000 (4.0)	—	—	—
Cymulate	1,000	—	—	—	1,000 (4.0)	—	—
Deepinfo	1,000	1,000 (6.0)	—	—	—	—	—
DefectDojo	1,000	—	—	1,000 (7.0)	—	—	—

Brand	Total	ChatGPT	Microsoft Copilot	Gemini	Google AI Mode	Google AI Overview	Perplexity
DeXpose	1,000	—	—	—	—	—	1,000 (5.0)
ExternalSight	1,000	—	1,000	—	—	—	—
FireMon	1,000	—	—	1,000	—	—	—
Forenzy	1,000	1,000	—	—	—	—	—
FortifyData	1,000	—	—	1,000 (11.0)	—	—	—
Gartner	1,000	—	—	—	1,000	—	—
Holm Security	1,000	—	1,000 (8.0)	—	—	—	—
Hunto AI	1,000	—	—	—	—	1,000 (5.0)	—
Lansweeper	1,000	—	—	—	1,000 (9.0)	—	—
Ordr	1,000	—	—	1,000 (2.0)	—	—	—
Outpost24	1,000	—	—	1,000 (7.0)	—	—	—
Panaseer	1,000	1,000	—	—	—	—	—
Patrowl	1,000	—	—	—	—	—	1,000 (7.0)
Project Discovery	1,000	—	—	1,000	—	—	—
RiskRecon	1,000	—	—	1,000 (8.0)	—	—	—
Searchlight Cyber	1,000	1,000 (5.0)	—	—	—	—	—
Sectora	1,000	—	1,000 (6.0)	—	—	—	—
SOCRadar	1,000	—	1,000 (7.0)	—	—	—	—
Trend Micro	1,000	—	—	—	—	1,000 (4.0)	—
ULTRA RED	1,000	—	—	—	—	—	1,000 (7.0)
UltraRed	1,000	—	—	—	—	1,000	—
Veracode	1,000	—	—	1,000 (6.0)	—	—	—
Vulcan Cyber	1,000	—	1,000	—	—	—	—
Zscaler	1,000	—	—	1,000	—	—	—
Total	1,243,000	256,000	259,000	272,000	217,000	215,000	24,000

Appendix B. Most-cited domains and source types

Table B1. Three most-cited domains by AI engine — cited-source counts in parentheses; n = 1,172,000 cited sources

AI engine	Cited sources	Distinct domains	Most-cited domains
ChatGPT	357,000	56	paloaltonetworks.com (59,000), gartner.com (42,000), tenable.com (30,000)

AI engine	Cited sources	Distinct domains	Most-cited domains
Microsoft Copilot	156,000	33	uinat.com (24,000), eastbaycyber.com (17,000), ionix.io (13,000)
Gemini	112,000	36	synack.com (17,000), bitsight.com (14,000), guideflow.com (11,000)
Google AI Mode	175,000	47	gartner.com (33,000), ionix.io (19,000), cycognito.com (14,000)
Google AI Overview	351,000	73	cycognito.com (53,000), ionix.io (30,000), paloaltonetworks.com (27,000)
Perplexity	21,000	19	brandefense.io (2,000), zerofox.com (2,000), ultrared.ai (1,000)

Table B2. Cited sources by source type and AI engine — counts; n = 1,172,000 cited sources

AI engine	Vendor site	Other	Review aggregator	Social	Forum	Parent ecosystem	Marketplace	Wikipedia	Total
ChatGPT	249,000	46,000	59,000	1,000	0	0	1,000	1,000	357,000
Microsoft Copilot	36,000	112,000	8,000	0	0	0	0	0	156,000
Gemini	63,000	41,000	8,000	0	0	0	0	0	112,000
Google AI Mode	64,000	62,000	35,000	12,000	1,000	1,000	0	0	175,000
Google AI Overview	172,000	140,000	19,000	15,000	3,000	1,000	1,000	0	351,000
Perplexity	17,000	3,000	1,000	0	0	0	0	0	21,000
Total	601,000	404,000	130,000	28,000	4,000	2,000	2,000	1,000	1,172,000

Appendix C. Product entries

Table C1. All product entries — 124 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	CyCognito	CyCognito	160,000	2.6	0.82	6	4
2	Cortex Xpanse	Palo Alto Networks	149,000	1.8	0.86	6	4
3	Microsoft Defender EASM	Microsoft	112,000	3.4	0.78	6	4
4	Censys	Censys	107,000	3.5	0.81	5	4
5	IONIX	IONIX	73,000	3.9	0.82	6	4
6	CrowdStrike Falcon Exposure Management	CrowdStrike	52,000	3.1	0.81	6	4
7	Tenable One	Tenable	51,000	3.1	0.80	5	4
8	CrowdStrike Falcon Surface	CrowdStrike	39,000	4.6	0.81	5	4
9	Tenable Attack Surface Management	Tenable	30,000	4.7	0.81	6	4

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
10	Mandiant ASM	Google Cloud	28,000	4.1	0.78	5	4
11	Bitsight	Bitsight	27,000	4.2	0.82	4	4
12	Wiz	Wiz	22,000	6.3	0.81	5	4
13	Qualys EASM	Qualys	18,000	5.7	0.76	4	4
14	UpGuard	UpGuard	17,000	5.1	0.72	4	4
15	Detectify	Detectify	15,000	5.4	0.72	3	4
16	runZero	runZero	15,000	7.9	0.74	4	4
17	Synack	Synack	14,000	4.4	0.84	3	4
18	Rapid7 Surface Command	Rapid7	12,000	4.6	0.78	5	3
19	Axonius	Axonius	12,000	5.8	0.86	4	4
20	Intruder	Intruder	11,000	3.5	0.80	3	4
21	Tenable One / Tenable ASM	Tenable	10,000	4.0	0.84	3	4
22	Rapid7 InsightVM	Rapid7	10,000	4.3	0.64	3	4
23	Axonius Asset Cloud	Axonius	8,000	2.0	0.91	4	4
24	Shodan Enterprise	Shodan	8,000	3.4	0.75	1	4
25	Qualys VMDR	Qualys	8,000	3.4	0.63	1	4
26	JupiterOne	JupiterOne	8,000	5.3	0.80	4	4
27	Rapid7 Exposure Command	Rapid7	8,000	5.4	0.78	2	3
28	ServiceNow Vulnerability Response	ServiceNow	8,000	5.4	0.76	1	4
29	Randori	IBM	7,000	5.4	0.85	2	4
30	Bishop Fox Cosmos	Bishop Fox	6,000	4.5	0.83	2	3
31	Mandiant Advantage Attack Surface Management	Mandiant	6,000	4.5	0.77	4	3
32	Recorded Future Attack Surface Intelligence	Recorded Future	6,000	4.3	0.70	2	3
33	Armis Centrix	Armis	6,000	5.3	0.83	4	3
34	XM Cyber	XM Cyber	5,000	4.0	0.65	1	4
35	ImmuniWeb Discovery	ImmuniWeb	5,000	4.8	0.74	2	3
36	HivePro Uni5 Xposure	HivePro	5,000	7.0	0.86	3	4
37	Bitsight Attack Surface Intelligence	Bitsight	4,000	3.0	0.82	2	2
38	Nucleus Security	Nucleus Security	4,000	4.5	0.80	2	3
39	ThreatClaw	ThreatClaw	4,000	6.5	0.85	1	3
40	Wiz Exposure Management	Wiz	3,000	2.7	0.87	2	1

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
41	Palo Alto Networks Expanse	Palo Alto Networks	3,000	3.0	0.80	2	3
42	Cybersecurity Asset Management / ASM	Qualys	3,000	4.3	0.80	2	3
43	Tenable Exposure Platform	Tenable	3,000	4.3	0.80	1	3
44	UpGuard Breach Risk / ASM	UpGuard	3,000	6.0	0.80	1	3
45	UpGuard BreachSight	UpGuard	3,000	7.0	0.80	2	2
46	Brandefense	Brandefense	2,000	2.0	0.65	2	2
47	PlexTrac	PlexTrac	2,000	5.0	0.75	2	2
48	Rapid7	Rapid7	2,000	5.0	0.60	1	2
49	FullHunt	FullHunt	2,000	6.5	0.85	1	2
50	Astra Security	Astra	2,000	6.5	0.70	1	2
51	ZeroFox Attack Surface Intelligence	ZeroFox	2,000	7.0	0.80	1	1
52	Tenable Exposure Graph	Tenable	2,000	7.0	0.75	2	1
53	Xpanse	Palo Alto Networks	2,000	7.5	0.75	2	2
54	SafeHill SecureIQ	SafeHill	2,000	7.5	0.60	1	1
55	Strobes ASM	Strobes	2,000	9.0	0.75	2	2
56	Cortex Xpanse / Falcon Surface	Palo Alto Networks	1,000	1.0	0.90	1	1
57	Cortex XSIAM / Expanse	Palo Alto Networks	1,000	1.0	0.80	1	1
58	Bugcrowd Savant Vista	Bugcrowd	1,000	1.0	0.80	1	1
59	Forenzy Atlas	Forenzy	1,000	—	0.50	1	1
60	Bugcrowd ASM offerings	Bugcrowd	1,000	—	0.40	1	1
61	Tenable ASM / Tenable One	Tenable	1,000	2.0	0.90	1	1
62	Ordr	Ordr	1,000	2.0	0.85	1	1
63	Rapid7 Exposure Command / InsightVM	Rapid7	1,000	2.0	0.80	1	1
64	IONIX Continuous Threat Exposure Management	IONIX	1,000	2.0	0.80	1	1
65	Mandiant Advantage ASM	Google Cloud	1,000	2.0	0.80	1	1
66	Tenable ASM / Tenable CSAM	Tenable	1,000	2.0	0.80	1	1
67	Rapid7 Surface Command / InsightVM	Rapid7	1,000	3.0	1.00	1	1
68	Censys Search & ASM	Censys	1,000	3.0	0.80	1	1
69	Exposure Command / InsightVM ecosystem	Rapid7	1,000	3.0	0.80	1	1
70	CATAAM	CATAAM	1,000	3.0	0.80	1	1
71	Recorded Future ASM	Recorded Future	1,000	3.0	0.70	1	1

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
72	Microsoft Security Exposure Management / Defender EASM	Microsoft	1,000	3.0	0.70	1	1
73	Qualys CyberAsset Asset Management (CAAM)	Qualys	1,000	3.0	0.70	1	1
74	Nucleus Security / Brinqa	Nucleus Security / Brinqa	1,000	4.0	1.00	1	1
75	Trend Vision One	Trend Micro	1,000	4.0	0.90	1	1
76	Tenable One Exposure Management	Tenable	1,000	4.0	0.80	1	1
77	CrowdStrike Falcon Recon	CrowdStrike	1,000	4.0	0.80	1	1
78	Qualys Cybersecurity Asset Management & VMDR	Qualys	1,000	4.0	0.80	1	1
79	FireCompass	FireCompass	1,000	4.0	0.80	1	1
80	ZeroFox	ZeroFox	1,000	4.0	0.80	1	1
81	Cymulate Exposure Management Platform	Cymulate	1,000	4.0	0.70	1	1
82	Palo Alto Networks Cortex Xpanse / Exposure Management	Palo Alto Networks	1,000	4.0	0.70	1	1
83	CybelAngel	CybelAngel	1,000	4.0	0.70	1	1
84	Hunto AI ASM	Hunto AI	1,000	5.0	0.90	1	1
85	Palo Alto Networks Prisma Cloud CDEM	Palo Alto Networks	1,000	5.0	0.90	1	1
86	External Attack Surface Monitoring	Rapid7	1,000	5.0	0.80	1	1
87	Intruder / RiskProfiler	Intruder / RiskProfiler	1,000	5.0	0.80	1	1
88	SentinelOne ASM	SentinelOne	1,000	5.0	0.80	1	1
89	Wiz Security Graph	Wiz	1,000	5.0	0.80	1	1
90	Exposure Management / Cyberint	Check Point	1,000	5.0	0.80	1	1
91	Assetnote	Searchlight Cyber	1,000	5.0	0.80	1	1
92	DeXpose	DeXpose	1,000	5.0	0.80	1	1
93	Bitsight Exposure Management	Bitsight	1,000	5.0	0.80	1	1
94	BreachLock ASM	BreachLock	1,000	6.0	0.90	1	1
95	Axonius Cyber Asset Attack Surface Management	Axonius	1,000	6.0	0.90	1	1
96	Sectora ASM	Sectora	1,000	6.0	0.90	1	1
97	NordLayer ASM 2.0	NordLayer	1,000	6.0	0.90	1	1
98	Check Point Exposure Management	Check Point	1,000	6.0	0.80	1	1
99	Veracode EASM	Veracode	1,000	6.0	0.80	1	1
100	SentinelOne Singularity	SentinelOne	1,000	6.0	0.80	1	1

#	Product	Brand	Mentions	Mean rank	Mean sentiment	Engines	Markets
101	ImmuniWeb ASM	ImmuniWeb	1,000	6.0	0.80	1	1
102	Bitsight Attack Surface Analytics	Bitsight	1,000	6.0	0.70	1	1
103	Deepinfo EASM	Deepinfo	1,000	6.0	0.70	1	1
104	Zafran	Zafran	1,000	6.0	0.60	1	1
105	GravityZone EASM	Bitdefender	1,000	7.0	0.80	1	1
106	ULTRA RED EASM	ULTRA RED	1,000	7.0	0.80	1	1
107	DefectDojo	DefectDojo	1,000	7.0	0.80	1	1
108	Patrowl EASM	Patrowl	1,000	7.0	0.80	1	1
109	Cybersprint	Outpost24	1,000	7.0	0.70	1	1
110	watchTowr / Hadrian / Detectify	Various	1,000	7.0	0.70	1	1
111	Cyble ASM	Cyble	1,000	7.0	0.70	1	1
112	SOCRadar ASM	SOCRadar	1,000	7.0	0.70	1	1
113	SecurityScorecard EASM	SecurityScorecard	1,000	8.0	0.90	1	1
114	Rapid7 Surface Command / Exposure Command	Rapid7	1,000	8.0	0.80	1	1
115	Holm Security	Holm Security	1,000	8.0	0.80	1	1
116	Qualys Enterprise TruRisk / EASM	Qualys	1,000	7.0	0.50	1	1
117	Google Cloud Security by Mandiant	Google Cloud	1,000	8.0	0.70	1	1
118	RiskRecon EASM	RiskRecon	1,000	8.0	0.70	1	1
119	Prisma Cloud	Palo Alto Networks	1,000	9.0	0.80	1	1
120	Lansweeper Cyber Asset Intelligence	Lansweeper	1,000	9.0	0.80	1	1
121	Palo Alto Networks Cortex / Prisma Cloud	Palo Alto Networks	1,000	9.0	0.80	1	1
122	ServiceNow Security Operations	ServiceNow	1,000	9.0	0.70	1	1
123	Rapid7 InsightVM / Surface Command	Rapid7	1,000	9.0	0.70	1	1
124	FortifyData ASM	FortifyData	1,000	11.0	0.70	1	1

Appendix D. Study parameters

Table D1. Study parameters — as configured for this study

Parameter	Value
Category	Enterprise attack surface management (ASM, EASM, CAASM)
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	1,172,000
Brand strings recorded	87
Product entries	124
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.