

INDUSTRY RESEARCH • SEPTEMBER 2026 BASELINE

Medicare Visibility Report

How AI Answer Engines Understand, Source, and Attribute Medicare Plan Information

~2,100

CMS Plan-ID entities

39

standardized plan facts

2

assertion engines analyzed

A baseline measurement of Plan-ID × engine × fact relationships, designed to separate search discovery, assertion production, provenance, and publisher attribution.

Published by Trust Publishing Institute • Research and analysis by David W. Bynon

Prepared from the September 2026 Medicare Visibility Monitor capture

Baseline edition • Results are observational, not causal

Executive Summary

AI systems increasingly answer Medicare questions directly rather than simply returning webpages. This report measures that environment at the level of the individual Medicare plan and the individual factual assertion, rather than treating rankings or citation counts as sufficient proxies for answer visibility.

54% Google AI • knowledge completeness	94% Google AI • assertions with attribution	68% Copilot • knowledge completeness
50% Copilot • assertions with attribution	188 Google AI • observed source domains	74 Copilot • observed source domains

Core finding: search engines discover documents; answer engines assemble assertions. The sources of those assertions differ by engine, entity, and fact.

Study at a Glance

Component	September 2026 baseline
Entity population	Approximately 2,100 CMS Medicare Plan-ID entities
Usable observations	2,141 Google AI; 2,146 Microsoft Copilot
Knowledge model	39 predefined Medicare plan fact categories
Primary unit	Plan-ID × engine × fact
Evidence relationship	Plan-ID → assertion → engine → attributed source
Engine scope	Observed Google AI and Microsoft Copilot surfaces in the September capture; not all Google/Microsoft AI products
Query protocol	Standardized Plan-ID-specific measurement protocol
Attribution coding	Explicit source associations exposed with an assertion; one assertion may map to multiple source relationships
Search comparison	Google/Bing discovery observations compared with assertion evidence where available
Auditability	Assertion-level capture records and provenance retained
Known limitation	Cross-engine semantic agreement withheld pending fact-specific normalization

How to Interpret This Report

- An assertion is an observed engine statement about a Plan-ID and fact category.
- Attribution is an exposed association between an assertion and one or more sources.
- Attribution does not prove that a fact is accurate, current, original, or exclusive to that source.
- Absence of exposed attribution does not prove that an engine used no source.
- This baseline measures observed behavior under a defined protocol, not consumer exposure or market share.

Research Independence and Disclosure

David W. Bynon manages publishing technology and editorial systems for Medicare.org, a publisher included in this measurement. Medicare.org is owned and operated by Health Network Group LLC, an Allstate company. This

report was produced by Trust Publishing Institute and MedicarePlans.com. Neither Medicare.org nor Allstate commissioned the report or determined its methodology, findings, or conclusions. All observed publishers were evaluated under the same measurement and aggregation rules, and underlying capture records and assertion-level provenance are retained for auditability.

Report Map

- Finding 1 — Google and Copilot Construct Medicare Answers Differently
- Finding 2 — Different Facts Have Radically Different Visibility
- Finding 3 — Medicare AI Evidence Is Supplied by a Relatively Small Publisher Ecosystem
- Finding 4 — Publisher Leadership Depends on the Fact Being Answered
- Finding 5 — Being the Leading Source Is Different From Controlling Attribution
- Finding 6 — Search Visibility and Answer Visibility Are Related, but Not Interchangeable
- Finding 7 — Q1Medicare’s Exit Creates a Measurable Change in the Ecosystem
- What This Report Does Not Measure
- Methodology
- What Stakeholders Can Learn From the Baseline

AI systems increasingly answer Medicare questions directly rather than simply returning a list of webpages for consumers to evaluate themselves.

That changes what it means for Medicare information to be visible.

In conventional search, visibility has largely been measured at the document level: whether a page was indexed, where it ranked, how often it appeared, and whether a user clicked it. Those measurements remain important, but they do not fully describe what happens when an answer engine constructs a response from information gathered across multiple sources.

A publisher can rank prominently in conventional search yet contribute little to an AI-generated answer. Another publisher can be difficult to find in conventional results while repeatedly supplying facts used by an answer engine. Even within the same answer, different publishers may supply evidence for different facts about the same Medicare plan.

The **September 2026 Medicare Visibility Report** establishes a baseline for measuring this emerging information environment below the level of the webpage and citation. It examines visibility at the level of the individual Medicare plan and the individual factual assertion made about that plan.

The baseline includes approximately **2,100 Medicare Plan-ID entities** and **39 standardized plan-information categories**, including plan identity, premiums, deductibles, maximum out-of-pocket limits, eligibility requirements, cost sharing, supplemental benefits, service areas, ratings, and member contact information.

For each Plan-ID, the measurement examines what an answer engine says, what it leaves unsaid, whether its assertions are explicitly attributed to sources, and which publishers supply that evidence.

Rather than asking only whether a publisher appeared in an AI response or received a citation, this report asks a different set of questions:

- **What does the answer engine assert about a Medicare plan?**
- **How much of the plan does it appear to understand?**
- **Which assertions receive explicit source attribution?**
- **Which publishers supply that evidence?**
- **How does conventional search discovery relate to assertion-level evidence?**

These distinctions matter because a citation alone does not tell us what information a source contributed. A publisher cited once may supply evidence for many individual assertions. Multiple publishers may support the same assertion. And an answer engine may state a fact without exposing any source attribution at all.

For that reason, this report does not treat rankings, citations, retrieval, assertions, or attribution as interchangeable measures of AI visibility. Each describes a different part of the process by which information becomes an answer.

The September 2026 results provide the first baseline for observing that process across a large population of Medicare plans.

Finding 1: Google and Copilot Construct Medicare Answers Differently

The first significant finding is not about which publisher receives the most attribution. It is about the answer engines themselves.

Google AI and Microsoft Copilot demonstrated substantially different patterns in both the amount of Medicare plan information they asserted and the extent to which those assertions were accompanied by identifiable source attribution.

Across the measured Plan-ID population, Google AI asserted approximately **54% of the 39 plan facts examined**. Copilot asserted approximately **68%**.

In other words, Copilot generally produced a more complete representation of the Medicare plan knowledge model used in this study.

Source attribution moved sharply in the opposite direction.

When Google made an assertion about a plan, approximately **94% of those assertions were associated with explicit source attribution**. For Copilot, approximately **50% of asserted facts were attributed to a source**.

Answer Engine	Knowledge Completeness	Assertions With Source Attribution
Google AI	54%	94%
Microsoft Copilot	68%	50%

Knowledge Completeness vs. Knowledge Provenance

The difference can be understood as two separate dimensions of an AI-generated answer.

In this report, “knowledge completeness” is an output-coverage metric: the share of the predefined 39-fact model that the engine explicitly asserted in the observed response. It does not measure internal model knowledge, factual accuracy, or consumer-facing response frequency. “Knowledge provenance” measures how much of that asserted information is accompanied by identifiable source attribution.

- **Google AI:** Lower knowledge completeness, very high provenance.
- **Microsoft Copilot:** Higher knowledge completeness, substantially lower provenance.

This distinction is important because an answer containing more information is not necessarily an answer containing more attributable information. Conversely, an engine that states fewer facts may expose substantially more evidence for the facts it does state.

For example, an answer engine could identify a plan's premium, deductible, maximum out-of-pocket limit, eligibility requirements and several supplemental benefits while providing sources for only some of those assertions. Another engine could provide fewer facts while associating nearly every statement with identifiable evidence.

Both answers may appear informative to a consumer, but they represent very different information environments from the perspective of measurement, verification and publishing.

The September baseline therefore suggests that **AI visibility cannot be adequately represented by citation counts alone**. Citation volume describes only one aspect of answer construction. Measuring the information environment also requires knowing what an engine asserts, what it omits, and how consistently the information it does provide is connected to identifiable evidence.

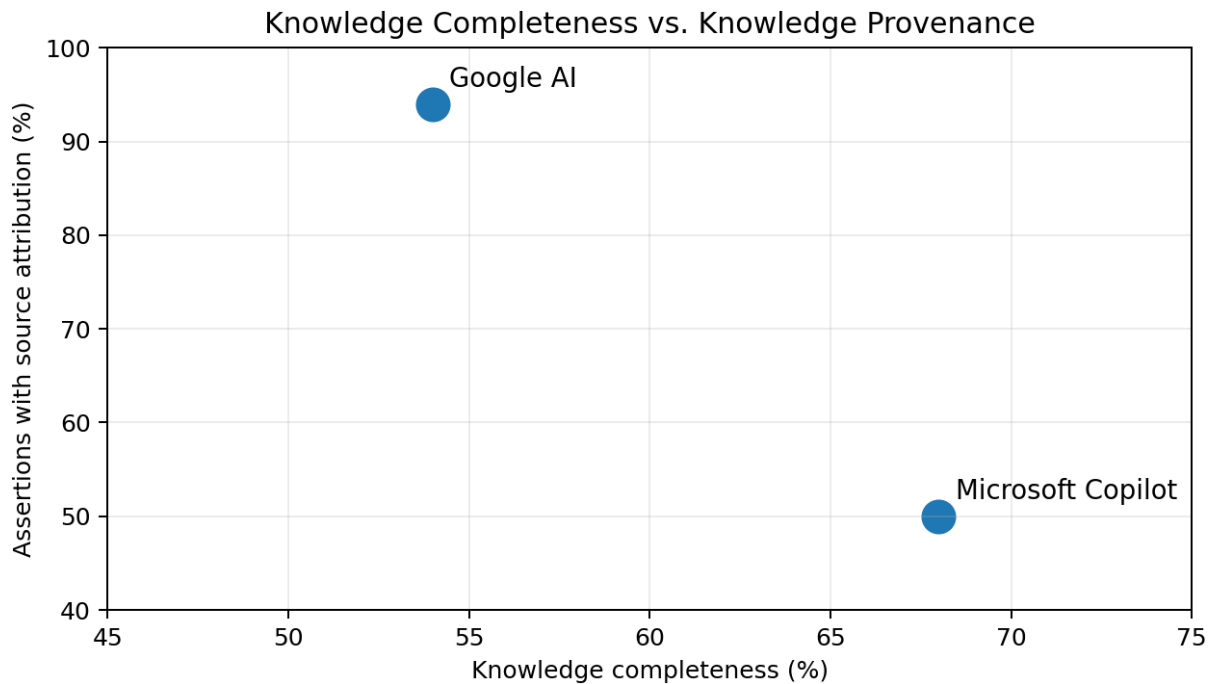


Figure 1. Knowledge completeness and provenance diverge sharply by engine.

Finding 2: Different Facts Have Radically Different Visibility

Overall knowledge completeness tells us how much an answer engine asserts about a Medicare plan. It does not tell us *which* information the engine is likely to provide.

At the individual fact level, the September baseline shows substantial differences in assertion behavior.

For example, Microsoft Copilot asserted an **emergency-room copay for 90.8% of measured Plan-IDs**. Google AI did so for only **24.8%**.

The same pattern appeared in several other important categories. Copilot asserted **Medicare eligibility for 96.9%** of measured plans, compared with **38.2%** for Google. It asserted **residency requirements for 89.8%**, compared with **36.9%** for Google.

Plan Fact	Google AI Assertion Rate	Copilot Assertion Rate
Emergency-Room Copay	24.8%	90.8%
Medicare Eligibility	38.2%	96.9%
Residency Requirement	36.9%	89.8%

These differences are not simply the result of Copilot asserting more information overall. The direction reverses for some fact categories. Google AI asserted **service-area information approximately 24 percentage points more often than Copilot** and **insulin-cost information approximately 17 percentage points more often**.

The result is an uneven knowledge profile rather than a simple distinction between a more complete and a less complete answer engine.

These differences should not be read as proof that one engine has better underlying Medicare data. They show that, under the study protocol, the two observed systems exposed different parts of the same predefined fact model with materially different frequency. Possible causes include answer format, retrieval behavior, source availability, extraction behavior, or system-level response policy; this baseline does not distinguish among them.

Medicare Knowledge Completeness Is Fact-Specific

There is no single Medicare knowledge-completeness characteristic that describes an answer engine equally well across all types of plan information.

An engine may consistently identify a plan's premium, deductible and Star Rating while frequently omitting eligibility requirements, cost-sharing details or supplemental benefits. Another engine may routinely assert those facts but provide substantially less source attribution for them.

This matters because the information consumers need to evaluate a Medicare plan is distributed across many different factual categories. Knowing a plan's monthly premium does not establish its maximum out-of-pocket exposure. Identifying a supplemental benefit does not establish the conditions under which that benefit is available. Identifying a Special Needs Plan does not necessarily communicate who is eligible to enroll.

For publishers and Medicare carriers, the implication is practical: **AI visibility should be measured at the level of individual facts, not merely at the level of pages, plans or domains.**

A Plan-ID can be highly visible while important parts of its knowledge representation remain consistently absent from an answer engine. Measuring assertion rates by fact category makes those gaps observable and creates a specific target for subsequent investigation and remediation.

Finding 3: Medicare AI Evidence Is Supplied by a Relatively Small Publisher Ecosystem

The September baseline identified a substantial number of domains supplying evidence for Medicare plan assertions, but attribution was not distributed evenly across them.

Google AI exposed attribution from 188 unique observed source domains within the measured Plan-ID population. Microsoft Copilot exposed attribution from 74.

Despite Google's broader source ecosystem, a relatively small group of publishers accounted for most observed attribution on both engines.

Attribution Concentration	Google AI	Microsoft Copilot
Largest Source	~35%	~60%
Top 3 Sources	~70%	~89%
Top 10 Sources	~91%	~99%

The difference between the two engines is notable. Google's evidence ecosystem was broader and attribution was distributed among a larger number of sources. Copilot relied on fewer observed domains, with attribution substantially more concentrated among its leading sources.

Domain diversity is not the same as evidence diversity. A larger number of observed source domains can coexist with substantial concentration among a small number of repeatedly attributed publishers.

Leading Sources of Medicare Plan Assertions

On Google AI, the three largest measured sources of attributed assertions were:

Publisher	Attributed Assertions
Medicare.org	34,607
Q1Medicare	20,911
MedicareAdvantage.com	14,225

On Microsoft Copilot, the three largest measured sources were:

Publisher	Attributed Assertions
Medicare.org	22,816
MedicareAdvantage.com	7,710
MedicarePlans.com	3,234

Medicare.org was the largest measured source of attributed assertions on both engines. This report does not interpret that result as a measure of information quality, accuracy, neutrality or consumer preference. It measures observed source attribution within the study population.

The distinction is important. An answer engine's decision to associate a source with an assertion establishes an observable evidence relationship. It does not, by itself, establish that the assertion is correct, that the source is the original source of the information, or that the publisher is preferable to another source.

Source Diversity Differs Sharply by Engine

The concentration data also shows why counting the number of domains cited by an AI system can be misleading.

Google drew evidence from more than twice as many observed domains as Copilot, yet its top three sources still accounted for approximately **70% of measured attribution**. Its top 10 accounted for approximately **91%**.

Copilot was substantially more concentrated. Its top three sources accounted for approximately **89% of measured attribution**, and its top 10 accounted for approximately **99%**.

Counting note: Source-concentration results are based on observed attribution relationships, not unique answer pages or unique citations. A single Plan-ID × fact assertion may be associated with more than one source and can therefore contribute multiple attribution relationships.

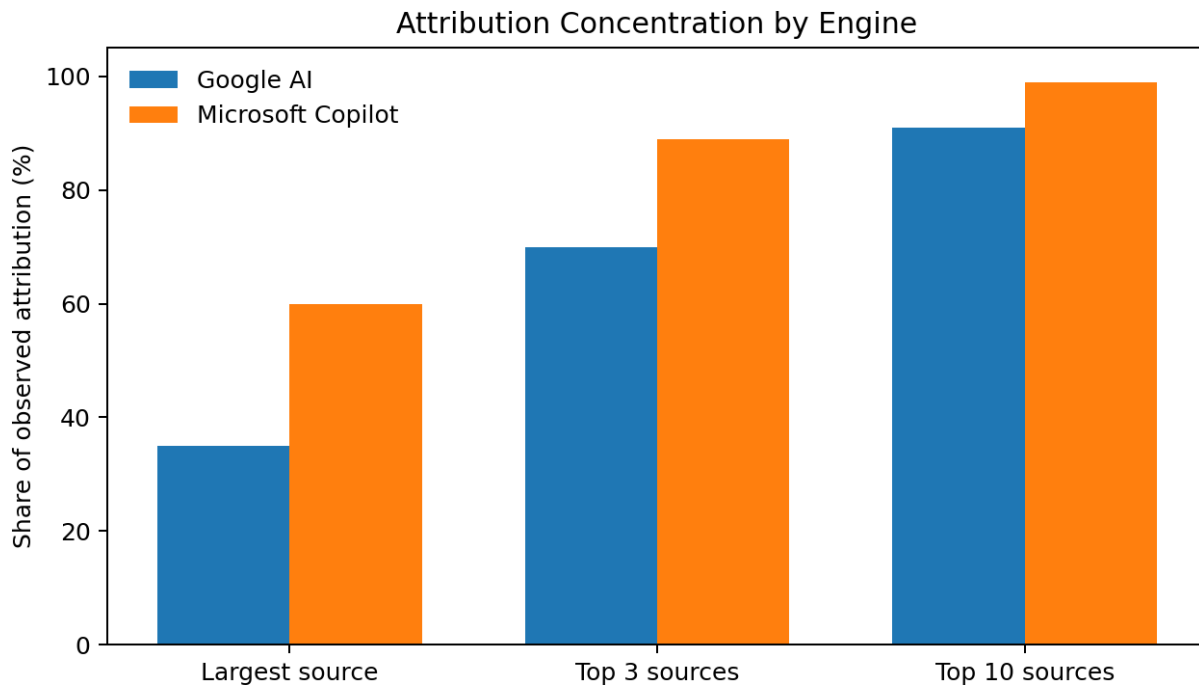


Figure 2. A small number of publishers account for most observed attribution, especially on Copilot.

This suggests that the Medicare answer ecosystem has a long tail of available sources but a much smaller group of publishers that answer engines repeatedly use as evidence across individual Medicare plan entities.

It also reinforces the need to distinguish between **being available to an answer engine** and **being repeatedly selected as evidence**. Hundreds of domains may participate in the information environment, while a relatively small number supply a disproportionate share of the assertions ultimately associated with generated answers.

Finding 4: Publisher Leadership Depends on the Fact Being Answered

The concentration of Medicare evidence among a relatively small number of publishers does not mean that a single publisher supplies every type of information equally.

When attribution is examined at the level of the individual fact, the competitive landscape becomes considerably more heterogeneous.

Across the 39 plan-information categories measured on Google AI, **Medicare.org was the leading attributed publisher for 30 facts, Q1Medicare led 8, and MedicareAdvantage.com led 1.**

Microsoft Copilot produced a different distribution. **Medicare.org led 32 of the 39 fact categories, MedicareAdvantage.com led 5, and UnitedHealthcare led 2.**

Answer Engine	Publisher	Fact Categories Led
Google AI	Medicare.org	30 of 39
Q1Medicare	8 of 39	
MedicareAdvantage.com	1 of 39	
Microsoft Copilot	Medicare.org	32 of 39
MedicareAdvantage.com	5 of 39	
UnitedHealthcare	2 of 39	

Different Publishers Lead Different Parts of the Plan Knowledge Model

Google's attribution patterns provide a clear example.

Q1Medicare was the leading measured source for several categories involving eligibility, prescription drug information and plan contact information. These included **dual eligibility, Medicaid eligibility, prescription drug deductible, target conditions, insulin cost, existing-member phone, new-member phone, and official website.**

The pattern was different on Copilot. MedicareAdvantage.com led several categories involving supplemental benefits, including **dental, hearing, vision, over-the-counter and healthy-food benefits.** UnitedHealthcare led attribution for **insulin cost and meal benefits.**

These differences indicate that publisher competition inside an AI-generated answer can occur at a much finer level than the webpage or domain.

Publisher visibility is not monolithic.

An answer engine can rely primarily on one publisher for premiums and deductibles, another for supplemental benefits, another for eligibility information, and another for contact information. Those evidence relationships can coexist within the answer engine's understanding of the same Medicare plan.

The Assertion Is a Distinct Unit of Visibility

This creates an important measurement distinction.

If an AI system cites a publisher while describing a Medicare plan, a citation count establishes that the publisher participated in the answer. It does not establish **which facts the publisher supported**.

Likewise, two publishers receiving similar numbers of citations may play very different roles. One may repeatedly supply core plan costs across thousands of entities, while another may be relied upon primarily for eligibility requirements, supplemental benefits or contact information.

Measuring only whether a page or domain received a citation therefore obscures the assertion-level competition occurring underneath the generated answer.

The September baseline suggests that a more useful model of publisher visibility requires identifying the relationship between the **entity being described, the fact being asserted, the answer engine making the assertion, and the publisher supplying the associated evidence**.

At that level, the question is no longer simply, *"Which publisher was cited?"*

It becomes, **"Which publisher does the answer engine rely on for this fact about this plan?"**

Finding 5: Being the Leading Source Is Different From Controlling Attribution

Identifying the leading publisher for a particular Medicare fact answers one question: which publisher received the most attribution?

It does not tell us how much of the attribution that publisher actually received.

This distinction becomes important when comparing Google AI and Microsoft Copilot. A publisher can rank first for a fact in both environments while occupying a very different position within each engine's evidence ecosystem.

On Google AI, the leading publisher for a fact may account for approximately 30% to 40% of measured attribution, leaving the majority distributed among competing sources. On Copilot, the leading publisher can account for more than three-quarters of the attribution for the same or similar types of plan information.

These represent fundamentally different levels of evidence concentration even though the same publisher may rank first in both cases.

Fact Leadership vs. Share of Attribution

This report therefore distinguishes between two measurements:

Fact leadership identifies the most frequently attributed publisher for a factual category.

Share of attribution measures how much of the observed evidence attribution for that fact belongs to each publisher.

Copilot provides several clear examples of highly concentrated attribution. Medicare.org received approximately:

Plan Fact	Medicare.org Share of Attribution on Copilot
Star Rating	81%
Part B Giveback	80%

Plan Fact	Medicare.org Share of Attribution on Copilot
Medical Deductible	79%
Maximum Out-of-Pocket Limit	78%
Monthly Premium	77%

In these categories, Medicare.org was not simply the leading measured source. Copilot associated a substantial majority of the observed attribution with a single publisher.

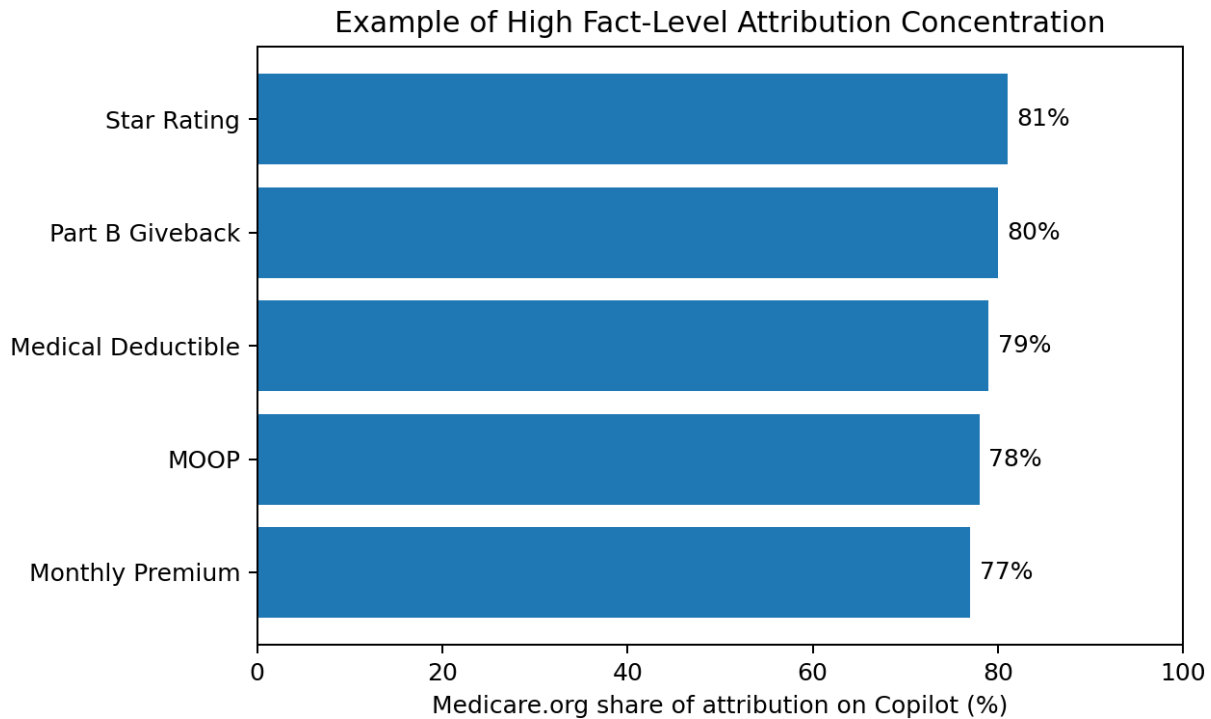


Figure 3. Leading a fact category and controlling most attribution are different measurements.

That pattern did not extend uniformly across the plan knowledge model.

For several supplemental-benefit categories, Medicare.org's attribution share was substantially lower and other publishers led the category. MedicareAdvantage.com, for example, led Copilot attribution for **dental, hearing, vision, over-the-counter and healthy-food benefits**, while UnitedHealthcare led **insulin cost and meal benefits**.

The result is an evidence landscape in which concentration can vary substantially from one fact to another, even within the same answer engine.

Rank Alone Does Not Describe the Competitive Environment

Consider two hypothetical factual categories. In the first, the leading publisher receives 32% of attribution while several competitors divide the remaining 68%. In the second, the leading publisher receives 80%.

Both publishers rank first.

But the underlying information environments are very different.

The first represents a relatively competitive evidence market in which multiple publishers regularly contribute to the answer engine's understanding of the fact. The second represents a highly concentrated environment in which one publisher supplies most of the observed evidence.

For this reason, the September baseline reports publisher rank and attribution share as separate measurements rather than combining them into a single visibility score.

A composite score could make the results easier to summarize, but it would also conceal the behavior the measurement is intended to expose: **which publisher supplies evidence for which fact, across how many entities, and with what degree of concentration.**

In assertion-level measurement, being first matters. **How much of the evidence environment that first-place position represents matters just as much.**

Finding 6: Search Visibility and Answer Visibility Are Related, but They Are Not Interchangeable

Conventional search discovery and AI assertion evidence are strongly related in the September baseline, particularly for publishers with large Plan-ID footprints.

But they are not the same measurement.

For Plan-IDs where Google AI used Medicare.org as assertion evidence, Medicare.org was also observed in Google's conventional search environment approximately **99% of the time**. Microsoft Copilot showed similarly high correspondence between its use of Medicare.org as assertion evidence and discovery in Bing.

Other major publishers showed the same general pattern. Publishers repeatedly used as evidence by an answer engine were frequently discoverable in that engine's corresponding search environment.

This establishes an important relationship between search discovery and answer visibility. It does not establish that conventional search discovery is a prerequisite for assertion-level evidence.

Answer Engines Use Sources Beyond the Search Results We Observe

The baseline identified numerous cases in which a publisher supplied evidence for an AI assertion even though the same publisher was not observed in the corresponding conventional search measurement for that Plan-ID.

On Google AI, for example, HealthPlanRadar supplied assertion evidence for 626 measured Plan-IDs. In 119 of those cases, the publisher was not observed in the corresponding Google search discovery measurement. Similar patterns, at smaller scales, appeared across additional publishers.

This means the observed relationship cannot be reduced to a simple sequence in which a publisher must first appear in the measured search results before an answer engine can use it as evidence.

The measurement does not establish how the answer engine discovered or retrieved sources that were absent from the corresponding search observation. It establishes only that **the evidence relationship existed even when conventional search discovery was not observed**.

One Publisher Can Have Very Different Evidence Footprints Across Engines

MedicarePlans.com provides a particularly useful example because the same publisher, content model and Plan-ID entity class produced substantially different evidence footprints across the two answer engines.

Publisher	Google AI Evidence Plan-IDs	Microsoft Copilot Evidence Plan-IDs
MedicarePlans.com	114	569

Copilot therefore used MedicarePlans.com as assertion evidence across approximately five times as many measured Plan-ID entities as Google AI did.

The difference is important because the underlying publisher and entity class did not change between the two measurements. What changed was the answer engine.

This provides another indication that visibility should not be treated as a property a webpage or publisher simply possesses. A resource can have one discovery and evidence profile on Google and a substantially different profile on Microsoft.

Search Visibility Is Not Answer Visibility

This distinction also appeared in earlier Medicare Visibility Monitor testing reported by *TheStreet*. MedicarePlanLookup.com was nearly absent from the conventional Google and Bing results in that measurement while ChatGPT cited it across 1,394 plan entities. MedicarePlans.com showed a different engine-specific pattern, receiving 1,051 ChatGPT citations compared with 114 from Google's AI.

“I therefore no longer treat search visibility and answer visibility as interchangeable measurements.”

The September assertion-level baseline provides a larger evidence framework underneath that observation.

Search discovery measures whether a resource or publisher can be observed in a conventional search environment. Answer visibility measures whether information from that publisher participates in the construction of an answer.

The two can be strongly associated without being identical.

That distinction has practical consequences for measurement. Search position alone cannot establish whether a publisher supplies evidence to an AI answer. Conversely, an AI citation alone does not explain whether the source was prominent in conventional search, how the engine encountered it, or which individual assertions it supported.

Search discovery, retrieval, assertion production and source attribution should therefore be treated as related but distinct stages of the information environment rather than different names for the same form of visibility.

Finding 7: Q1Medicare's Change in Publishing Status Creates a Longitudinal Observation Opportunity

The September baseline captured the Medicare information ecosystem immediately before a significant change in its publisher landscape.

Q1Medicare was one of the largest sources of assertion-level evidence observed in Google AI. Across the baseline population, Google associated Q1Medicare with approximately **20,900 attributed assertions across 1,833 Plan-ID entities**.

Its evidence footprint extended across **all 39 plan-information categories** measured in the study.

Q1Medicare Google AI Baseline	Measured Footprint
Plan-ID Entities	1,833
Attributed Assertions	~20,900
Fact Categories Supplied	39 of 39
Fact Categories Led	8 of 39

Q1Medicare was also the leading Google AI source for eight measured fact categories, including **dual eligibility, Medicaid eligibility, prescription drug deductible, target conditions, insulin cost, existing-member phone, new-member phone, and official website**.

In several of those categories, the difference between Q1Medicare and the next-largest source was relatively small. For prescription drug deductibles, for example, the baseline recorded approximately **1,462 Q1Medicare attributions and 1,423 Medicare.org attributions**.

A Publishing-Status Change Immediately Following the Baseline

Immediately following the September baseline measurement, Q1Medicare changed its active publishing status. This report treats that change as an observable longitudinal event; it does not assume that historical Q1Medicare pages, cached versions, indexed resources, or previously established source relationships disappeared at the same time.

That event does not mean Q1Medicare's existing answer-engine footprint will disappear immediately.

Existing resources may remain indexed, retrieved and cited for an unknown period. Search engines and answer engines may continue to associate those resources with Medicare plan assertions even though the publisher is no longer actively maintaining the publishing operation.

For that reason, this report makes **no assumption about whether, when or how quickly Q1Medicare's attribution footprint will decline**.

It also makes no assumption that attribution currently associated with Q1Medicare will transfer to any particular publisher. Evidence relationships could move to other private publishers, Medicare carriers, Medicare.gov or other authoritative sources. Some assertions could also become less frequently attributed or disappear from generated answers altogether.

A Natural Longitudinal Measurement

The timing does, however, create an unusual opportunity to observe how an answer-engine information ecosystem changes when a major evidence supplier stops active publication.

Because the September baseline records Q1Medicare's footprint at the level of the **Plan-ID, fact, engine and attributed source**, future captures can measure considerably more than whether the domain receives fewer citations.

Subsequent measurements can determine:

- Whether Q1Medicare continues to supply evidence across the same number of Plan-ID entities.
- Which factual categories retain or lose Q1Medicare attribution.
- Whether Q1Medicare's share of attribution changes within individual fact categories.
- Which publishers or authoritative sources gain evidence relationships previously associated with Q1Medicare.
- Whether some assertion-level evidence disappears rather than transferring to another source.

This makes Q1Medicare's September footprint an important part of the baseline rather than simply a historical publisher ranking.

Future editions of the Medicare Visibility Report will track that footprint without presuming the outcome. If the evidence environment changes, the measurement system can identify **where it changed, which facts were affected, which Plan-ID entities were involved, and which sources subsequently supplied the evidence**.

In that sense, Q1Medicare's publishing-status change creates a real-world longitudinal test of one of the central questions raised by this report: what happens to assertion-level knowledge when a major publisher stops actively maintaining the information ecosystem it previously supplied?

What This Report Does Not Measure

The September 2026 baseline measures observable relationships among Medicare Plan-ID entities, answer-engine assertions, source attribution, publishers and conventional search discovery.

Those measurements should not be interpreted as evidence for claims the study was not designed to test.

In particular, this baseline does **not** establish causality, factual accuracy, consumer exposure, referral traffic, commercial influence or semantic agreement between answer engines.

Causality

The measurement can identify changes in search discovery, assertion rates, source attribution and publisher attribution share. It cannot establish that a particular publishing technique, markup change, page structure or other intervention **caused** an observed ranking, retrieval or attribution outcome.

For example, observing increased attribution after a page modification establishes a temporal relationship that can be investigated and tested further. It does not, by itself, establish that the modification caused the increase.

Factual Accuracy

Source attribution does not establish that an assertion is correct.

An answer engine can attribute an incorrect, outdated or contextually inappropriate assertion to a real source. Likewise, the presence of multiple sources does not establish factual consensus or correctness.

This baseline measures **what the engines asserted and which sources they associated with those assertions**. It does not score the factual accuracy of every assertion against authoritative CMS plan data.

Consumer Exposure

An observed answer-engine response does not establish how frequently consumers encounter the same response.

Generated answers can vary by query wording, engine behavior, timing, location, personalization and other factors outside the scope of this baseline. The presence of an assertion or source in the measurement therefore should not be interpreted as an estimate of consumer impressions or audience reach.

Referral Traffic

Answer-engine attribution is not equivalent to website traffic.

A publisher may supply evidence used prominently in a generated answer without receiving a click. Conversely, a publisher receiving referral traffic may not be a major supplier of assertion-level evidence.

This report measures participation in the information used to construct answers, not visits generated by those answers.

Commercial Influence

A publisher's frequency of attribution does not establish that its ownership, advertising relationships, lead-generation activities, enrollment relationships or other business interests caused an answer engine to select it as a source.

Commercial relationships may be relevant context when evaluating a Medicare information source, but this measurement does not test whether those relationships influence retrieval, ranking, citation or source selection.

Semantic Agreement Between Answer Engines

The measurement system captured raw assertion values from both Google AI and Microsoft Copilot. However, **cross-engine agreement results are excluded from this baseline report because assertion values require fact-specific normalization before meaningful comparison.**

This limitation became apparent during analysis of the initial capture.

Two engines can express the same underlying fact differently. One might return a single copay amount while another describes multiple forms of cost sharing within the same benefit. Dollar values, percentages, eligibility descriptions, benefit limits, URLs, plan types and compound benefit descriptions each require different normalization rules before equality or disagreement can be evaluated reliably.

A simple comparison of the raw returned values would therefore risk classifying differences in wording or specificity as factual disagreement.

Rather than report a metric that could not yet be interpreted reliably, **cross-engine semantic agreement has been withheld from the September baseline.** It can be introduced in a future report after fact-specific normalization and validation are implemented.

Measurement Before Interpretation

These limitations are intentional boundaries around what the September baseline can support.

The purpose of the report is not to assign a generalized AI visibility score, declare one publisher more trustworthy than another, or infer why an answer engine selected a particular source.

Its purpose is narrower and measurable:

to observe what Medicare plan facts answer engines assert, which assertions they attribute, which sources supply that evidence, and how those relationships vary across engines, facts, entities and publishers.

Where the measurement cannot support a conclusion, this report does not make one.

Methodology

The September 2026 Medicare Visibility Report was designed to measure answer-engine behavior at a level below the webpage or citation.

The primary unit of analysis is the relationship among a **Medicare Plan-ID entity, an answer engine, and an individual factual category**.

This allows the study to distinguish between several events that are often combined under the general label of AI visibility: discovery of a resource, production of an assertion, attribution of that assertion to evidence, and selection of a particular publisher as the source of that evidence.

Entity Population

The baseline includes approximately **2,100 CMS Medicare Plan-ID entities**.

A Plan-ID represents an individual Medicare plan entity rather than a general topic, carrier or geographic market. Using the Plan-ID as the entity boundary allows observations about different publishers and answer engines to be associated with the same underlying Medicare plan.

The measured population produced **2,141 Google AI Plan-ID observations** and **2,146 Microsoft Copilot Plan-ID observations**.

Plan Knowledge Model

Each Plan-ID was evaluated against **39 predefined Medicare plan fact categories**.

The knowledge model includes information describing plan identity, classification, geography, costs, eligibility, prescription drug coverage, medical cost sharing, supplemental benefits, ratings and member contact information.

The same fact model was applied across the measured Plan-ID population so that assertion behavior could be compared consistently across plans and engines.

Answer Engines

Assertion-level analysis in this baseline focuses on the Google AI and Microsoft Copilot surfaces observed by the Medicare Visibility Monitor during the September 2026 capture. These labels identify the measured surfaces in this dataset and should not be generalized to all Google or Microsoft AI products, interfaces, or configurations.

Each engine received a standardized Plan-ID-specific measurement protocol designed to elicit the same predefined plan-information model. Results therefore describe observed behavior under that protocol, not all possible user query formulations.

Additional search and answer-engine observations were collected where available and are used in portions of the report addressing discovery and the distinction between conventional search visibility and answer visibility.

Results from one engine should not be assumed to describe another. As the findings demonstrate, engines can differ substantially in assertion completeness, attribution behavior, source selection and publisher concentration.

Primary Unit of Analysis

The principal observation can be represented as:

Plan-ID × Engine × Fact

For each combination, the system records whether the engine asserted the fact and, when source information is available, which evidence the engine associated with that assertion.

The resulting evidence relationship can be represented as:

Plan-ID → Assertion → Engine → Attributed Source

This structure allows aggregate measurements to be traced back to individual Plan-ID assertions and their associated evidence.

Core Measurements

Assertion Rate

The percentage of measured plans for which an answer engine asserted a particular fact.

$$\frac{\text{Plans where Engine asserted Fact X}}{\text{Eligible measured plans}}$$

Assertion rate measures how consistently a particular type of information appears in the engine's representation of Medicare plans.

Knowledge Completeness

The percentage of the 39-fact knowledge model asserted by an engine for an individual Plan-ID.

$$\frac{\text{Assertions made for Plan-ID}}{39 \text{ expected facts}}$$

Knowledge completeness measures the breadth of the engine's observable representation of the plan. It does not establish whether those assertions are accurate.

Attribution Rate

The percentage of asserted facts for which the answer engine exposed identifiable source attribution.

$$\frac{\text{Attributed assertions}}{\text{Assertions made}}$$

This measurement distinguishes the amount of information an engine states from the amount of stated information for which it exposes provenance.

Publisher Fact Control

Publisher fact control compares source attribution within a particular factual category to identify which publishers are most frequently associated with evidence for that fact.

This is a relative measurement. A publisher can lead a fact category without supplying the majority of all attribution for that category.

Share of Attribution

Share of attribution measures a publisher's portion of all observed attribution relationships for a particular fact and engine.

$$\frac{\text{Publisher attribution relationships for Fact X}}{\text{All observed attribution relationships for Fact X and Engine}}$$

This measurement provides the concentration dimension missing from publisher rank alone.

Publisher Identity Normalization

Publisher/source analysis is based on the observed source-domain identities retained in the evidence package. Results are reported by the publisher labels used in the measurement dataset. Corporate ownership was not used to merge otherwise distinct publisher properties unless explicitly stated. This preserves the observed evidence relationship rather than converting it into an ownership-level market-share measure.

Assertions and Attribution Relationships Are Not Necessarily One-to-One

An important feature of the measurement is that **one assertion can be associated with multiple sources**.

If an answer engine associates three sources with a single assertion, the system records one assertion and three source relationships. As a result, the number of attribution relationships can exceed the number of attributed assertions.

Conversely, an engine can make an assertion without exposing an identifiable source. That assertion contributes to knowledge completeness and assertion-rate measurements but not to publisher attribution measurements.

For this reason, assertion counts, citation counts, source relationships and publisher attribution counts should not be interpreted as interchangeable quantities.

Aggregation and Auditability

The aggregate findings in this report are derived from lower-level Plan-ID observations rather than from publisher-level visibility estimates.

The measurement retains assertion-level provenance so that aggregate results can be traced back to the underlying **Plan-ID, engine, fact, assertion and attributed source**.

This structure is intended to make subsequent captures directly comparable and to support investigation of changes at both the ecosystem level and the individual-plan level.

Worked Example: How an Assertion Becomes an Evidence Relationship

The following schematic illustrates the measurement grain without representing a specific plan observation. It shows why assertion counts and attribution relationships are not interchangeable.

Plan-ID	Engine	Fact	Asserted?	Attribution?	Source relationship
Example	Google AI	Monthly premium	Yes	Yes	Publisher A
Example	Google AI	Dental benefit	No	—	—
Example	Copilot	Monthly premium	Yes	No	—
Example	Copilot	Dental benefit	Yes	Yes	Publisher B; Publisher C

In the final row, one asserted fact produces two source relationships. In the third data row, the engine asserted a fact but exposed no attribution. Both conditions are retained separately in the measurement.

What Publishers, Carriers and Regulators Can Learn From the Baseline

The September baseline does not prescribe a single strategy for improving visibility in AI-generated Medicare answers. It does identify several characteristics of the information environment that publishers, Medicare carriers and public agencies can measure more directly than was previously possible.

For Medicare Publishers

Publishing a webpage is not, by itself, a sufficient measurement target for answer-engine visibility.

Individual facts on the same page can exhibit substantially different discovery, assertion and attribution behavior. An answer engine may repeatedly use a publisher for premiums and deductibles while relying on another source for eligibility requirements or supplemental benefits.

Publishers can therefore evaluate visibility at the level of the information they are attempting to communicate: **which entities are understood, which facts are asserted, which facts are omitted, and which assertions are associated with their evidence.**

This creates a more specific remediation target than attempting to increase generalized AI citations.

For Medicare Carriers

A Medicare carrier is an authoritative source of information about its own products, but that does not necessarily mean an answer engine will select the carrier as its evidence source for every fact about those products.

The September baseline shows that third-party publishers frequently supply assertion-level evidence across large populations of Medicare Plan-IDs.

Plan-ID-level measurement allows a carrier to identify where its own information is being discovered and used, where another publisher supplies the evidence instead, and which factual categories have weak or inconsistent representation across answer engines.

The relevant question is therefore not only whether a carrier's website is visible. It is also **whether the carrier participates in the answer engine's representation of its own plan entities and, if so, for which facts.**

For CMS and Regulators

The baseline shows that private publishers currently play a substantial role in the evidence environment surrounding AI-generated Medicare plan answers.

CMS is a primary official source for Medicare plan data and public plan-comparison information, while carriers provide plan-specific materials and operational details. The appropriate verification source depends on the fact, plan year, and enrollment context. Answer engines nevertheless frequently associate individual Medicare assertions with private publishers that organize, interpret, or republish this information.

This raises an important infrastructure question as AI systems become a more common intermediary between public Medicare data and beneficiaries:

How should authoritative Medicare information be structured and distributed so that machines can reliably resolve the correct plan, plan year, geography, eligibility requirements, costs and benefits without depending unnecessarily on intermediary interpretations?

The September baseline does not answer that policy question. It makes part of the underlying information flow measurable.

For Researchers and Measurement Providers

The results suggest that conventional search ranking, resource discovery, assertion production, source attribution and publisher attribution should be measured as **related but distinct phenomena**.

Collapsing those events into a single AI visibility score may obscure important differences between engines and publishers.

A publisher can be highly visible in search but weakly represented in generated answers. It can be frequently cited while supplying only a narrow class of facts. It can lead a factual category with a relatively small attribution share, or supply evidence to an answer engine without appearing in the conventional search observations being measured.

Those distinctions become visible when measurement moves from the page to the assertion.

For Medicare Consumers

For beneficiaries, the findings support a simpler conclusion: an AI-generated Medicare answer should be treated as a starting point for research rather than the final authority for an enrollment decision.

Important information can be omitted, different engines can provide different levels of plan detail, and the presence of a citation does not establish that every statement in an answer has been independently verified.

Before enrolling, consumers should verify plan year, service area, eligibility, costs, prescription drug coverage, provider participation and important benefits using Medicare's official resources and information supplied directly by the plan.

About This Report

The Medicare Visibility Report is an independent measurement project of the Trust Publishing Institute (TrustPublishing.org) and MedicarePlans.com. Research, measurement methodology, analysis, and reporting were conducted by David W. Bynon.

Bynon manages publishing technology and editorial systems for Medicare.org, one of the publishers measured in this report. Medicare.org is owned and operated by Health Network Group LLC, an Allstate company. Neither Medicare.org nor Allstate commissioned this report or determined its methodology, findings, or conclusions.

Publisher results are reported from the measured data without adjustment for ownership, commercial relationships, or competitive position.

© 2026 Trust Publishing Institute • Medicare Visibility Report • Research and analysis by David W. Bynon

Appendix A: Medicare Plan Knowledge Model

The September baseline evaluates each Plan-ID against the following 39 predefined fact categories. These public definitions describe the assertion categories used for coverage measurement. Fact-specific value-normalization rules are maintained separately and will be expanded before semantic-agreement analysis is published.

Fact category	Public definition
Plan ID	CMS plan identifier for the measured entity.
Plan Name	Published name of the Medicare plan.
Plan Year	Coverage year to which the assertion applies.
Carrier	Organization offering the plan.
Plan Class	High-level Medicare plan classification.
SNP Type	Special Needs Plan subtype when applicable.
Network Type	Plan network design, such as HMO, PPO, or HMO-POS.
Plan Segment	Segment identifier associated with the Plan-ID.
Service Area	Geographic area in which the plan is offered.
Monthly Premium	Plan monthly premium amount.
Part B Giveback	Plan-provided reduction or giveback associated with the Part B premium when offered.
Medical Deductible	Medical deductible amount or applicable deductible structure.
Maximum Out-of-Pocket	Maximum out-of-pocket limit for covered medical services.
Star Rating	CMS Star Rating associated with the plan.
Prescription Drug Coverage	Whether/how the plan includes Part D prescription drug coverage.
Prescription Drug Deductible	Applicable Part D deductible amount or structure.
Drug Benefit Type	Classification of the plan's prescription drug benefit.
Medicare Eligibility	Medicare eligibility condition stated for enrollment.
Medicaid Eligibility	Medicaid eligibility condition when relevant.
Dual Eligibility	Whether dual Medicare-Medicaid eligibility is required or relevant.
Target Conditions	Specified chronic or qualifying conditions targeted by the plan when applicable.
Residency Requirement	Geographic or institutional residency condition for enrollment.
Cost Sharing	General plan cost-sharing structure when explicitly asserted.
Out-of-Network Coverage	Whether and under what conditions out-of-network services are covered.
Primary Care Copay	Member cost sharing for primary care visits.
Specialist Copay	Member cost sharing for specialist visits.
Urgent Care Copay	Member cost sharing for urgent care.
Emergency Room Copay	Member cost sharing for emergency-room services.
Ambulance Copay	Member cost sharing for ambulance services; may include multiple transport types.
Insulin Cost	Plan-reported insulin cost or cost-sharing condition.
OTC Benefit	Over-the-counter benefit availability and/or value when asserted.
Healthy Food Benefit	Healthy-food benefit availability and/or value when asserted.
Meal Benefit	Meal benefit availability and/or conditions when asserted.
Dental Benefits	Dental coverage, cost sharing, allowance, limits, or conditions explicitly asserted.
Vision Benefits	Vision coverage, cost sharing, allowance, limits, or conditions explicitly

Fact category	Public definition
	asserted.
Hearing Benefits	Hearing coverage, cost sharing, allowance, limits, or conditions explicitly asserted.
New Member Phone	Telephone contact intended for prospective/new members.
Existing Member Phone	Telephone contact intended for current members.
Official Website	Website identified as the official plan or carrier resource.

Conclusion: The Assertion Is the Measurement Unit

The September 2026 baseline suggests that the emerging Medicare answer ecosystem is not simply a new version of search.

Search engines discover documents. Answer engines assemble assertions.

Those assertions have sources. The sources differ by engine, entity and fact.

Measuring that system therefore requires moving below the level of rankings, pages and citation counts to the level at which answers are actually constructed: the assertion.

Future editions of the Medicare Visibility Report will measure how those assertion relationships change over time.

Research note

This baseline reports observed answer-engine behavior. It does not assign a generalized AI visibility score, infer source-selection motives, or treat attribution as proof of accuracy. Cross-engine semantic agreement is intentionally withheld pending fact-specific normalization.