



CONVENTION · BA-C-8

Version 1.0 · Stable

# A query-intent taxonomy for visibility measurement

---

Answer engines respond differently to different kinds of question, and an entity's visibility on one kind does not carry to another.

|           |                                                |
|-----------|------------------------------------------------|
| DOCUMENT  | BA-C-8                                         |
| VERSION   | 1.0 (Stable)                                   |
| EFFECTIVE | 2026                                           |
| SUBJECTS  | Measurement & statistics; Conventions & policy |
| LICENSE   | CC-BY-4.0                                      |
| AUTHOR    | Barkhausen AI                                  |

**ABSTRACT**

Answer engines respond differently to different kinds of question, and an entity's visibility on one kind does not carry to another. This convention classifies the query space a credible visibility study must cover into five intent classes — navigational, informational, comparative, constraint-based, and recommendation-seeking — defined by the goal a user's information need expresses, and states what a study must disclose about its coverage. Each class behaves differently under measurement: navigational visibility is near-degenerate, informational mentions are incidental, comparative needs are Share of Voice's natural home, constraint-based needs are where Discovery Depth applies, and recommendation-seeking needs yield the least stable and most commercially consequential recommendation sets. A study must disclose which classes it covers and in what proportion, must not generalize a class-specific result across classes, and should report per-class estimates where classes are mixed — all at the level of classes, never the concrete queries.

**CONTENTS**

|   |                                          |   |
|---|------------------------------------------|---|
| 1 | 1. Scope and the level of classification | 3 |
| 2 | 2. Five intent classes                   | 3 |
| 3 | 3. Coverage and disclosure requirements  | 6 |
| 4 | 4. Boundary cases                        | 6 |
| 5 | 5. Conformance                           | 7 |
| 6 | 6. Limitations                           | 7 |

A visibility study measures an entity's presence in answers to some set of questions, and the questions are not interchangeable. A question that names a single known entity, a question that asks for background on a topic, and a question that asks which option to choose exercise an engine in different ways and make an entity mention mean different things. A study that measures visibility on one kind of question and reports it as visibility in general has drawn a conclusion its evidence does not support. This convention classifies the query space a credible visibility study must reckon with, and states what a study must disclose about which parts of that space it covers.

The classification is analytic: it partitions questions by the kind of goal a user brings to them, so that a study can state its coverage and a reader can judge it. It is concept-level throughout. It says nothing about which concrete questions a study should ask — the construction of a query set is a study's own matter, and the specific questions are never the subject of this convention — only about the classes such a set must be described in terms of.

**Normative note.** The key words **MUST**, **MUST NOT**, **SHOULD**, and **MAY** carry their established normative meanings. They govern how a study describes and discloses the intent coverage of its query set; they do not prescribe which questions a study asks. Every requirement is written so that conformance can be checked from a study's disclosure alone, at the level of classes and proportions, without any concrete query being published.

## 1 1. Scope and the level of classification

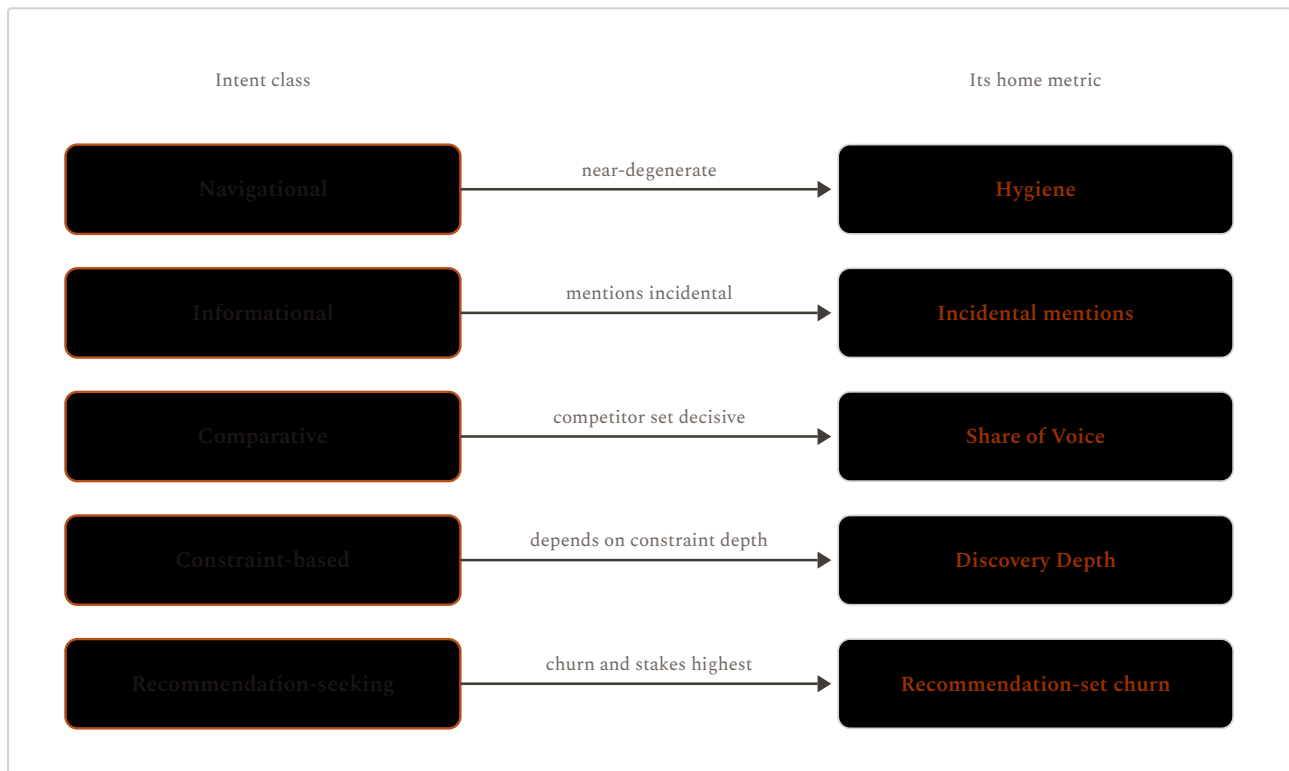
This convention applies to any study that estimates entity visibility in answer engines and reports a result about a set of questions. It classifies **information needs** — what a user is trying to find out (BA-C-7) — not surface phrasings, and not the answers returned. Intent is a property of the need, and the same need keeps its class across the many phrasings that express it; a study therefore assigns a class per information need, once, and not per phrasing or per run.

Classification is disclosed at the level of the class, never the concrete query. A study states which classes its query set covers and in what proportion; it does not, and under this convention need not, publish the questions themselves. This mirrors the phrasing-disclosure rule of BA-C-3 and disclosure items 1 and 8 of BA-C-4 §2, which together require a study to characterize its queries — what they ask (item 1), and how many phrasings with what coverage (item 8) — without exposing the specific wordings. Coverage is a property a reader can assess from counts and classes; it does not require the query list, and this convention does not ask for it.

## 2 2. Five intent classes

This convention defines five classes, partitioning information needs by the kind of goal each expresses. The table names them; the subsections give the normative definition of each and the reason each behaves differently under measurement. The examples throughout are deliberately generic.

| CLASS                  | THE USER IS TRYING TO                            | WHY IT BEHAVES DIFFERENTLY FOR MEASUREMENT                    |
|------------------------|--------------------------------------------------|---------------------------------------------------------------|
| Navigational           | reach one specific, already-known entity         | Visibility is near-degenerate; mostly hygiene-sensitive       |
| Informational          | learn about a topic, with no entity in mind      | Entity mentions are incidental; Share of Voice is noisy       |
| Comparative            | weigh a set of entities against one another      | Share of Voice's natural home; the competitor set is decisive |
| Constraint-based       | satisfy a need narrowed by successive conditions | Discovery Depth's home; result depends on constraint depth    |
| Recommendation-seeking | be told which entity or entities to choose       | Recommendation-set churn is highest; stakes are highest       |



**Figure 1.** Each of the five intent classes paired with the visibility measure that is its natural home: navigational needs with hygiene, where Visibility Probability saturates near one and carries little information; informational needs with incidental mentions; comparative needs with Share of Voice; constraint-based needs with Discovery Depth; and recommendation-seeking needs with recommendation-set churn, the least stable and highest-stakes of the five. A class-specific result does not carry across classes (§3).

Barkhausen AI · BA-C-8

## 2.1 Navigational

A **navigational** need seeks a specific entity the user already has in mind — reaching a particular known organization, product, or page, not discovering a new one. The textbook example is a user who asks for a named national weather service by name. For measurement, the target entity's visibility on its own navigational queries is near-degenerate: the entity is almost always present, so Visibility Probability sits close to one and carries little information, and what variation exists is mostly a matter of hygiene — whether the engine resolves the name to the right entity at all (the lower rungs of BA-C-1). Navigational

queries are poor instruments for competitive visibility, because a competitor is rarely a candidate answer to a request for a named third party. A study that measures an entity mostly on its own navigational queries will report a high, stable, and largely uninformative number.

## 2.2 Informational

An **informational** need seeks knowledge about a topic, with no particular entity in view — how something works, what a term means, why a phenomenon occurs. The textbook example is “how does a thunderstorm form.” Entity mentions in answers to informational needs are incidental: an answer may name an entity as the source of a fact or as an example, but naming any particular one is not the point of the answer, and small changes in phrasing or retrieval can add or drop such mentions. Share of Voice computed over informational answers is therefore noisy and easily over-read — an entity that appears in an informational answer has been cited as a source or an illustration, which is not the same as being recommended. Informational coverage measures a real and distinct thing, whether an entity is drawn on as a source; a study **MUST NOT** let a strong informational result stand in for a recommendation result.

## 2.3 Comparative

A **comparative** need weighs a set of entities against one another — the user has, or the question implies, two or more candidates and wants them compared. The textbook example is “compare two named laptop models.” This is the natural home of Share of Voice: the answer’s business is to allocate attention across a set of entities, so an entity’s share of the mentions in comparative answers is a meaningful competitive quantity rather than an incidental count. The measurement hazard specific to this class is the competitor set. A comparative result is defined only relative to the set of entities in play, and the same entity’s share can be made to look strong or weak by the choice of comparators; a study reporting comparative visibility **MUST** disclose the competitor set it measured against, as BA-C-2 §3 already requires for Share of Voice, because a share without its denominator set is uninterpretable.

## 2.4 Constraint-based

A **constraint-based** need is one progressively narrowed by qualifying conditions — the user adds requirements until the field of acceptable answers shrinks to those that satisfy all of them. The textbook example is a request for a camera that is then qualified by price, then by weight, then by lens compatibility. This is the home of Discovery Depth (BA-C-2 §4): the quantity of interest is the constraint depth at which an entity first enters the answer’s set of candidates. Constraint-based needs behave differently because each added constraint changes the candidate set, often sharply, and an entity visible under a loose need can vanish under a tighter one, or appear only once the need is narrow enough to select it. A study that measures at one constraint depth has measured one point on this curve; it **MUST NOT** report that point as the entity’s visibility for the broader family of needs, because the curve is exactly what constraint-based measurement exists to expose.

## 2.5 Recommendation-seeking

A **recommendation-seeking** need is an open request to be told which entity or entities to choose — “which should I get,” “who should I use for this” — with no candidate set supplied by the user. The textbook example is “which coffee maker should I buy.” This is the commercially load-bearing class, because its answer is a recommendation set (BA-C-7) that directly shapes a choice, and it is also the least stable. The public evidence is specific: a 2026 study of production retrieval-augmented recommendation systems, running roughly 12,000 queries across deployed OpenAI and Anthropic models, found that

rewording a request while adding a qualifying constraint cut the overlap between the two answers' recommendation sets to a Jaccard similarity of about 0.14, far below the overlap of about 0.50 to 0.61 that simply repeating an identical request reproduced [1]. Recommendation-set membership is thus both the most consequential visibility outcome and the most sensitive to phrasing, which makes single-observation claims about it the least defensible of any class and the sampling requirements of BA-C-3 most acute here.

### 3 3. Coverage and disclosure requirements

The classes exist to be disclosed. A study's choice of which classes to cover, and in what proportion, determines what its visibility figure means, and a reader can judge the figure only if that choice is stated.

**Disclose the coverage.** A study **MUST** disclose which of the five classes its query set covers and in what approximate proportion. The disclosure is at the level of classes and shares — for instance, that a query set is predominantly recommendation-seeking with a minority of comparative needs — and it does not require, and this convention does not ask for, the concrete queries behind those shares (BA-C-3; BA-C-4 §2 item 8).

**Do not generalize across classes.** A study **MUST NOT** report a result obtained on one class as a result about another, or about visibility in general. “Visible for informational needs” is not “recommended”; a strong Share of Voice on comparative needs is not presence on the open recommendation-seeking needs that most directly drive a choice. Where a single headline figure is unavoidable, the study **MUST** state which class or mix of classes it was computed on.

**Report per class where classes are mixed.** Where a query set spans more than one class, a study **SHOULD** report its estimates per class rather than only in aggregate, because a pooled figure over an unstated mixture of classes is a weighted average whose weights — the class proportions — are themselves a study-design choice a reader cannot see. Per-class estimates let a reader reweight to the mixture they care about; a single pooled number does not.

**Classify at the level of the need.** Classification is assigned per information need, not per phrasing and not per answer, and is disclosed at the class level (BA-C-7; BA-C-3). Two phrasings of one need share its class; the class is a property of the goal, not of the words or of what the engine happened to return.

### 4 4. Boundary cases

Three cases sit at the edges of the scheme and are handled explicitly.

**Multi-intent needs.** A single need can carry more than one intent — a question that asks both what the options are and how to choose among them is comparative and recommendation-seeking at once. A study **SHOULD** assign such a need to the class that governs the result it reports, and disclose that the need is mixed; it **MUST NOT** silently count a multi-intent need under whichever class flatters the result. The classes partition an idealized query space; real needs occasionally straddle a boundary, and the requirement is to say so.

**Intent drift within a conversation.** A user's intent can move over the course of a multi-turn session — an informational exchange that turns into a recommendation request. This convention classifies single, self-contained needs and does not model within-conversation drift; multi-turn intent dynamics are out of scope, consistent with the single-turn framing of the measurement conventions, and are noted here only so that the scope boundary is explicit rather than assumed.

**Locale-dependence of the intent mix.** The mix of intents a population brings to a topic is not universal; it varies by region, language, and market. This convention defines the classes, which are stable, and not the mixture, which is not. A study that measures in one locale **MUST NOT** present its class coverage as representative of another, because the proportion of, say, recommendation-seeking to informational needs is itself locale-dependent — a property of the population, not of the taxonomy.

## 5 5. Conformance

A study conforms to this convention when it discloses the intent-class coverage of its query set at the class-and-proportion level, refrains from generalizing a class-specific result beyond its class, and reports per-class estimates where its query set mixes classes. Conformance is a property of the study’s disclosure and reasoning, checkable without any concrete query being published, and it concerns coverage and interpretation only — it does not certify the validity of the underlying estimates, which is governed by BA-C-2 and BA-C-3. A study **MAY** state conformance as: *query-intent coverage disclosed per BA-C-8 v1.0*. Partial conformance **MUST** be stated as named exceptions rather than implied by silence.

## 6 Limitations

This taxonomy is analytic, not empirical. It partitions the query space by the goal a need expresses; it makes no claim about how real query traffic divides among the classes, on any engine or in any market. Any statement about the real-world share of, for instance, recommendation-seeking needs would require traffic data this convention does not have and does not assert. The classes are a lens for describing coverage, not a measurement of it.

The partition is also a model, and needs do not always fall cleanly into one class. The boundary cases of §4 are the routine exceptions; there will be others. The five classes are chosen because each behaves distinctly enough under measurement to be worth separating, not because the space of goals has exactly five natural kinds. As answer engines and the ways people query them change, a class may prove to need splitting — recommendation-seeking, in particular, spans a wide range of stakes and answer formats — and such a change will arrive through a new version of this convention rather than through silent reinterpretation of the classes defined here.

Finally, a taxonomy of intent does not by itself make a query set representative. Covering all five classes says nothing about whether the needs chosen within each class are the ones that matter to a given question; representativeness within a class is a separate matter, governed by the query-selection and phrasing requirements of BA-C-2 and BA-C-3. This convention ensures a study can state and defend its coverage of the intent space; it does not, on its own, guarantee that coverage is complete.

## REFERENCES

1. Jack, Lehman, Maloney, Xu; arXiv:2605.27440. *Paraphrase Brittleness in Production Retrieval-Augmented Commercial Recommendation: Reproducibility Below the Rerun-Stability Baseline* (2026). <https://arxiv.org/abs/2605.27440> Accessed 2026-07-10. [archived]

---

## HOW TO CITE

Barkhausen AI (2026). A query-intent taxonomy for visibility measurement.  
<https://barkhausen.ai/conventions/query-intent-taxonomy/>

Published under the Creative Commons Attribution 4.0 International (CC-BY-4.0).



Barkhausen AI · Est. MMXXVI · *Every leap begins as a crackle.*