

RAG Signal Methodology

The 7-Dimension Signal Scoring Algorithm

RAG Signal | Bora Kurum

August 2026 | Version 1.0 | Open Source (Apache 2.0)

1. Overview

RAG Signal engineers brand visibility inside AI retrieval systems using **Adaptive RAG** — a five-phase process that continuously adjusts to model updates and competitive changes:

MAP → BUILD → WEIGHT → REINFORCE → MEASURE

This document details the **WEIGHT** phase: the Signal Scoring Engine, its seven dimensions, scoring criteria, weightings, and interpretation guidelines.

2. The Signal Scoring Engine

Every brand signal is scored across seven dimensions. The weights reflect how each dimension influenced citation outcomes in RAG Signal's 2025 controlled experiment (63 prompts × 4 LLMs, test–retest correlation 0.94).

Dimension	Weight	Scoring criterion
Source Authority	25%	Trustworthiness of sources surrounding your claims: publication quality, domain history, third-party corroboration.
Factual Consistency	20%	Claims internally consistent and corroborated across independent surfaces; no contradictions.
Entity Linkage	15%	Density and clarity of entity definitions: category, geography, product, people, proof in retrievable form.
Cross-Model Persistence	12%	Cited consistently for the same prompt across ChatGPT, Claude, Perplexity, Gemini.
Temporal Freshness	12%	Recency of content relative to query; presence of date stamps and update history.
Citation Frequency	10%	How often the brand is referenced across the open web and in model corpora.
Competitive Diff	6%	How specifically the content answers this query versus competitors.

3. Interpretation Guidelines

Dimension scores are normalized to 0–100 per prompt cluster, then combined with the weights above.

Composite score	Interpretation
0–20	Structurally invisible: entity definitions missing, no third-party references.
20–50	Retrievable but not preferred: cited occasionally, loses to competitors on specificity.
50–75	Strong: consistently cited, some model gaps.
75+	Dominant: cited in most prompts across all models; the reference answer.

3.1 Per-model calibration

Weights are *not* static. Retrieval behavior differs by model: in our data, Perplexity weights external citation anchors more heavily, Gemini is the most freshness-sensitive, and Claude rewards corroboration (factual consistency). We re-calibrate the weighting per model and per prompt cluster in every engagement.

4. Measurement

Visibility is tracked with three metrics:

- **Citation Rate** — percentage of buyer prompts where a model names the brand.
- **Citation Delta** — change between measurement points (30/60/90 days).
- **Cross-Model Persistence** — citation consistency across the four models.

Each prompt is run three times per model and averaged; a test–retest analysis over two weeks produced a correlation of 0.94.

5. References

- RAG Signal, *The 2025 Controlled Experiment*, <https://ragsignal.com/research/2025-controlled-experi>
- RAG Signal, *Adaptive RAG Architecture whitepaper* (Apache 2.0), <https://ragsignal.com/whitepaper/>
- RAG Signal, *Signal Weighting Explained*, <https://ragsignal.com/insights/signal-weighting-explained>
- RAG Signal, *State of AI Visibility 2026*, <https://ragsignal.com/research/state-of-ai-visibility-2026/>

Published under Apache License 2.0. Reproduce, audit, and challenge — the methodology is meant to be checked.